

UNIVERZITA KOMENSKÉHO V BRATISLAVE
RÍMSKOKATOLÍCKA CYRILOMETODSKÁ
BOHOSLOVECKÁ FAKULTA

Etické výzvy a morálne aspekty súčasných systémov umelej inteligencie

Dizertačná práca

Študijný program: katolícka teológia

(jednoodborové štúdium, doktorandské III. st., externá forma)

Študijný odbor: teológia

Katedra: RKCMBF.KST – Katedra systematickej teológie

Školiteľ: ThDr. Ing. Vladimír Thurzo, PhD.

Bratislava, 2022

ThLic. Ing. Peter Šantavý

Zadanie

Čestné vyhlásenie

Vyhlasujem, že som dizertačnú prácu vypracoval samostatne a uviedol som všetku použitú literatúru.

Pod'akovanie

Vďačný Bohu, zo srdca ďakujem môjmu školiťovi, ThDr. Ing. Vladimírovi Thurzovi, PhD., za cenné rady a pripomienky, erudovanosť, odborné vedenie, trpezlivosť i povzbudenie pri písaní dizertačnej práce.

Motto

„Umelá inteligencia“ (artificial intelligence - AI) prináša a bude čoraz viac prinášať veľké zmeny do života ľudí. Ponúka totiž mimoriadne možnosti...

...teraz teda viac ako inokedy musíme zaistiť taký prístup, v ktorom sa rozvoj AI nesústreďuje len na technológiu, ale berie ohľad na dobro ľudstva a životného prostredia, na náš spoločný domov a jeho ľudských obyvateľov, ktorí sú navzájom nerozlučne prepojení.

Výzva na etiku v oblasti umelej inteligencie

Abstrakt

Bibliografická identifikácia	ŠANTAVÝ Peter. <i>Etické výzvy a morálne aspekty súčasných systémov umelej inteligencie</i>. Univerzita Komenského v Bratislave. Rímskokatolícka cyrilometodská bohoslovecká fakulta. RKCMBF.KST – Katedra systematickej teológie. Školiteľ: ThDr. Ing. Vladimír Thurzo, PhD. Stupeň odbornej kvalifikácie: doktor
Abstrakt	V rámci rodiacej sa informačnej spoločnosti ako dôsledok technologického rozvoja a paradigmatickej zmeny vystupujú do popredia nové fenomény, ktorých opodstatnenosť a potenciál v sebe nesie silný morálny akcent. Systémy umelej inteligencie patria nielen k podstatným fenoménom informačnej a znalostnej spoločnosti, ale sa im prisudzuje nesmierny potenciál v budúcnosti. Zámerom práce je predstaviť súčasné technológie umelej inteligencie, diskutovať ich limity a riziká v kontexte etických princípov i morálnych dôsledkov a navrhnúť etické zásady pre zabezpečenie dôveryhodných a na dobro človeka orientovaných systémov umelej inteligencie.
Cieľ práce	Na základe analýzy aktuálneho stavu rozvoja systémov umelej inteligencie a existujúcich etických pravidiel snaha navrhnúť základnú štruktúru zásad, ktorá musí byť podchytená pri etickom návrhu, realizácii i využívaní akýchkoľvek systémov umelej inteligencie.
Kľúčové slová	umelá inteligencia, virtuálny svet, kybernetický priestor, informačná spoločnosť, informačné technológie, etika, morálne princípy

Abstract

Bibliographic identification	ŠANTAVÝ Peter. <i>Ethical Issues and Moral Aspects of Current Artificial Intelligence Systems</i>. Comenius University in Bratislava. Roman Catholic Faculty of Theology. Department of Systematic Theology. Supervisor: ThDr. Ing. Vladimír Thurzo, PhD. Degree of professional qualifications: philosophiae doctor
Abstract	Within the emerging information society, as a consequence of technological development and paradigmatic change, new phenomena are coming to the fore, the validity and potential of which carry a strong moral emphasis. Artificial intelligence systems are not only one of the essential phenomena of the information and knowledge society, but are seen as having enormous potential in the future. The aim of this thesis is to introduce current AI technologies, discuss their limits and risks in the context of ethical principles and moral implications, and propose ethical principles for ensuring trustworthy and human-centred AI systems.
Objective	To propose the basic principles to be upheld during ethics proposal, execution and usage of any AI systems, based on the analysis of the current state of AI systems development and existing ethical rules.
Keywords	artificial intelligence, virtual world, cyberspace, information society, information technologies, ethics, moral principles

Obsah

Zadanie.....	2
Čestné vyhlásenie.....	3
Podakovanie.....	4
Motto.....	5
Abstrakt.....	6
Abstract.....	7
Obsah.....	8
Zoznam obrázkov.....	11
Zoznam skratiek a symbolov.....	12
Predslov.....	13
Úvod.....	19
1. Základy umelej inteligencie.....	29
1.1. Úsvit umelej inteligencie.....	31
1.2. Definícia pojmov.....	35
1.3. Základné vlastnosti a delenie systémov umelej inteligencie.....	37
1.4. Metódy umelej inteligencie – ich anarchia i filozofia rozdelenia.....	39
1.5. Systémy inšpirované činnosťou mozgu na úrovni neurónov.....	42
1.6. Algoritmy strojového učenia.....	46
1.6.1. Učenie s učiteľom (supervised machine learning).....	46
1.6.2. Učenie bez učiteľa (unsupervised machine learning).....	47
1.6.3. Učenie formou odmeňovania (reinforcement learning).....	48
1.7. Stačí strojové učenie, alebo hľadáme ďalej?.....	48
1.8. Striedanie ročných období.....	53

1.9. Jednoduché veci sú ťažké.....	56
1.10. Transhumanizmus a umelá inteligencia.....	58
1.11. Masívne zavádzanie prvkov umelej inteligencie.....	60
2. Limity a riziká súčasných systémov umelej inteligencie.....	62
2.1. Vybrané limity a rizikové faktory systémov umelej inteligencie.....	63
2.1.1. Neurónová sieť ako čierna skrinka.....	64
2.1.2. Zraniteľnosti systémov strojového učenia.....	67
2.2. Bezpečnosť procesov umelej inteligencie.....	77
2.3. Kybernetická bezpečnosť.....	80
2.4. Technologická komplexnosť a potrebná infraštruktúra.....	83
2.5. Dôsledky pre spoločnosť.....	87
2.6. Umelá inteligencia prichádza do spravodajských služieb.....	99
2.7. Systémy umelej inteligencie v armáde.....	109
2.7.1. Vojenské spravodajstvo.....	111
2.7.2. Modelovanie technológií, konfliktov a operácií.....	112
2.7.3. Podpora pre velenie.....	114
2.7.4. Trenažéry, simulátory a výcvik.....	114
2.7.5. Autonómne zbraňové systémy.....	115
2.7.6. Skupinové riadenie bojových prostriedkov a autonómnych systémov.....	120
2.7.7. Vedenie vojny v kybernetickom priestore.....	121
2.7.8. Ďalšie etické a právne aspekty.....	129
2.8. Môžeme umelej inteligencii dôverovať?.....	136
3. Umelá inteligencia a etika.....	138
3.1. Aké etické výzvy prináša umelá inteligencia.....	139
3.2. Angažovanosť a aktivity na poli etiky umelej inteligencie.....	155
3.3. Legislatívne kroky v oblasti etiky umelej inteligencie.....	161

3.4. Postoj Cirkvi.....	173
4. Navrhnuté riešenie etických problémov ANI.....	181
4.1. Základný postoj vo svetle Zjavenia.....	181
4.2. Interdisciplinárny rámec.....	184
4.3. Všeobecné návrhy.....	185
4.3.1. Umelá inteligencia zameraná na človeka.....	185
4.3.2. Dôveryhodná umelá inteligencia.....	186
4.3.3. Etické požiadavky na dôveryhodné systémy umelej inteligencie.....	188
4.3.4. Oblasť implementácie etických princípov.....	191
4.4. Špecifické odporúčania.....	192
4.4.1. Oblasť pokročilého riadenia štátu, spravodajstva a plošného dohľadu.....	193
4.4.2. Systémy AI vo vojenskej oblasti.....	194
4.5. Priestor pre angažovanie sa Cirkvi.....	198
4.5.1. Misia zjednocovať, usmerňovať a propagovať etické aktivity.....	198
4.5.2. Akcent na univerzálne bratstvo a sociálne priateľstvo.....	200
5. Vízia silnej a všeobecnej umelej inteligencie.....	202
5.1. „Trhliny v inteligencii“ pokročilých systémov AI.....	203
5.2. Ontologické otázky.....	209
Záver.....	213
Prílohy.....	224
Zoznam použitej literatúry.....	227
Slovník termínov.....	249

Zoznam obrázkov

Obrázok č. 1: Vzťah medzi silnou a slabou, všeobecnou a úzko špecializovanou umelou inteligenciou	39
Obrázok č. 2: A – neurón v mozgu, B – jednoduchý perceptron	43
Obrázok č. 3: Viacvrstvová neurónová sieť	45
Obrázok č. 4: Hlboká neurónová sieť	46
Obrázok č. 5: Vzťah medzi dátovou vedou, počítačovou vedou a umelou inteligenciou	53
Obrázok č. 6: Pravdepodobnosť výskytu niektorých situácií, s ktorými sa môže autonómne vozidlo stretnúť v prevádzke	71
Obrázok č. 7: Správne a nesprávne klasifikované obrázky sieťou AlexNet	73
Obrázok č. 8: Príklady šumu, ktoré konvolučné siete vyhodnocujú ako kategórie objektov	74
Obrázok č. 9: Výpočtový výkon systémov AI sa v poslednej dekáde zvyšuje exponenciálne	84
Obrázok č. 10: Dva pohľady na singularitu v oblasti umelej inteligencie	93
Obrázok č. 11: Päť úrovní automatizácie vozidla	98

Zoznam skratiek a symbolov

AI	artificial intelligence (umelá inteligencia)
ANI	artificial narrow intelligence (narrow AI)
AGI	artificial general intelligence (general AI)
ASI	artificial super intelligence
CNN	konvolučné neurónové siete (convolutional neural networks)
ConvNets	konvolučné neurónové siete (convolutional neural networks)
GEB	zaužívaný akronym pre knihu <i>Gödel, Escher, Bach: an Eternal Golden Braid</i> , ktorú v roku 1979 napísal Douglas Hofstadter
IKT	informačné a komunikačné technológie
IS	informačná spoločnosť
LAWs	lethal autonomous weapons – smrtiace autonómne zbraňové systémy
NhRP	Non-human Rights Project

Predslov

Keď v roku 1979 Douglas Hofstadter¹ napísal knihu *Gödel, Escher, Bach: an Eternal Golden Braid*² (známu tiež pod skratkou GEB), určite nečakal, koľko súčasných vedcov a pedagógov venujúcich sa umelej inteligencii privedie k tejto problematike práve jeho kniha³.

Možno i preto bol Douglas Hofstadter v roku 2014 pozvaný do centrál Google, aby tam prednášal vybraným vedcom z oblasti umelej inteligencie. A keďže viacerí z nich na GEB vyrastali, tridsaťpäť rokov od jej napísania bol dosť dlhý čas na rekapituláciu vízií, ktorú nemohol urobiť nik povolanejší, než sám autor.

Zvedavému auditóriu miesto očakávaného nadšenia, predstavených vízií, povzbudenie, či rád do vedeckej práce legendárny Douglas Hofstadter prezentoval obavy a zdesenie.

V roku 1979 bola téma umelej inteligencie vzrušujúcou, no veľmi vzdialenou budúcnosťou; vzdialenou až do tej miery, že akákoľvek jej možná realizácia a riziká boli brané ako niečo, čo sa týka budúcich generácií. Douglas Hofstadter – narozdiel od viacerých iných odborníkov – veril v realizovateľnosť umelej inteligencie napodobňujúcej inteligenciu človeka, no v jeho podaní išlo viac o víziu, než o projekt, ktorý by sa mohol v najbližších

1 Douglas Richard Hofstadter je americký filozof, spisovateľ, vysokoškolský pedagóg, informatik a fyzik, ktorý sa zaoberá interdisciplinárnymi témami, ako napr. vedomie, preklad, tvorivosť a porovnávanie ľudského rozumu a umelej inteligencie. Je legendou v oblasti umelej inteligencie.

Douglas Hofstadter [on-line]. [cit. 30. júla 2020].

Dostupné na internete: <https://en.wikipedia.org/wiki/Douglas_Hofstadter>

2 Jednoducho povedané, GEB je súhrnom intelektuálnych záujmov Douglasa Hofstadtera (matematika, umenie, hudba, reč, humor a slovné hry), ktoré sú premietnuté do fundamentálnej otázky, ako sa inteligencia, uvedomenie a pocity sebauvedomenia, vlastné ľudským bytostiam, môžu tak zásadne vynoriť z neinteligentného, nevedomého substrátu biologických buniek. A tiež ako k inteligencii a sebauvedomeniu môžu eventuálne dospieť počítače.

HOFSTADTER, D. *Gödel, Escher, Bach: an Eternal Golden Braid*. New York: Basic Books, 1979.

3 Za všetkých môžeme spomenúť Dr. Melanie Mitchell, bývalú doktorandku a spolupracovníčku Douglasa Hofstadtera na University of Michigan, ktorá v súčasnosti pôsobí ako profesorka počítačových vied na Portland State University. O svojej ceste k AI a podobnej skúsenosti viacerých vedcov z rôznych tímov z Google AI píše vo svojej knihe *Artificial Intelligence*.

MITCHELL, M. *Artificial Intelligence*. Farrar, Straus and Giroux. 2019, s. 5-6.

desaťročiach reálne uskutočňovať.⁴

V závere GEB Hofstadter uvádza desať otázok a špekulácií o umelej inteligencii.

Jednou z nich bola i otázka: Bude existovať šachový program, ktorý by porazil *kohokoľvek*?

Hofstadter odpovedá: nie, nebude existovať (maximálne tak porazí dobrého amatéra). Na porazenie šachového majstra by muselo ísť o počítačový systém s univerzálnou inteligenciou.⁵

Hofstadter vychádzal z analýz svojho priateľa Eliota Hearsta, šachového šampióna a profesora psychológie. Hearst sa venoval rozdielom medzi ľudským šachovým expertom a počítačovým šachovým programom, tvrdiac, že zatiaľčo sa človek rozhoduje na základe rozpoznávania vzorov a abstrahovania, program koná na základe masívneho vyhľadávania budúcich ťahov hrubou výpočtovou silou⁶. Podľa Hearsta bez všeobecnej inteligencie a jej zodpovedajúcej úrovni abstrakcie konceptov nie je možné dosiahnuť úroveň šachového veľmajstra.

Osemdesiate a deväťdesiate roky minulého storočia však priniesli veľký pokrok nielen vo výpočtovom výkone počítačov, ale i v šachových programoch pracujúcich v princípe stále len *hrubou silou*. Neustále zlepšovanie vyústilo v roku 1997 do svetoznámeho duelu medzi Deep Blue machine od IBM a úradujúcim svetovým šampiónom Garry Kasparovom. Deep Blue v zápase šiestich hier Kasparova porazil.⁷

Šachové majstrovstvo, ktoré bolo kedysi považované za vrchol ľudskej inteligencie,

4 Sám sa vyjadroval, že realizácia človeku podobnej umelej inteligencie je vzdialená na *sto Nobelových cien* (citované v ROTA G.-C. *Indiscrete Thoughts*. Boston: Berkhäuser, 1997, s. 22).

5 HOFSTADTER, Gödel, Escher, Bach: an Eternal Golden Braid, s. 678.

6 Je to zjednodušene a skôr principiálne vyjadrené, keďže už od roku 1949 Arthur Samuel, inžinier v IBM a jeden z priekopníkov v oblasti umelej inteligencie (práve on šíril termín strojové učenie / machine learning), vyvíjal program pre hru dámy, v rámci ktorého boli položené viaceré základy z budúcich metód umelej inteligencie, napr. strom hry, vyhodnocovacie funkcie, učenie sa samou hrou s načrtnutými základmi pre učenie posilňovaním (reinforcement learning). Viac napríklad v publikácii: SAMUEL, A. L. *Some Studies in Machine Learning Using the Game of Checkers*. In: *IBM Journal of Research and Development*. 1959, č. 3, s. 210-229.

7 *Deep Blue versus Garry Kasparov* [on-line]. [cit. 30. júla 2020].
Dostupné na internete: <https://en.wikipedia.org/wiki/Deep_Blue_versus_Garry_Kasparov>

podľaňlo postupom využívajúcim hrubú výpočtovú silu...⁸

Pre Douglasa Hofstadtera bola hudba *baštou človečenstva* – jednoducho doménou i výkladnou skriňou ľudského ducha. Preto – po zlyhaní šachovej vízie – prichádza na rad ďalšia z desiatich otázok, ktorými kedysi Douglas Hofstadter zavŕšil GEB: Bude vedieť počítač skomponovať *nádhernú* hudbu?

Hofstadter predpovedá: áno, no nebude to čoskoro. Ďalej vysvetľuje: hudba je jazykom emócií, a pokiaľ programy nedisponujú emóciami podobne komplexnými, ako sú ľudské, nie je možné, aby program skomponoval čokoľvek nádherné, či krásne. Môžu existovať falzifikáty a povrchné imitácie vychádzajúce z poznania a syntaxe už vytvorenej hudby, no to určite nestačí na tvorivú činnosť, ktorej ovocím je krásna hudba. Pre samotného Douglasa Hofstadtera bola táto úvaha tak dôležitá, že – ako sám hovorieval – na toto

8 MITCHELL, *Artificial Intelligence*, s. 8.

Treba dodať, že šach nie je vhodným základom pre meranie schopností umelej inteligencie, či už vzhľadom na reálne vypočítateľný počet možných ťahov alebo algoritimizáciu šachových pravidiel (Deep Blue – narozdiel napr. od Samuelovho programu na hru dámy – neobsahoval žiadne pokročilejšie metódy umelej inteligencie). Oveľa náročnejšou skúškou je čínska abstraktná strategická hra GO, ktorá je oveľa zložitejšia než šach. GO umožňuje vytvoriť 2×10^{170} pozícií, čo je oveľa viac, než je atómov v známom vesmíre (10^{80}). Jednoduchá algoritimizácia a hrubá výpočtová sila tu zlyhávajú.

I napriek tomu v rokoch 2015 – 2017 umelá inteligencia AlphaGo od Googlu porazila viacerých najlepších hráčov sveta. Jej predposledná verzia AlphaGo Zero, využívajúc pokroky v AI, porazila pôvodnú AlphaGo 100:0 a najnovšia verzia AlphaZero je považovaná za najlepšieho hráča Go a pravdepodobne i šachu (narozdiel od Deep Blue však všetky verzie AlphaGo pokročilé metódy umelej inteligencie využívajú – ide o použitie známeho vyhľadávania pomocou stromu Monte-Carlo a GPT, t.j. Generative Pre-Trained Transformer).

AlphaGo [on-line]. [cit. 30. júla 2020].

Dostupné na internete: <<https://en.wikipedia.org/wiki/AlphaGo>>

O AlphaZero sa niekedy hovorí ako o prekročení slabých umelých inteligencií (o slabej a silnej umelej inteligencii budeme pojednávať v ďalšom texte).

TUROŇ, J. *Alpha Zero: soumrak slabých umělých inteligencí* [on-line]. [cit. 30. júla 2020].

Dostupné na internete: <<https://www.osel.cz/9702-alpha-zero-soumrak-slabych-umelych-inteligenci.html>>

I keď treba povedať, že sa skôr prikláňame k názoru Dr. Mitchellovej a mnohých iných odborníkov na umelú inteligenciu, tvrdiac, že ani pri takých systémoch, ako sú AlphaGo, GPT-3 a pod. sa stále nedá hovoriť o prekračovaní hraníc slabej umelej inteligencie, aj keď treba povedať, že seriózný vývoj v tejto oblasti neustále prebieha (napr. projekt Alexandria, COMET, atď.).

tvrdenie by vsadil aj život.⁹

I toto tvrdenie ohľadom umelej inteligencie však bolo v priebehu ďalších desaťročí otrasené. V deväťdesiatych rokoch sa Hofstadter stretol s programom EMI (Experiments in Musical Intelligence), ktorý napísal hudobník, skladateľ a profesor hudby David Cope. Pôvodným zámerom pre napísanie EMI bolo vytvoriť softvérovú podporu pre Copeho komponovanie hudby, ktorá by bola schopná automaticky vytvárať časti kompozície v osobitnom Copeho štýle. Napriek tomu sa EMI preslávila produkovaním diel v štýle klasických skladateľov, akými boli Bach a Chopin.

EMI *komponovala* využívajúc veľkú množinu pravidiel, ktoré vytvoril Cope a ktoré boli zamerané na zachytenie všeobecnej syntaxe komponovania. Tieto pravidlá boli následne aplikované na mnohé príklady z konkrétnych diel skladateľov s cieľom vytvoriť nové dielko v konkrétnom štýle konkrétného skladateľa.

Sám Hofstadter to na stretnutí v centrále Google emotívne opisuje takto: „Sadol som si k svojmu klavíru a hral som jednu z mazuriek vytvorených programom EMI podľa štýlu Chopinových mazuriek. Neznelo to presne ako Chopin, avšak znelo to veľmi podobne ako Chopin a ako súvislá hudba, takže som sa cítil *hlboko* znepokojený.“¹⁰

Toto veľké znepokojenie u Hofstadtera vychádzalo z jeho prežívania hudby ako emočného dotyku s osobou, ktorá ju skomponovala. Hudbu vnímal ako možnosť priameho prístupu do duše skladateľa, a preto pre neho na svete neexistovalo nič ľudskejšieho ako vyjadrenie sa hudbou. Myšlienka, že povrchná manipulácia so vzormi môže priniesť veci, ktoré znejú, akoby vychádzali z ľudského srdca, bola pre neho veľmi znepokojujúca a jeho úvahy to úplne rozbilo.¹¹

Toto znepokojenie bolo ešte umocnené, keď následne požiadal osadenstvo prestížnej Eastman School of Music v Rochesteri, v New Yorku, aby porovnali konkrétnu mazurku od Chopina a mazurku skomponovanú programom EMI. Hofstadter bol šokovaný, keď mnohí z fakulty označili kompozíciu od EMI ako pravé dielo Chopina a to skutočné Chopinovo dielo ako menej dokonalý výtvor EMI.¹²

9 HOFSTADTER, Gödel, Escher, Bach: an Eternal Golden Braid, s. 676.

10 MITCHELL, Artificial Intelligence, s. 9.

11 por. MITCHELL, Artificial Intelligence, s. 9.

12 HOFSTADTER, D. Staring Emmy Straight in the Eye – and Doing My Best Not to Flinch. In: DARTNELL, T. Creativity, Cognition, and Knowledge. Westport, Conn.: Praeger, 2002, s. 67-100.

Douglas Hofstadter túto skúsenosť zaujímavo hodnotí: „Bol som zhrozený z EMI, nenávidel som ju a cítil som sa ňou extrémne ohrozený. EMI hrozilo zničením toho, čo som si najviac cenil na ľudskosti. Myslím, že EMI bolo najtypickejším príkladom obáv, ktoré som mal z umelej inteligencie.“¹³

Medzi spoločnosťami, ktoré sa venujú výskumu a vývoju umelej inteligencie, zaberá (aj vzhľadom na svoje akvizície technologických start-up-ov) Google vedúce postavenie.¹⁴ O to viac je pre Hofstadtera znepokojujúce vedomie, že práve Google si osvojil víziu Raya Kurzweila¹⁵ o *singularite*¹⁶, v ktorej umelá inteligencia so schopnosťou učiť sa, zlepšovať sa a samostatne sa vyvíjať rýchlo dosiahne a následne prevýši inteligenciu človeka.

Hoci Hofstadter – a nielen on – vážne pochybuje o predpoklade singularity v oblasti umelej inteligencie, vidiac pokroky, ktoré sa v oblasti vývoja umelej inteligencie za posledné desaťročia dosiahli, pociťuje obavy a rozrušenie z možného naplnenia Kurzweilovej predikcie. Na stretnutí v centrále Google to komentuje takto: „Bol som vydesený týmto scenárom. Veľmi skeptický, no v rovnakom čase, mysliac si, že i keď predpokladaný časový harmonogram je nereálny, možno majú pravdu. A potom budeme úplne zaskočení.

13 MITCHELL, *Artificial Intelligence*, s. 10.

14 Výskum, rozvoj a aplikácia umelej inteligencie má v portfóliu Google (resp. v súčasnosti Alphabet) veľmi široké zastúpenie: autonómne vozidlá, rozpoznávanie reči a inteligentná strojová asistencia, pochopenie prirodzeného jazyka a metódy prekladu, počítačovo generované umenie, riešenie logických hier (AlphaGO až AlphaZero), inteligentné humanoidné roboty a pod.

15 Ray Kurzweil je svetoznámy vynálezca a kontroverzný futurista, ktorý predstavil myšlienku singularity umelej inteligencie. Google zamestnal Kurzweila, aby pomohol túto víziu zrealizovať.
por. MITCHELL, *Artificial Intelligence*, s. 4.

16 Pojem singularita sa v mnohých oblastiach matematiky a fyziky používa na vyjadrenie stavu, ktorý je zvláštny, nedefinovaný, odchyľujúci sa z trendu, či množiny, stav prekračujúci očakávajúce parametre, nekonečný, či v daných podmienkach neriešiteľný (napr. aktuálne fyzikálne modely, ktorými nedokážeme popísať singulárne body, akými sú napr. Veľký tresk a čierne diery).

Pod termínom technologická singularita sa myslí teoretický bod vo vývoji vedeckej civilizácie (znalostnej spoločnosti), v ktorom sa technologický pokrok zrýchli do nekonečna a prevýši všetky predpovede.

Singularita [on-line]. [cit. 3. augusta 2020].

Dostupné na internete: <<https://sk.wikipedia.org/wiki/Singularita>>

Singularitou v umelej inteligencii je myslený stav, ktorý nastane, ak umelá inteligencia so schopnosťou učiť sa, zlepšovať sa a samostatne sa vyvíjať rýchlo dosiahne a následne prevýši inteligenciu človeka.

Teda situácia, keď sa počítačové systémy stanú inteligentnejšími než ľudia.

por. MITCHELL, *Artificial Intelligence*, s. 10.

Budeme si myslieť, že sa nič nedeje a zrazu, skôr než si to uvedomíme, budú počítače múdrejšie než my.“ A keď sa to stane, „budeme nahradení, staneme sa reliktom. Zostaneme v prachu.“ Končiac svoj príhovor dodáva: „Možno sa to stane, ale nechcem, aby sa to stalo priskoro. Nechcem, aby moje deti zostali v prachu.“¹⁷

Priamo na inžinierov, vývojárov a vedcov z Google sa v závere Douglas Hofstadter obracia, hovoriac: „Považujem to za veľmi hrozivé, veľmi znepokojujúce, veľmi smutné a tiež považujem za hrozné, desivé, bizarné, bezradné i zarážajúce to, že ľudia sa slepo a akoby v ošiali hrnú do vytvárania týchto vecí.“¹⁸

Sú obavy Douglasa Hofstadtera opodstatnené? Alebo sú len ovocím prekročenia jeho vízií o rozvoji umelej inteligencie spôsobom, ktorý svojho času pre neho, ale i pre druhých nebolo možné predikovať? V tejto práci sa pokúsime uchopiť, pomenovať a možno i definovať viaceré aspekty z oblasti umelej inteligencie, aby sme si mohli ozrejmiť základné hodnoty a postoje, vďaka ktorým i systémy umelej inteligencie môžu slúžiť nášmu proaktívnemu postoju k rozvoju sveta vo svetle evanjelia a pozvania, ktoré sme od Pána dostali.

17 MITCHELL, *Artificial Intelligence*, s. 10-11.

18 MITCHELL, *Artificial Intelligence*, s. 10-11.

Úvod

Prežívame dejinný úsek, ktorý sa hrdí takými prívlastkami ako digitálna disrupcia, rodiaca sa informačná a znalostná spoločnosť, zmena paradigmy, permanentná technologická revolúcia, atď. V tomto kontexte takmer každá technologická novinka sa pýši dodatkom, že v jej útrobach je implementovaná umelá inteligencia. Nájdeme a vieme vymenovať veľa oblastí, v ktorých prvky umelej inteligencie zohrávajú stále dôležitejšiu úlohu: v lekárskej diagnostike a liečbe, v riadení technologických procesov, predikcii vývoja, spracovaní obrazu, analýze a syntéze reči, kybernetickej bezpečnosti¹⁹, doprave, zábave, komunikačnej technike, genetickom výskume, jadrovej fyzike, astrofyzike, atď.

Javí sa, že umelá inteligencia nie je len *buzzword*, ktorý má pomôcť technologickým firmám presadiť sa a zarobiť, či okrášliť stratégie lídrov dnešného sveta, ale ide o reálny koncept a integrálnu súčasť technologickej budúcnosti našej civilizácie. Systémy umelej inteligencie vo viacerých oblastiach posúvajú úroveň poznania a skúmanie míľovými krokmi vpred. Dokážu byť extrémne nápomocné pri záchrane ľudských životov, ochrane zdravia i zaradení sa do reálneho života po ťažkých chorobách a úrazoch. Stávajú sa takmer nepostrádateľnými asistentmi v bežnom živote, v doprave, komunikácii, zábave. Prinášajú extrémny posun vpred pri analýze a spracovaní vedeckých dát. Pridávajú ďalší stupeň ochrany pri bezpečnostných systémoch, v obrane, v boji so zločinom a pod. Takmer v každej oblasti ľudskej činnosti a života človeka nájdeme niečo, v čom sa použitie prvkov umelej inteligencie stáva veľmi osožným.

Využitie umelej inteligencie má však aj svoju temnú stranu a riziká: jednou z privilegovaných oblastí vývoja umelej inteligencie sú autonómne zbraňové systémy a samostatné nasadenie v boji; využitie prvkov umelej inteligencie v kybernetickej kriminalite je už teraz nočnou morou informačnej bezpečnosti; zneužitie umelej inteligencie pri rozpoznávaní tvárí, komplexnom monitoringu a kategorizácii ľudí, detekcii a predikcii vývoja v spoločnosti je svätým grálom akéhokoľvek autoritárskeho režimu; vytváranie psychologických, zdravotných, sociologických, či iných profilov ľudí spájaním informácií z verejných zdrojov, sociálnych sietí a metadát dokáže vytvoriť nekompromisnú verejnú

19 Sám autor k veľkej spokojnosti využíva umelú inteligenciu v oblasti kybernetickej bezpečnosti a ochrany dát.

sondu do ľudskej duše a bezprecedentne odhaliť súkromie človeka²⁰; zneužitie falošnej identity, či vytvorenie falošných informácií o človeku²¹ môže danú osobu priviesť k totálnemu spoločenskému i osobnému kolapsu...

V kontexte vyššie uvedeného – ak také osobnosti, ako napr. Ray Kurzweil, predikujú umelej inteligencii schopnosť konkurovať, ba až nahradiť ľudskú bytosť a iní, ako napr. Douglas Hofstadter, vidia reálne riziká vyplývajúce z potenciálu napredovania v jej vývoji, **pridávame sa k hlasom, ktoré volajú po etickom zhodnotení, zvážení morálnych aspektov a definovaní pravidiel pre vývoj, používanie i samotné fungovanie systémov umelej inteligencie.**

Uvedomujeme si, že až na pár výnimiek, sa v súčasnosti umelou inteligenciou nazývajú informačné systémy, ktoré majú pramálo spoločné s víziou Kurzweila a obavami Hofstadtera²², t.j. s víziou človeku konkurujúcej *skutočnej* umelej inteligencie. Nech však ide o akýkoľvek stupeň, či schopnosti týchto systémov, treba k nim zodpovedne pristupovať, aby sme v spoločnosti nielenže znovu nemuseli boľavo objavovať to známe *dobrý sluha a zlý pán*, ale aj nakoniec nezistili, že z umelej inteligencie dobrého sluhu vo svojej podstate ani nedokážeme reálne vytvoriť²³.

Pri úvahe o dobrom sluhovi a zlom pánovi sme si podvedome dovolili ignorovať diskusiu o uznaní práv osoby pre pokročilé systémy umelej inteligencie, tzv. superinteligencie. Je to

20 Systémy umelej inteligencie sú na základe analyzovaných dát a metadát častokrát schopné vytvoriť psychologický profil osoby, odhaliť jeho sexuálnu orientáciu, detekovať psychické problémy a pod.

21 Napr. vytvorenie falošnej videokonferencie (jedna z foriem tzv. DeepFake), na ktorej sa verejný činiteľ priznáva k veciam, ktoré nevykonal; alebo vytvorenie falošného videa, na ktorom človek s vysokým morálnym kreditom v spoločnosti koná nemravné veci; a pod.

22 Vo väčšine implementácií skôr ide o IA (intelligent assistance) a nie o AI (artificial intelligence).

23 Toto riziko sa reálne zvažuje pri debatách o všeobecnej umelej inteligencii (rôzne formy superinteligencie) nielen v rovine matematických teorémov popisujúcich komplexnosť takých systémov (napr. Riceov teorém), ale – ako prirodzený a nutný rozmer zodpovedného využitia superinteligencie – aj v rovine etickej a filozofickej.

CHOI, Q. CH. *Superintelligent AI May Be Impossible to Control; That's the Good News* [on-line]. [cit. 26. januára 2021].

Dostupné na internete: <<https://spectrum.ieee.org/tech-talk/robotics/artificial-intelligence/super-artificialintelligence>>

ALFONSEC, M., CEBRIAN, M. et al. *Superintelligence Cannot be Contained: Lessons from Computability Theory* [on-line]. [cit. 26. januára 2021].

Dostupné na internete: <<https://jair.org/index.php/jair/article/view/12202>>

však diskusia, ktorá sa určitým spôsobom dotýka najhlbšieho jadra Hofstadterových obáv ohľadom umelej inteligencie a s ňou súvisiacej fiktívnej alebo reálnej singularity. Ak sa nad tým zamyslíme – podobne, ako to po opisovanej prednáške v centrále Google urobila prof. Melanie Mitchellová²⁴ – Hofstadterove obavy nie sú ničím novým. Svet, ktorý nerozumie technológiám umelej inteligencie a je naviac živený katastrofickými scenármi science fiction, je týchto obáv už niekoľko desaťročí plný.

Novými sú skôr dva fenomény:

- obavy spojené s rizikami pokročilej umelej inteligencie čoraz viac vyjadrujú ľudia, ktorí sú v tejto oblasti doma a častokrát patria k prvotriednym odborníkom, či technologickým vizionárom.²⁵ Netreba však obavy spájať len s víziou nadľudskej

24 MITCHELL, *Artificial Intelligence*, s. 11-12.

25 Už v roku 2014 svetoznámy fyzik Stephen Hawking vyhlásil: „Vývoj úplnej umelej inteligencie by mohol znamenať koniec ľudskej rasy.“ A vysvetľuje: „Ľudia, ktorí sú obmedzení pomalou biologickou evolúciou, nemohli by (s umelou inteligenciou) súťažiť a boli by nahradení.“

CELLAN-JONES, R. *Stephen Hawking Warns Artificial Intelligence Could End Mankind*. In: *BBC News*. [on-line]. 2014, 2. 12. [cit. 5. augusta 2020].

Dostupné na internete: <<https://www.bbc.com/news/technology-30290540>>

Elon Musk, vizionár a zakladateľ spoločností Tesla a SpaceX, v tom istom roku varuje pre umelou inteligenciu hovoriac, že ide pravdepodobne o „naše najväčšie existenciálne ohrozenie“ a že „umelou inteligenciou privoláme démona.“

MCFARLAND, M. *Elon Musk: „With Artificial Intelligence, We Are Summoning the Demon“*. In: *Washington Post*, 2014, 24. 10.

Elon Musk pokračoval v tejto zmysle aj vo svojich varovaniach z r. 2019 a 2021.

Elon Musk pritom v rámci Tesly vyvíja jednu z najpokročilejších umelých inteligencií pre autonómne vozidlá a v rámci svojich vizionárskych aktivít (Neuralink Corporation) vytvára rozhranie pre rozšírenie možností človeka, ako napr. prepojenie ľudského mozgu a umelej inteligencie pre dosiahnutie symbiózy, ktorú by sme poľahky mohli zaradiť do oblasti transhumanizmu.

Neuralink [on-line]. [cit. 5. augusta 2020].

Dostupné na internete: <<https://en.wikipedia.org/wiki/Neuralink>>

Bill Gates, spoluzakladateľ softvérového gigantu Microsoft, dopĺňa: „V tomto súhlasím s Elonom Muskom i viacerými ďalšími a nemôžem pochopiť, prečo niektorí ľudia nie sú znepokojení.“

Hi Reddit, I'm Bill Gates and I'm back for my third AMA. Ask me anything. In: *Reddit*. [on-line]. 2015, 28. 1. [cit. 7. augusta 2020].

Dostupné na internete:

<https://www.reddit.com/r/IAmA/comments/2tzjp7/hi_reddit_im_bill_gates_and_im_back_for_my_third/>

Významné varovanie je vyjadrené aj v knihe *Human Compatible* od v súčasnosti jednej z najväčších autorít v oblasti umelej inteligencie, prof. Stuarta Russella.

umelej inteligencie – oveľa aktuálnejšie je riešenie mnohých výziev (medzi nimi i etických) spojených s prvkami umelej inteligencie, ktoré sú aktuálne vyvíjané a začínajú sa masovo používať naprieč rôznymi oblasťami modernej spoločnosti²⁶;

- riziká spojené s pokročilou umelou inteligenciou sú často spojené s disproporciou medzi vnímaním ľudstva a umelou inteligenciou²⁷, a tak nastavujú zrkadlo nášmu pohľadu na podstatu ľudského bytia²⁸: kým sme, čo nás definuje, do akej miery sme redukovateľní na technologické vyjadrenie, aká je hodnota a podstata ľudského bytia, atď... Teda otázky skôr filozofické, antropologické a tiež i teologické²⁹, ktoré sú tým aktuálnejšie, čím viac sa posúvame vo vývoji systémov umelej inteligencie dopredu a čím viac rastie náš apetít po vytvorení umelej inteligencie schopnej

Vo svojej práci rieši i tzv. „zásadnú chybu v srdci umelej inteligencie, ktorá – ak nebude vyriešená – môže prerásť do katastrofy“ (pozri komentár od Iana Sample z Guardianu v rámci Books of the Year).

RUSSELL, S. *Human Compatible*. Penguin Books, 2020, s. III.

A Judea Pearl, počítačový vedec a filozof, autor jedného z princípov využívaných v systémoch umelej inteligencie i známej knihy *The Book of Why*, dodáva: „Human Compatible ma primálo osvojiť si Russelove starosti ohľadom našej schopnosti ovládať náš prichádzajúci výtvor – super inteligentné stroje. Keďže na rozdiel od alarmistov a futuristov mimo odboru je Russell vedúcou autoritou v oblasti umelej inteligencie.“

RUSSELL, *Human Compatible*, s. IV.

- 26 Viacerí významní autori (Mitchel Kapor, Rodney Brooks, Gary Marcus,...) hovoria: „áno, mali by sme sa ubezpečiť, že systémy umelej inteligencie sú bezpečné a nepoškodzujú ľudí, ale akékoľvek správy o nadľudskej AI sú veľmi prehnané.“

Por. MITCHELL, *Artificial Intelligence*, s. 13.

- 27 Vedomie, že mnohí vývojári systémov pokročilej umelej inteligencie podceňujú inteligenciu ľudí, nás poľahky môže priviesť k opačnému extrému – podceňovaniu výkonu, možností a dôvery či nádejí vkladaných do súčasných technológií umelej inteligencie.

Por. MITCHELL, *Artificial Intelligence*, s. 12-13.

- 28 Z tohto vlastne prameňa Hofstadterove obavy – či sa jeho predstava o ľudskom bytí a jeho hodnote nerozpadne do ním spomínaného *prachu*: nie je to o umelej inteligencii, ktorá sa stáva oveľa múdrejšou, príliš invazívnou, veľmi škodlivou alebo priveľmi použiteľnou. Je to zhrozenie z uvedomenia si, že inteligencia, kreativita, emócie a možno dokonca i vedomie by mohli byť príliš ľahko reprodukovateľné a vytvárané balíkom trikov, vďaka ktorým povrchný súbor algoritmov spojený s hrubou výpočtovou silou môže vysvetliť ľudského ducha.

A dodáva: „Ak by sa takéto mysle nekonečnej jemnosti a komplexnosti a emocionálnej hĺbky (myslí tým veľikánov ľudstva) dali trivializovať malým čipom, zničilo by to môj pohľad na človečenstvo.“

MITCHELL, *Artificial Intelligence*, s. 11-12.

- 29 A tým zároveň ide o aj o otázky etické, legislatívne, právne, sociologické, psychologické, medicínske,...

konkurovať ľudskému bytiu³⁰.

V skutočnosti sa však nachádzame vo svete, ktorý presahuje akademické úvahy o potenciáli a rizikách pokročilej umelej inteligencie. Viaceré moderné armády vyvíjajú a koketujú s nasadením autonómnych bojových systémov (USA, Rusko, Čína, Izrael,...)³¹, v niektorých krajinách sa pripravuje legislatíva pre autonómne vozidlá³², prvá krajina uznala práva pre inteligentného robota³³, prvým systémom, ktorý ešte pred prepuknutím pandémie ochorenia Covid-19 informoval, že prichádza nová epidémia koronavírusu, bola kanadská lekárska umelá inteligencia³⁴, lieky, ktoré by mohli účinkovať voči SARS-CoV-2 navrhla taktiež umelá inteligencia³⁵, začínajú sa realizovať lekárske operácie vedené

30 Účastníci Hofstadterovej prednášky v centrále Google boli prakticky všetko ľudia, ktorých snahou bolo naplnenie Kurzweilovej vízie o umelej inteligencii schopnej v konečnom dôsledku prevýšiť inteligenciu človeka.

Por. MITCHELL, *Artificial Intelligence*, s. 12-13.

31 Konkrétnu citáciu neuvádzame, pretože v tejto oblasti je neustále prichádzajúcich správ také množstvo a v takej odbornej granularite, že pre utvorenie uceleného obrazu by sme mohli uviesť niekoľko desiatok citácií...

32 Krajiny, v ktorých už prebieha testovanie autonómnych automobilov, postupne pripravujú svoju legislatívu. Kalifornia ako jeden z prvých federálnych štátov USA zakomponovala do svojho zákonníka o vozidlách (Vehicle code) oddiel 16.6, (Divison 16.6. Autonomous Vehicles [38750]), ktorý upravuje oblasť autonómnych vozidiel.

Autonómne vozidlá [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<http://akgunis.sk/autonomne-vozidla/>>

33 28. októbra 2017 udelila Saudská Arábia na podujatí Future Investment Initiative conference občianstvo robotovi Sofii.

Robot Sophia speaks at Saudi Arabia's Future Investment Initiative [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://youtu.be/dMrX08PxUNY>>

PAROULKOVÁ, V. *Lidská práva pro roboty? Evropská unie chystá právní statut tzv. umělé osoby* [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://plus.rozhlas.cz/lidska-prava-pro-roboty-evropska-unie-chysta-pravni-statut-tzv-umele-osoby-6598059>>

34 MIHULKA, S. *První varování před koronavirem Wuhan poslala umělá inteligence BlueDot*. [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://www.osel.cz/11002-prvni-varovani-pred-koronavirem-wuhan-poslala-umela-intelligence-bluedot.html>>

35 BECK, B. R. et al. *Predicting commercially available antiviral drugs that may act on the novel coronavirus (2019-nCoV), Wuhan, China through a drug-target interaction deep learning model*. In: *bioRxiv*. [online]. 2020, s. 2020.01.31.929547. [cit. 6. augusta 2020]. DOI: 10.1101/2020.01.31.929547.

systémami umelej inteligencie, japonská lekárska umelá inteligencia s veľkou presnosťou a prekračujúcou schopnosťou najlepších špecialistov detekuje vzácne typy rakoviny, umelá inteligencia sa podieľa na špičkových fyzikálnych výskumoch³⁶, prakticky celá špička firiem špecializujúcich sa na kybernetickú bezpečnosť masívne implementuje prvky umelej inteligencie do svojich riešení³⁷, pokroky v rozpoznávaní reči v spojení s ďalšími prvkami umelej inteligencie umožnili vznik inteligentných asistentov³⁸, systémy spracovania a rozpoznávania obrazu kategorizujú, titulujú, filtrujú, či inak spracúvajú snímky a obrazový materiál v takmer všetkých cloudových riešeniach a koniec koncov primitívnejšie systémy umelej inteligencie sú súčasťou prakticky každého spracovania obrazu, navrhovania trás v mapových systémoch a marketingových nástrojov v elektronických médiách, mobilných zariadeniach a e-shopoch...

V kontexte informačnej spoločnosti sa teda už ne bavíme o tom, či a kedy prvky umelej inteligencie nasadiť – veď sú tu medzi nami a ich nasadenie sa neustále rozširuje – ale ako, to znamená za akých podmienok, pre aké ciele, akým spôsobom a s akými dôsledkami by mala byť umelá inteligencia súčasťou nášho sveta. Súčasťou sveta, v ktorom má človeku a spoločnosti slúžiť (vizionár by povedal *koexistovať*), ľudia i spoločnosť majú mať jasno v tom, ako vyzerá etický návrh, realizácia a využívanie systémov umelej inteligencie, pričom samotná umelá inteligencia v svojom autonómnom a adaptívnom fungovaní „dokáže“ etické hodnoty a nastavené *morálne normy dodržať*³⁹.

Dostupné na internete: <<https://www.biorxiv.org/content/10.1101/2020.01.31.929547v1>>

- 36 Napr. HE, S. et al. *Learning to predict the cosmological structure formation*. In: *Proceedings of the National Academy of Sciences*. [online]. 2019, roč. 116, č. 28, s. 13825. [cit. 6. augusta 2020]. DOI: 10.1073/pnas.1821458116. Dostupné na internete: <<https://www.pnas.org/content/116/28/13825>>
SUMMER, T. *The first AI universe sim is fast and accurate—and its creators don't know how it works*. [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://phys.org/pdf480780725.pdf>>

- 37 Miesto uvedenia všetkých významných implementátorov umelej inteligencie (CheckPoint, Kaspersky, Sophos, Cisco, Eset, IBM, Fortinet,...) v oblasti kybernetickej bezpečnosti by sme oveľa ľahšie vymenovali tých, ktorí tak ešte nekonajú.
- 38 Spomeňme Siri od Apple, Cortanu od Microsoftu, Now a Assistant z dielne Google, Alexu od Amazonu... *Virtual assistant* [on-line]. [cit. 6. augusta 2020].
Dostupné na internete: <https://en.wikipedia.org/wiki/Virtual_assistant>
- 39 Viac krát sa stalo, že systémy s umelou inteligenciou museli byť vypnuté, lebo po uvedení do prevádzky sa vyvíjali nesprávnym smerom. Napr.:

V súčasnosti preto silnejú hlasy vyjadrujúce potrebu skúmať, navrhnúť, prijať a realizovať etický rámec vývoja, používania a fungovania umelej inteligencie – či už ide o oblasť vedeckého bádania⁴⁰ alebo vývoja a realizácie⁴¹. Od petícií a otvorených listov adresovaných OSN i niektorým vládam badať prechod k snahe systematicky túto oblasť podchytiť a preskúmať na národnej i medzinárodnej úrovni⁴².

Naliehavosť tejto témy vyjadruje aj záujem, ktorý jej prikladá Katolícka cirkev. Po viacerých

KRAFT, A. *Microsoft shuts down AI chatbot after it turned into a Nazi*. [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://www.cbsnews.com/news/microsoft-shuts-down-ai-chatbot-after-it-turned-into-racist-nazi/>>

HAMILTON, I. A. *Amazon built an AI tool to hire people but had to shut it down because it was discriminating against women*. [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://www.businessinsider.com/amazon-built-ai-to-hire-people-discriminated-against-women-2018-10>>

GRIFFIN, A. *Facebook's artificial intelligence robots shut down after they start talking to each other in their own language*. [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://www.independent.co.uk/life-style/gadgets-and-tech/news/facebook-artificial-intelligence-ai-chatbot-new-language-research-openai-google-a7869706.html>>

- 40 Viaceré iniciatívy poukazujúce na riziká a vyzývajúce k prijatiu potrebných pravidiel. Prehľad viacerých otvorených listov a petícií je uvedený na stránke:

Compilation of open letters against autonomous weapons. [on-line]. [cit. 19. augusta 2020].

Dostupné na internete: <<https://autonomousweapons.org/compilation-of-open-letters-against-autonomous-weapons/>>

- 41 Angažovanosť vyúsťujúca v kontexte aktuálnej politickej a spoločenskej situácie do viacerých aktivít vývojárov, napr. protest softvérovej inžinierky Laury Nolanovej proti zapojeniu Google do projektu Maven. MCDONALD, H. *Ex-Google worker fears 'killer robots' could cause mass atrocities*. [on-line]. [cit. 19. augusta 2020].

Dostupné na internete: <<https://www.theguardian.com/technology/2019/sep/15/ex-google-worker-fears-killer-robots-cause-mass-atrocities>>

V kontexte udalostí rokov 2019-2021 – napr. nepokoje v USA a v Honkongu – sa viaceré korporácie pracujúce v oblasti aplikácie prvkov umelej inteligencie do bezpečnostných systémov rozhodlo obmedziť svoje aktivity, či zrušiť niektoré projekty a kontrakty.

- 42 Prehľad dokumentov do roku 2019 je sumarizovaný napríklad v tabuľke *Ethical guidelines for AI by country of issuer*, ktorá je súčasťou štúdie:

JOBIN, A., IENCA, M., VAYENA, E. *Artificial Intelligence: the global landscape of ethics guidelines*. [on-line]. [cit. 19. augusta 2020].

Dostupné na internete: <<https://arxiv.org/pdf/1906.11668>>

širšie koncipovaných aktivitách⁴³ Pápežská akadémia pre život vo februári 2020 zorganizovala konferenciu *renAIssance 2020*, ktorá sa venovala etike v oblasti umelej inteligencie. Jedným z výsledkov konferencie bolo aj podpísanie *Výzvy na etiku v umelej inteligencii* s cieľom „akcentovať etický prístup k umelej inteligencii a podporovať zmysel pre zodpovednosť medzi organizáciami, vládami a inštitúciami s cieľom vytvoriť budúcnosť, v ktorej digitálne inovácie a technologický pokrok slúžia ľudským schopnostiam a tvorivosti, a nie ich postupnému nahrádzaniu“⁴⁴. Signatármi výzvy boli okrem Pápežskej akadémie pre život aj talianska vláda, FAO a také spoločnosti ako Microsoft, IBM a pod.

Zámerom konferencie bol nielen vstup do rozpravy o etických kritériách umelej inteligencie v jednotlivých oblastiach, ale aj diskusia so snahou o premostenie pohľadov jednotlivých krajín, firiem či záujmových skupín na oblasť etiky v problematike umelej inteligencie a s cieľom stanoviť konkrétne etické kritériá. Sumarizáciu tohoto úsilia vyjadruje vyššie spomenutý záverečný dokument, ktorého súčasťou je i formulovanie šiestich princípov pre použitie umelej inteligencie pre dobro človeka a bez strachu zo zneužitia.

Predložený náčrt problematiky umelej inteligencie v optike morálnych aspektov a etiky môže slúžiť ako uvedenie do tejto práce, v ktorej by sme ponajprv chceli predstaviť zjednodušenou formou základy umelej inteligencie v kapitole *Základy umelej inteligencie*. Pôvodne sme chceli uviesť, že znalý čitateľ môže túto kapitolu bez problémov preskočiť. Uvedomili sme si však, že mnohé, v tejto kapitole uvedené, postrehy a implikácie presahujú jednoduché vovedenie do problematiky AI a skôr sa priamo týkajú témy našej práce.

V kapitole *Limity a riziká súčasných systémov umelej inteligencie* sa pokúsime poukázať

43 Napr. Hackaton 2018, ktorý sa však netýkal priamo umelej inteligencie, ale v duchu hesla *bridge the gap* skôr širšiemu využitiu informačných technológií na premostenie a podporu marginalizovaných, ktorí trpia napr. digitálnym rozdelením (digital divide) a pod. Viac o digital divide napr. v kapitole *Digitálne diferencovanie ľudstva (digital divide)* práce:

ŠANTAVÝ, P.. *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi* [licenciátska práca]. [on-line]. Bratislava: RKCMBF UK, 2017, s. 30-32. [cit. 19. augusta 2020].

Dostupné na internete:

<https://peter.santavy.cloud/docs/Analýza_virtualneho_sвета_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

44 *Rome Call for AI Ethics*. [on-line]. [cit. 19. augusta 2020].

Dostupné na internete: <<http://www.academyforlife.va/content/pav/en/events/intelligenza-artificiale.html>>

na limity a riziká existujúcich systémov AI, popísať technologické výzvy v oblasti bezpečnosti procesov umelej inteligencie a s tým súvisiacej kybernetickej bezpečnosti a akcentovať viaceré dôsledky v psychologickkej, spoločenskej i vojenskej oblasti, ktorých zneužitie môže znamenať vážne ohrozenie pre jednotlivcov i spoločnosť.

V kapitole *Umelá inteligencia a etika* sumarizujeme etické dôsledky limitov, rizík i možností existujúcich systémov AI, zmapujeme súčasné aktivity na poli etiky a regulácie umelej inteligencie a načrtujeme aktuálny pohľad na etické otázky súvisiace s vývojom a využívaním týchto systémov tak z pohľadu osobností a spoločností zaangažovaných v tejto oblasti, ako aj z pohľadu štátnych inštitúcií a Cirkvi.

Kapitola *Navrhnuté riešenie etických problémov ANI* na základe analýzy súčasného stavu rozvoja systémov AI i predpokladaných možností ich ďalšieho vývoja a na základe porovnania i zhodnotenia existujúcich etických pravidiel a právnych noriem ponúkne pretavenie zistených skutočností do návrhu základnej štruktúry všeobecných i špecifických etických zásad, ktoré musia byť aplikované pri etickom návrhu, realizácii a využívaní systémov slabej umelej inteligencie (ANI).

I keď na poli vývoja uvedomelej umelej inteligencie⁴⁵ nebadat' až tak veľké pokroky, ako by sa z rôznych mediálnych správ mohlo zdať, kapitola *Vízia silnej a všeobecnej umelej inteligencie* sa dotkne niektorých etických problémov aj v tejto vzrušujúcej a vizionárskej oblasti. Pravdu povediac – v tom rozsahu, v akom sa v tejto práci umelej inteligencii venujeme – debate o ľudskom bytíu podobnej umelej inteligencii, resp. superinteligencii sa v súčasnosti asi nedá vyhnúť.

Hlavným cieľom práce je na základe analýzy aktuálneho stavu rozvoja systémov umelej inteligencie a existujúcich etických pravidiel snaha navrhnúť základnú štruktúru zásad, ktorá musí byť podchytená pri etickom návrhu, realizácii i využívaní akýchkoľvek systémov umelej inteligencie.

Z hľadiska metodológie používame v dizertačnej práci viacero metód:

- analyticko-syntetickú metódu:
 - sumarizovanie aktuálneho stavu (technologického) rozvoja AI;
 - analýza známych limitov a rizík zlyhania i problematických dôsledkov nasadenia systémov AI.

⁴⁵ Známej aj pod skratkou AGI (artificial general intelligence).

- analýza existujúcich etických pravidiel v oblasti AI.
- komparatívnu metódu:
 - porovnanie jednotlivých oblastí AI z pohľadu možnosti implementácie etických pravidiel;
 - porovnanie potencionálnych schopností sofistikovaných systémov AI s ľudským modelom správania, schopnosťou myslieť a vedomím.
- aplikáciu a návrh:
 - na základe analýzy súčasného stavu rozvoja systémov AI i predpokladaných možností ich ďalšieho vývoja a na základe porovnania i zhodnotenia existujúcich etických pravidiel predkladáme aplikáciu zistených skutočností do návrhu základnej štruktúry etických noriem, ktorá musí byť podchytená pri etickom návrhu, realizácii a využívaní akýchkoľvek systémov umelej inteligencie (ANI);
 - vklad do polemík o realizácii uvedomelej umelej inteligencie (AGI) na základe predloženej analýzy súčasných systémov AI a ich konfrontácie s ľudskou schopnosťou myslenia a sebauvedomenia.

V rámci aplikovaného výstupu v práci riešime analýzu aktuálneho stavu rozvoja systémov AI, komplexný popis limitov i rizík zlyhania a z toho prameniacych dôsledkov pre človeka i spoločnosť, a v neposlednom rade diskutujeme etické princípy a navrhované právne normy, resp. regulácie.

K základnému výskumu patrí porovnanie fenoménu uvedomelej umelej inteligencie a ľudského bytia v kontexte osoby a skutočnej schopnosti myslieť, návrh interdisciplinárneho rámca a spôsobu prístupu k riešeniu problémov ANI, stanovenie základných etických princípov a noriem, ktoré musia byť základom pre etický návrh, realizáciu a využívanie systémov slabej umelej inteligencie vo všeobecnosti, a tiež formulácia špecifických odporúčaní a parciálnych usmernení pre osobitné oblasti nasadenia systémov AI.

1. Základy umelej inteligencie

*Umelá inteligencia je ako Colombova žena, nikto ju nevidel a všetci o nej hovoria.*⁴⁶

I keď umelú inteligenciu pokladáme za niečo nové a bytostne sa viažuce na súčasnú modernú – primárne informačnú – spoločnosť, ide o fenomén, ktorý oslovoval mysliteľov a vizionárov už niekoľko storočí dozadu, možno od počiatkov Prvej priemyselnej revolúcie⁴⁷, keďže jedným z podstatných prvkov priemyselnej revolúcie bolo nahradenie ľudskej práce strojmi. Činnosť, ktorú dovtedy manuálne vykonávali ľudia, odrazu vykonávali – a efektívnejšie vykonávali – stroje. S konkrétnou paradigmatickou zmenou vtedajšej spoločnosti sa tak začali otvárať nové horizonty aj v pohľade na zariadenia, ktoré môže človek vytvoriť.

Zaujímavým predstaviteľom tohto snaženia môže byť Kempelenov⁴⁸ Turek, mechanický stroj zostrojený v roku 1771, ktorý bol vydávaný za šachový automat. Stroj, ktorý bol prezentovaný ako automat, dokázal zohrať dobrú šachovú partiu proti ľudskému protihráčovi. V skutočnosti však išlo o mechanické zariadenie ukrývajúce skutočného šachistu, pričom podvod vyšiel na povrch až v roku 1857, t.j. niekoľko rokov po požiari, pri ktorom stroj zhorel (r. 1854).⁴⁹

Nerozoberajúc mechanické prevedenie Turka⁵⁰ je možné si uvedomiť **ideu vytvorenia**

46 doc. Ing. Ľuboš MAGDOLEN, CSc., SJF STU.

47 V rámci Prvej priemyselnej revolúcie sa ako dôsledok pokroku vo výrobných nástrojoch (osobitne v textilných manufaktúrach) a následne vynálezu parného stroja začína industrializácia a masívne využívanie strojov, ktoré nahrádzali ručnú prácu.

Priemyselná revolúcia [on-line]. [cit. 20. augusta 2020].

Dostupné na internete: <https://sk.wikipedia.org/wiki/Priemyseln%C3%A1_revol%C3%BAcia>

48 Barón Wolfgang Kempelen, pochádzajúci z nemecky hovoriacej rodiny, bol uhorský vysoký štátny úradník, polytechnik a vynálezca, ktorý sa narodil 23. januára 1734 v Bratislave.

Wolfgang Kempelen [on-line]. [cit. 20. augusta 2020].

Dostupné na internete: <https://sk.wikipedia.org/wiki/Wolfgang_Kempelen>

49 *Turek (stroj)* [on-line]. [cit. 20. augusta 2020].

Dostupné na internete: <[https://sk.wikipedia.org/wiki/Turek_\(stroj\)](https://sk.wikipedia.org/wiki/Turek_(stroj))>

50 V súčasnosti Amazon ponúka službu *Mechanical Turk* ako „trhovisko pre prácu, ktorá vyžaduje ľudskú inteligenciu“. Mechanical Turk od Amazonu spája záujemcov s úlohami ťažko realizovateľnými počítačmi so zamestnancami, ktorí za poplatok tieto úlohy vykonávajú. Takto sa napr. pripravujú veľké množiny dát (sorting and labeling in datasets) pre učenie a testovanie systémov umelej inteligencie.

stroja, ktorý hrá šach a dokáže sa v tejto hre, ktorá bola dlhé roky považovaná za prejav inteligencie, postaviť človekovi.

Tieto tendencie boli následne umocnené s príchodom Druhej a osobitne Tretej priemyselnej revolúcie, v rámci ktorej nové *mysliace stroje* (t.j. počítačové systémy) boli stále viac schopné vykonávať koncepčné, riadiace i administratívne funkcie a koordinovať tok výroby, od ťažby surovín po predaj a distribúciu konečných výrobkov i služieb.⁵¹ Technologické zariadenia a informačné systémy sa tak v podstatnej miere začali podieľať na fungovaní spoločnosti a jej premene na spoločnosť informačnú⁵².

Príchod Tretej priemyselnej revolúcie (známej tiež ako digitálnej), ktorý nastal v polovici 20. storočia, znamenal aj úsvit umelej inteligencie a jej systematického bádania.

Digitálna revolúcia, ktorá bola počiatkom informačného veku a premeny spoločnosti na spoločnosť založenú na vedomostiach, sa v kontexte permanentnej technickej revolúcie⁵³ podieľa na formovaní v súčasnosti nastupujúcej Štvrtej priemyselnej revolúcie (známej ako *Industry 4.0*). Je charakterizovaná zlúčením technológií, ktoré stierajú hranice

MITCHELL, *Artificial Intelligence*, s. 85.

51 Druhá priemyselná revolúcia spadá do obdobia druhej polovice 19. storočia až po 1. svetovú vojnu. V tomto čase sa masívnym spôsobom rozrástla priemyselná výroba a prichádza k elektrifikácii spoločnosti.

Tretia priemyselná revolúcia, ktorá sa datuje niekedy od polovice 20. storočia a je tiež známa ako digitálna, je obdobím nástupu digitálnych technológií, počítačov a moderných foriem spracovania informácií i komunikácie. Digitálna revolúcia bola začiatkom informačného veku a premeny spoločnosti na spoločnosť založenú na vedomostiach.

HUMBER, M. *Technology and Workforce: Comparison between the Information Revolution and the Industrial Revolution* [on-line]. Berkeley: University of California, 2007, s. 2-3. [cit. 20. augusta 2020]. Dostupné na internete: <<http://infoscience.epfl.ch/record/146804/files/InformationSchool.pdf>>

52 ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi* [on-line], s. 19-20. [cit. 21. augusta 2020].

Dostupné na internete:

<https://peter.santavy.cloud/docs/Analyza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

53 ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi* [on-line], s. 23-24. [cit. 20. augusta 2020].

Dostupné na internete:

<https://peter.santavy.cloud/docs/Analyza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

medzi fyzickými, digitálnymi a biologickými sférami s dôrazom na celkovú transformáciu spoločnosti na spoločnosť znalostnú a informačnú, pričom dôležitú úlohu zohráva nástup systémov a prvkov umelej inteligencie do takmer všetkých oblastí fungovania spoločnosti.⁵⁴

1.1. Úsvit umelej inteligencie

Ako už bolo spomenuté, úsvit umelej inteligencie nastával ruka v ruke s príchodom digitálnej revolúcie a s rozvojom elektronických výpočtových systémov.

Vtedajší vizionári a inventori v oblasti informačných technológií prekračovali základný i aplikovaný výskum v oblasti matematiky a informatiky pohľadom skôr filozofickým či futurologickým so základnou otázkou, kam až môže vývoj počítačových systémov zájsť.

Keďže informačné a komunikačné technológie (IKT) sa čoraz viac etablovali ako *mysliace stroje* (vysvetlenie a významové zaradenie je spomenuté v predchádzajúcej kapitole), k úvahám o umelej inteligencii bol už len malý krok, ktorý viacerí s nadšením a entuziazmom neváhali vykonať.

Medzi dôležité artefakty tohto obdobia bezpochyby patrí tzv. **Turingov test**, ktorého cieľom bolo zistiť, či je nejaký stroj inteligentný, resp. dokáže myslieť. Turingov test je založený na jednoduchej premise: ak dokáže človek aspoň päť minút konverzovať s nejakým respondentom bez toho, že by zistil, že ide o stroj (nie človeka), tak tento stroj (systém, počítač,...) úspešne prejde testom.⁵⁵ Inak povedané, dokáže stroj napodobniť ľudské myšlienky, dokáže myslieť?

54 *Prečo Industry 4.0* [on-line]. [cit. 20. augusta 2020].

Dostupné na internete: <<http://industry4.sk/o-industry-4-0/co-je-industry-4-0/>>

55 V Turingovom teste osoba v úlohe pýtajúceho sa dáva otázky dvom respondentom, s ktorými nie je v priamom kontakte, pričom jedným respondentom je človek a druhým stroj. Komunikácia je výlučne textová, trvá päť minút a rozsah konverzačných tém nie je obmedzený.

Ak pýtajúci sa spolu s hodnotiacou komisiou (hodnotiť odpovede môže viac osôb) nedokáže podľa odpovedí respondentov rozlíšiť, ktorý z nich je človek a ktorý stroj, potom tento stroj možno považovať za inteligentný. Prelomovou hranicou sa považuje 30% úspešnosť, teda minimálne traja z desiatich hodnotiteľov musia byť presvedčení, že komunikácia prebieha s človekom a nie strojom, prípadne nedokážu rozlíšiť človeka od stroja.

por. *Turingov test* [on-line]. [cit. 29. januára 2021].

Dostupné na internete: <https://sk.wikipedia.org/wiki/Turingov_test>

Test sa prvý krát podarilo zvládnuť v roku 2014 počítačovému programu Eugene, ktorý simuloval konverzáciu trinásť ročného chlapca.⁵⁶ A po ňom nasledovali ďalšie úspešné programy, čo len dokazuje, že Turingov test nie je dostatočným kritériom pre umelý systém, ktorý rozmýšľa ako človek. **Základným problémom – a jeho podstata sa týka prakticky takmer všetkých systémov umelej inteligencie dneška – je skutočnosť, že i keď nejaký systém dokáže odpovedať tak, akoby konverzácii rozumel, neznamená to, že je inteligentný a že konverzácii aj skutočne rozumie.**^{57 58 59}

Každopádne **Alan Turing⁶⁰, priekopník a otec modernej počítačovej vedy, ktorý svoj**

56 *First Turing Test success marks milestone in computing history* [on-line]. [cit. 29. januára 2021].

Dostupné na internete: <<https://phys.org/news/2014-06-turing-success-milestone-history.html>>

57 Realitu tohoto problému vystihuje aj argument čínskej izby, ktorý bol predložený filozofom Johnom Searlom v roku 1980. Ide o myšlienkový experiment názorne vyjadrujúci, že **schopnosť zmysluplne odpovedať na položené otázky nie je dostatočná na preukázanie schopnosti otázkam porozumieť a myslieť**.

SEARLE, R. J. *Minds, Brains, and Programs* In: *The Behavioral and Brain Sciences*. [on-line].

Cambridge University Press, 1980, zv. 3. [cit. 30. januára 2021].

Dostupné na internete: <<https://web.archive.org/web/20071210043312/http://members.aol.com/NeoNoetics/MindsBrainsPrograms.html>>

58 Nielen program Eugene, ale aj oveľa sofistikovanejšie systémy AI, napr. Watson od IBM, vytvárajú chyby, ktoré sú však často iné, než robia v podobných situáciách ľudia. Sú také „neľudské“, ľudsky nepochopiteľné. Je však zaujímavé, že jedným z faktorov, ktorý dokázal hodnotiteľov Eugene presvedčiť, bola schopnosť robiť chyby, ktoré sú podobné ľudským chybám.

59 V súčasnosti sa pri systémoch spracovania jazyka, inteligentných asistentoch a pod. používa na meranie schopnosti porozumieť komunikácii tzv. SQuAD (The Stanford Question Answering Dataset), ktorý bol v roku 2016 predstavený výskumníkmi Stanfordskej univerzity.

RAJPURKAR, P. et al. *SQuAD: 100 000+ Questions for Machine Comprehension of Text*, In: *Proceeding of the 2016 Conference on Empirical Methods in Natural Language Processing*. 2016, 2382-92.

60 Alan Turing (1912-1954) bol britský matematik, logik a kryptograf, ktorý súbežne (nezávisle od seba) s matematikom Johnom von Neumannom vytvoril princípy dnešných elektronických počítačov. Turing sa už počas II. svetovej vojny preslávil skonštruovaním zariadenia, ktoré umožnilo prekonať nemecké šifrovacie zariadenie Enigma. Povojnovú prácu na konštrukcii digitálnych počítačov rozšíril o štúdium neurónových sietí, uplatnenie matematických metód v biológii a skúmanie vízie mysliacich strojov.

Alan Turing [on-line]. [cit. 30. januára 2021].

Dostupné na internete:

<https://en.wikipedia.org/wiki/Alan_Turing#Pattern_formation_and_mathematical_biology>

The Turing Digital Archive [on-line]. [cit. 30. januára 2021].

Dostupné na internete: <<http://www.turingarchive.org/>>

test predstavil v roku 1950, ním akoby položil základy umelej inteligencie a odštartoval niečo, čo sa už nedalo zastaviť.

Za skutočný zrod umelej inteligencie sa však považuje **Dartmouthský seminár** – dvojmesačný výskumný projekt zaoberajúci sa umelou inteligenciou, ktorý sa uskutočnil v lete 1956 v Dartmouth college v USA.⁶¹ V podstate išlo o pracovný seminár, ktorý organizoval mladý matematik John McCarthy v spolupráci s Marvinom Minskym (bývalým kolegom zo štúdií, ktorý s ním zdieľal fascináciu inteligentnými počítačmi a neskôr sa stal zakladateľom laboratória umelej inteligencie na MIT), Claudem Shannonom (kolegom z Bell Labs/IBM a vynálezcom teórie informácií) a Nathanielom Rochesterom (pionierom v oblasti elektroniky a architektom viacerých počítačov IBM).

Prvé použitie, či **vynájdenie termínu *umelá inteligencia*** sa pripisuje práve Johnovi McCarthymu, ktorý sa tak chcel vyhraniť voči príbuznej, čerstvo sa rozvíjajúcej oblasti kybernetiky.^{62 63} Je pritom zaujímavé, že použitie prívlastku umelá sa vtedy nepáčilo prakticky nikomu na Dartmouthskom seminári, keďže mali na mysli pravú, skutočnú inteligenciu. Avšak v pracovných debatách ju bolo treba nejako nazvať a tak ju nazvali umelá inteligencia (artificial intelligence).⁶⁴ I táto drobnosť vyjadruje počiatočné snahy uchopiť novú oblasť, ktorá sa s rozvojom elektronických počítačov, matematických metód a poznania v oblasti biológie javila ako síce neznáma, no veľmi lákavá.

61 *Dartmouth workshop* [on-line]. [cit. 1. februára 2021].

Dostupné na internete: <https://en.wikipedia.org/wiki/Dartmouth_workshop>

62 Kybernetika je vo všeobecnosti veda o skúmaní, riadení, regulácii a komunikácii v dynamických systémoch v technike, biologickej sfére a spoločnosti. Predmetom kybernetiky je spracovanie informácií pre potreby riadenia a jej metódami sú systémový prístup a modelovanie pri riešení problémov.

Kybernetika [on-line]. [cit. 3. februára 2021].

Dostupné na internete: <<https://sk.wikipedia.org/wiki/Kybernetika>>

63 Na margo vyhranenia sa voči kybernetike však treba povedať, že umelá inteligencia a kybernetika vo svojom rozvoji majú styčné plochy a spoločné oblasti, napr. robotika, automatizácia, počítačové videnie a pod. Autor tejto práce, absolvujúc štúdium technickej kybernetiky na Elektrotechnickej fakulte SVŠT (aktuálne pôsobiacej pod názvom Fakulta elektrotechniky a informatiky STU), v rámci už vtedajšieho kybernetického odboru navštevoval aj základy neurónových sietí a vo svojej diplomovej práci riešil spracovanie obrazu metódami matematickej morfológie, t.j. metódami, ktoré využíva nielen robotika a automatizované systémy, ale napr. i vstupná vrstva konvolučných sietí používaných na rozpoznávanie obrazu v moderných systémoch umelej inteligencie.

64 NILSSON, N. J., MCCARTHY, J. *A Biographical Memoir*. Washington D.C.: National Academy of Sciences, 2012.

V podkladoch, ktoré boli súčasťou žiadosti o finančnú podporu na realizáciu seminára adresovanej the Rockefeller Foundation, organizátori predstavili návrh, že „**všetky aspekty učenia, resp. akékoľvek iné črty inteligencie môžu byť v princípe tak detailne popísané, že je možné vytvoriť stroje, ktoré by inteligenciu mohli simulovať**“.⁶⁵

Detailné predstavenie tohoto návrhu v sebe zahŕňalo celú množinu tém, ktoré mali byť na seminári diskutované: **strojové spracovanie ľudskej reči, neurónové siete, strojové učenie, abstraktné koncepty a myslenie, kreativita...** Treba povedať, že uvedené témy definujú oblasti skúmania a rozvoja umelej inteligencie dodnes.

Z dnešného uhľa pohľadu – vnímajúc problémy, s ktorými sa pri tvorbe systémov umelej inteligencie potýkajú už celé generácie odborníkov z rozličných oblastí a uvedomujúc si, že najvýkonnejšie počítače v roku 1956 boli milión krát pomalšie, než dnešné bežné chytré mobilné telefóny – je zarážajúci optimizmus, ktorý účastníci seminára ohľadom dosiahnutia umelej inteligencie mali: „myslíme si, že významný pokrok v riešení jedného alebo viacerých z týchto problémov sa dá dosiahnuť, ak starostlivo vybraná skupina vedcov bude na tom spoločne pracovať počas jedného leta“.⁶⁶

Z dnešného uhľa pohľadu vieme, aký nereálny a idealistický pohľad účastníci seminára mali.⁶⁷ Nech sa však v nasledovných rokoch a desaťročiach dialo čokoľvek, v roku 1956

65 MCCARTHY, J. et al.: *Proposal for the Dartmouth Summer Research Project in Artificial Intelligence*. In: *AI Magazine*. [on-line]. 1955, 27(4). [cit. 3. februára 2021].

Dostupné na internete: <<https://doi.org/10.1609/aimag.v27i4.1904>>

66 MCCARTHY, et al.: *Proposal for the Dartmouth Summer Research Project in Artificial Intelligence*. [on-line]. [cit. 3. februára 2021].

Dostupné na internete: <<https://doi.org/10.1609/aimag.v27i4.1904>>

67 McCarthy na začiatku šesťdesiatych rokov zakladá Standfordský projekt umelej inteligencie „s cieľom vytvorenia plne inteligentného stroja v rámci dekády“.

MORAVEC, H. *Mind Children: The Future of Robot and Human Intelligence*. Cambridge, Mass.: Harvard University Press, 1988, s. 20.

Budúci laureát Nobelovej ceny v tom istom čase predikuje: „Do dvadsať rokov budú stroje schopné vykonávať akúkoľvek prácu, ktorú môže vykonávať človek“.

SIMON, H. A. *The Shape of Automation for Men and Management*. New York: Harper & Row, 1965, s. 90.

Minského predikcia hovorila podobne nádejne: „v rámci jednej generácie budú problémy spojené s vytvorením umelej inteligencie v podstate vyriešené“.

MINSKY, M. L. *Computation: Finite and Infinite Machines*. Upper Saddle River, N.J.: Prentice Hall, 1967,

prakticky započal rozvoj nového vedného odboru, ktorý na Dartmouthskom seminári dostal nielen meno, ale i solídny náčrt základných cieľov.⁶⁸

V kontexte zamerania tejto práce je zaujímavé primárne inšpirovanie sa realizácie AI biologickými systémami a porovnávanie AI s ľudskou inteligenciou. I keď prvotný pohľad bol čiste biologicko-matematický, moderný rozvoj systémov AI a dosah ich použitia v rôznych oblastiach života priniesol veľké množstvo otázok, ktoré sa dotýkajú *koexistencie* AI a človeka, resp. spoločnosti. Ide teda o otázky psychologické, sociologické a etické. Inak povedané, **ak chceme riešiť pokročilé systémy umelej inteligencie a ich použitie, nie je možné stavať len na matematických a biologických základoch, ale musíme si uvedomiť, že akákoľvek inteligencia – ak sa má takto nazývať – tieto základy bytostne presahuje.**

1.2. Definícia pojmov

Už debata o termíne umelá inteligencia naznačuje otázky, ktoré si od čias Dartmouthského seminára bádatelia v oblasti AI znovu a znovu kladú: Koľko nám toho ešte zostáva zdolať do vytvorenia skutočne inteligentného stroja? Bude vytvorenie plnej umelej inteligencie vyžadovať reverzné inžinierstvo ľudského mozgu v celej jeho komplexite, alebo nájdeme skratku, nejakú super šikovnú množinu dnes ešte neznámych algoritmov, ktoré sa budú podieľať na vytvorení plnej umelej inteligencie? A čo je podstatné – **vieme vôbec definovať inteligenciu a čo to vlastne znamená „plná inteligencia“?**

Voltaireova výzva „definujte svoje termíny ... lebo inak nikdy jeden druhému neporozumieme“ je výzvou pre kohokoľvek, kto v oblasti umelej inteligencie tvorí, alebo o nej hovorí.⁶⁹ Pretože už základné pojmy ako inteligencia, myslenie, poznanie, vedomie a city sú pre potreby bádania v oblasti umelej inteligencie zle definované. Marvin Minsky

s. 2.

I napriek úžasnému pokroku – osobitne v posledných dekádach - žiadna z týchto predpovedí sa doteraz nenaplnila...

68 Štyria z účastníkov, John McCarthy, Marvin Minsky, Allen Newell a Herbert Simon sa nazývajú „veľká štvorka“ zakladateľov.

69 Súčasnosť dáva tomu za pravdu, keď sa spojenie umelá inteligencia používa ako buzzword – módny termín – zahŕňajúci všetko možné od inteligentných asistentov a expertných systémov, cez jednoduché algoritmy a aplikácie umelej inteligencie, autonómne a robotické systémy, až po sofistikované systémy typu AlphaGO a IBM Watson.

preto pre tieto pojmy razí termín „suitcase word“.⁷⁰ Každý z týchto pojmov je ako veľký kufor, obsahujúci spleť rôznych významov a možností. Preto i umelá inteligencia participuje na tomto probléme a nadobúda tak rôzne významy podľa konkrétnych kontextov.

Ako ľudia tak nejako podvedome vieme rozlíšiť, čo je a čo nie je inteligentné (napr. človek vs. kameň na ceste,...) a vieme takto aj inteligenciu porovnávať (napr. človek vs. bocian,...). Dokonca máme vytvorené aj merítka pre ľudskú inteligenciu (IQ) a navyše vieme rozlišovať jej rôzne dimenzie: emocionálnu, verbálnu, logickú, sociálnu, atď.

Takže inteligenciu kategorizujeme binárne (niečo je alebo nie je inteligentné), kontinuálne (jeden objekt je inteligentnejší ako druhý) a multidimenzionálne (inteligencia v rôznych oblastiach). Pojem inteligencia teda môžeme v Minského poňatí vnímať ako *na prasknutie naplnený kufor*.

Doterajší vývoj systémov AI však ukazuje, že uvedené rozlišovanie v oblasti AI je vo veľkej miere ignorované a dôraz je skôr kladený na dva rôzne smery vývoja – vedecký a praktický.

Na strane vedeckého vývoja sa výskumníci snažia pochopiť prirodzený, t.j. biologický mechanizmus inteligencie a tento následne implementovať v počítačových systémoch.

Praktický smer vývoja sa snaží skôr jednoduchou a pragmatickou cestou vytvoriť také počítačové systémy, ktoré dokážu vykonávať úlohy rovnako alebo lepšie ako ľudia, pričom pramálo záleží na tom, či tieto programy myslia spôsobom ľudského myslenia.

Správa z roku 2016 stanfordskej 100-ročnej štúdie o umelej inteligencii⁷¹ **definuje umelú inteligenciu ako oblasť počítačových vied, ktorá študuje vlastnosti inteligencie syntetizovaním inteligencie.**⁷²

70 MINSKY, M. L. *The Emotion Machine: Commonsense Thinking, Artificial Intelligence, and the Future of the Human Mind*. New York: Simon & Schuster, 2006, s. 95.

71 Stanfordská 100-ročná štúdia o umelej inteligencii (AI100) je projekt, ktorý sa začal v decembri 2014 s ambíciou pravidelne vyhodnocovať AI panelom odborníkov raz za 5 rokov. Ide o aktivitu zameranú na pravidelnú analýzu súčasného stavu rozvoja umelej inteligencie a jej vplyvu na ľudí.

72 SIMON, H. A. *Artificial Intelligence: An Empirical Science*. In *Artificial Intelligence* 77. 1995, č. 2, s. 95-127.

One Hundred Year Study on Artificial Intelligence. [on-line]. AI100, 2016, s. 13. [cit. 14. novembra 2021]. Dostupné na internete: <<https://ai100.stanford.edu/2016-report>>

100-ročná štúdia o umelej inteligencii (AI100) v správe z roku 2021 dopĺňa alternatívnu definíciu AI: **umelá inteligencia je o vytváraní systému, ktorý vykazuje také správanie, o ktorom si myslíme, že vyžaduje inteligenciu.**⁷³

Aktuálna neexistencia presnej definície môže byť aj výhodou pre akcelerovanie celého odvetvia umelej inteligencie v rôznych smeroch vývoja a aplikovania systémov AI. Neexistencia presnej definície AI a viac menej ignorancia základných kategorizácií sú živnou pôdou pre veľmi rôznorodý a kreatívny rozvoj tejto oblasti, keďže bádateľov vedie len hrubý zmysel pre toto smerovanie a snaha etablovať sa, či dostať sa do tejto oblasti.⁷⁴

1.3. Základné vlastnosti a delenie systémov umelej inteligencie

Na základe doterajšej debaty ohľadom definície a kategorizácie inteligencie – systémy, ktoré by boli schopné vykazovať inteligentné správanie, musia spĺňať dve základné vlastnosti:

- **autonómnosť** – schopnosť samostatne konať, t.j. schopnosť systému vykonávať úlohy v komplexnom prostredí bez neustáleho vedenia používateľom.
- **adaptívnosť** – schopnosť sa prispôsobovať, t.j. schopnosť zlepšovať svoj výkon (a schopnosti) učením sa zo skúseností.

Z predchádzajúcej kapitoly so pripomeňme, že o inteligencii môžeme hovoriť v kategóriách kontinuity (jeden objekt je inteligentnejší ako druhý) a viacerých dimenzií (inteligencia v rôznych oblastiach).

Preto i systémy môžu byť adaptívne a autonómne len v určitej oblasti, t.j. sú schopné riešiť určité úlohy „inteligentným spôsobom“, pričom v ostatných oblastiach zlyhávajú. Takéto systémy nazývame **úzko špecializované umelé inteligencie (narrow AI)**⁷⁵. Ide teda o vysoko špecializované systémy, ktoré sú optimalizované na zvládnutie konkrétnej úlohy, resp. množiny úloh.⁷⁶

73 *One Hundred Year Study on Artificial Intelligence*. [on-line]. AI100, 2021, s. 78. [cit. 15. novembra 2021]. Dostupné na internete: <<https://ai100.stanford.edu/2021-report>>

74 MITCHELL, *Artificial Intelligence*, s. 20.

One Hundred Year Study on Artificial Intelligence. [on-line]. AI100, 2016, s. 12. [cit. 14. novembra 2021]. Dostupné na internete: <<https://ai100.stanford.edu/2016-report>>

75 Používa sa skratka **ANI – Artificial Narrow Intelligence**.

76 IBM Deep Blue napr. dokáže poraziť najlepších šachistov sveta, no ak by tento systém mal byť použitý

Všeobecná umelá inteligencia (general AI)⁷⁷ je systém, ktorý dokáže zvládnuť akúkoľvek intelektuálnu úlohu. V podstate ide o umelú inteligenciu, ktorá je na úrovni človeka.⁷⁸

Všetky metódy a systémy umelej inteligencie, ktoré dnes používame, spadajú do kategórie špecializovaných systémov AI (narrow AI), pričom súčasný rozvoj v tejto oblasti napreduje míľovými krokmi.

Všeobecná umelá inteligencia – i napriek bombastickým titulkom v novinách a niektorým futurologickým predpovediam – patrí v súčasnosti do oblasti sci-fi.

Analogicky rozlišujeme umelú inteligenciu ako *silnú* a *slabú* na základe rozlišovania medzi inteligentnými systémami a systémami, ktoré inteligentne konajú.⁷⁹

Silná umelá inteligencia (strong AI) je skutočne inteligentná a „uvedomelá“, t.j. skutočne aj rozumie tomu, či rieši a vykonáva. Uvedomelá AI je súčasne aj všeobecnou, pretože má schopnosť generalizovať, t.j. zovšeobecňovať a prenášať, či adaptovať naučené schopnosti na iné úlohy (čo mimochodom patrí k základom ľudského myslenia).

Slabá umelá inteligencia (weak AI) vykazuje inteligentné správanie na základe modelov a aplikovaných metód i dát, na ktorých sa učí (je natrénovaná). Ide teda o systémy, ktoré sú zamerané na riešenie konkrétnych úloh a sú závislé na ľudskom vstupe a konfigurácii.⁸⁰ Slabá AI reprezentuje systémy, ktoré aktuálne máme, t.j. počítačové systémy, ktoré

v autonómnych vozidlách, nedokázal by sa ani rozbehnúť na rovnej ceste.

77 Používa sa skratka **AGI – Artificial General Intelligence**.

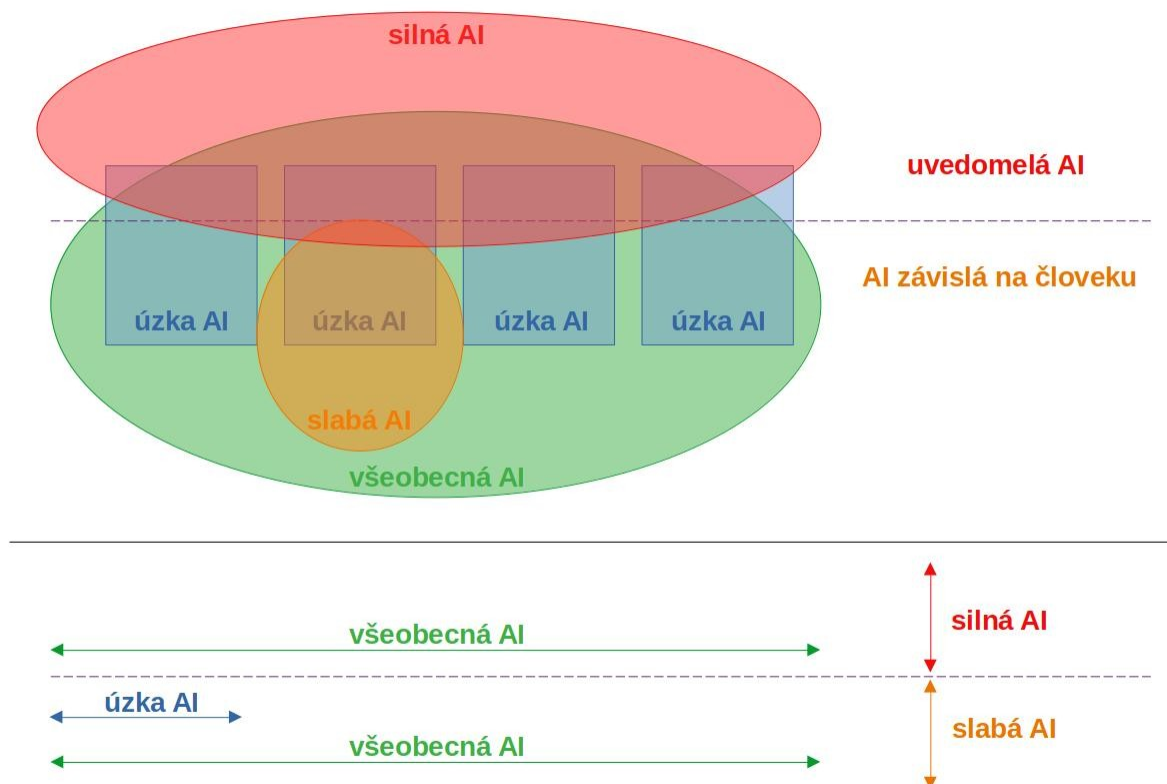
78 Niekedy sa z okruhu AGI samostatne vyčleňuje ešte **ASI – Artificial Super Intelligence**, t.j. umelá inteligencia, ktorá by bola naprieč všetkými oblasťami inteligentnejšia ako človek. Úvahy o ASI sa mnohokrát spájajú s konceptom *singularity v oblasti AI*, v ktorej umelá inteligencia so schopnosťou učiť sa, zlepšovať sa a samostatne sa vyvíjať rýchlo dosiahne a následne prevýši inteligenciu človeka. Viac o vízii Raya Kurzweila a o singularite je uvedené v poznámke č. 16.

79 Tento rozdiel akcentuje zvýraznený odstavec a poznámky o základnom probléme v kapitole 1.1.: „Základným problémom – a jeho podstata sa týka prakticky takmer všetkých systémov umelej inteligencie dneška – je skutočnosť, že i keď nejaký systém dokáže odpovedať tak, akoby konverzácií rozumel, neznamená to, že je inteligentný a že konverzácií aj skutočne rozumie.“

80 Vytvorenie skutočne dobre fungujúceho systému AI vyžaduje skutočných odborníkov, schopných aplikovať správne metódy i správne nakonfigurovať a parametrizovať daný systém. Pre neznalých to môže vyzeráť ako mágia, či voodoo:-) Por. MITCHELL, *Artificial Intelligence*, s. 98.

vykazujú inteligentné správanie.

Nasledujúci obrázok vyjadruje dve znázornenia vzájomného vzťahu a pomeru medzi silnou a slabou, všeobecnou a úzko špecializovanou umelou inteligenciou.



Obr. 1. Vzťah medzi silnou a slabou, všeobecnou a úzko špecializovanou umelou inteligenciou.

1.4. Metódy umelej inteligencie – ich anarchia i filozofia rozdelenia

Ak sme trošku medzi riadkami čítali v kapitole 1.1. o úsvite umelej inteligencie a osobitne o Dartmouthskom seminári, môžeme tušiť, že už od počiatkov výskumu existoval nespočet pohľadov na to, akým spôsobom treba uchopiť problematiku umelej inteligencie a akými cestami by sa jej výskum a vývoj mal uberať. Ani nasledovné desaťročia na rozsahu tejto spleťosti neubrali. A preto, vzhľadom na našu nedostatočnosť chápania inteligencie ako takej⁸¹ a z nej prameniacej neschopnosti uchopiť vytvorenie všeobecnej umelej inteligencie, miesto jasnej cesty vývoja AI sa vývoj a následný pokrok dosahuje v takej

⁸¹ Pojednávali sme o tom v kapitole 1.2.

rozmanitosti a spletnosti metód AI, že môžeme hovoriť o anarchii metód.⁸²

Výskum umelej inteligencie tak obsahuje širokú množinu prístupov, ktorá – na prvý pohľad – môže byť pre vývoj systémov AI brzdou. Máme však za to, že v kontexte našej nedostatočnosti chápania inteligencie ako takej, je tento prístup asi jediný možný. Na jednej strane sú tak preskúmané smery a metódy, ktoré k cieľu nevedú, na strane druhej súčasné úspechy v oblasti rozvoja AI stavajú na schopnosti kombinovať rôzne metódy a prístupy,⁸³ prípadne využiť staršie, javiace sa ako neúspešné, avšak zdokonalené natoľko, že prinášajú kýžené ovocie.⁸⁴

I napriek uvedenej anarchii však môžeme aspoň principiálne – filozoficky rozdeliť túto širokú množinu metód do dvoch za nimi stojacich prístupov:

- symbolická AI
- subsymbolická AI

Symbolická umelá inteligencia ide cestou vytvárania umelej inteligencie na báze ľudského myslenia, t.j. pojmov, slov, fráz (= symboly) a vzťahov medzi nimi. Symbolické systémy na základe definovaných pravidiel a postupov („ak niečo, tak potom toto“) môžu jednotlivé symboly spracovávať a vykonávať priradené úlohy.⁸⁵ Pre symbolickú AI sa význam jednotlivých symbolov odvíja od spôsobu ich kombinácie, vzájomných vzťahov a operácií, ktoré môžu byť nad nimi vykonávané.⁸⁶

Počas prvých dvoch či troch dekád výskumu umelej inteligencie symbolické systémy prevládali. Jednou z prvých implementácií symbolickej AI bol General Problem Solver

82 LEHMAN, J., CLUNE, J., RISI, S.: *An Anarchy of Methods: Current Trends*. In: *How Intelligence Is Abstracted in AI*. IEEE Intelligent Systems. 2014, 29, č. 6, s. 56-62.

83 Napr. realizácia konvolučných sietí.

84 Jeden konkrétny príklad z praxe: „The return of patch-based self-supervision! It never worked well and you had to bend over backwards with ResNets (I tried). Now with ViT, very simple patch-based self-supervised pre-training rocks! First BeIT, now Masked AutoEncoders i1k=87.8%“ [on-line]. Twitter. [cit. 18. novembra 2021].

Dostupné na internete: <<https://twitter.com/giffmana/status/1459092079020285976>>

HE, K., CHEN, X., XIE, S., LI, Y., DOLLÁR, P., GIRSHICK, R.: *Masked Autoencoders Are Scalable Vision Learners* [on-line]. Facebook AI Research (FAIR). [cit. 18. novembra 2021].

Dostupné na internete: <<https://arxiv.org/abs/2111.06377>>

85 MITCHELL, *Artificial Intelligence*, s. 21.

86 MITCHELL, *Artificial Intelligence*, s. 23.

(GPS)⁸⁷, ktorý vedel riešiť také problémy, ako napr. hlavolam Misionári a kanibali⁸⁸. Jedným z vrcholov symbolických systémov sú expertné systémy,⁸⁹ ktoré využívali (a dodnes využívajú) človekom vytvorené pravidlá pre riešenie úloh spojených s lekárskou diagnostikou či právnymi rozhodnutiami, alebo zadania v spravodajských službách⁹⁰ atď.

Pri snahe riešiť zložitejšie úlohy, ktoré mali napr. nepresné vstupy, či vyžadovali rozpoznávanie objektov v reálnych obrazových scénach s vizuálnym šumom a pod., však symbolické systémy zlyhávali.

Popritom treba uviesť, že symbolické systémy umelej inteligencie stále existujú, v niektorých oblastiach AI sa stále využívajú a ovocie ich doterajšieho vývoja, resp. kombinácia s rôznymi algoritmiami subsymbolických metód je zvažovaná ako reálny vklad do diskusie o vývoji nových systémov ANI a AGI.⁹¹

Subsymbolická umelá inteligencia a jej vznik boli inšpirované pokrokom v neurovede. Subsymbolický prístup k AI sa snaží uchopiť naše myšlienkové procesy, ktoré by sme mohli nazvať niekedy nevedomými, či automatickými, a ktoré sú základom tzv. rýchleho vnímania (fast perception), čo využívame napr. pri rozpoznávaní tvárí alebo identifikácii hovorených slov.

Subsymbolické programy umelej inteligencie tak neobsahujú súbor presných softvérových postupov na úrovni logického myslenia, ale sú tvorené len stohom rovníc, pre nezainteresovaného len húštinou ťažko interpretovateľných operácií s číslami.

Subsymbolické systémy sú navrhnuté tak, aby sa učili vykonávať úlohy na základe dát.⁹²

Väčšina moderných implementácií systémov umelej inteligencie, medzi ktoré patria

87 NEWELL, A., SIMON, H. A. *GPS: A Program That Simulates Human Thought*. Santa Monica, California: P-2257, Rand Corporation, 1961.

88 *Missionaries and cannibals problem* [on-line]. [cit. 18. novembra 2021].
Dostupné na internete: <https://en.wikipedia.org/wiki/Missionaries_and_cannibals_problem>

89 I keď pre autora už v tých časoch boli neurónové siete a subsymbolická AI oveľa prítiažlivejšia, na konci deväťdesiatych rokov bol súčasťou tímov, ktoré pre reálne použitie zvažovali nasadenie expertných systémov ako v tom čase jasnej a reálnej voľby.

90 Využívajú sa napr. systémy na inteligentné vyhľadávanie a hlbokú analýzu dát (data mining) alebo na rozpoznávanie objektov (supervector machine).

91 Por. MITCHELL, *Artificial Intelligence*, s. 24.

92 Por. MITCHELL, *Artificial Intelligence*, s. 24.

systémy strojového učenia a neurónové siete, vychádza zo subsymbolického prístupu.

Symbolická AI – veľmi zjednodušene povedané – sa pomocou matematickej logiky snaží emulovať procesy myslenia⁹³, subsymbolická AI – znovu zjednodušujúc – sa snaží emulovať činnosť mozgu na úrovni neurónov.

1.5. Systémy inšpirované činnosťou mozgu na úrovni neurónov

Ako sme v predchádzajúcej kapitole uviedli, väčšina moderných implementácií systémov umelej inteligencie vychádza zo subsymbolického prístupu, ktorý je inšpirovaný činnosťou mozgu na úrovni neurónov.

Tento prístup si môžeme objasniť na jednom z prvých príkladov subsymbolického, mozgom inšpirovaného programu AI, ktorý bol na sklonku 50-tych rokov minulého storočia pod názvom **perceptron** vyvinutý psychológom Frankom Rosenblattom.⁹⁴

I keď by sa mohlo zdať, že v dnešnej dobe ide z pohľadu počítačovej vedy o príklad z praveku, perceptron bol nielen míľnikom v oblasti umelej inteligencie, ale predovšetkým principiálne ovplyvnil vývoj subsymbolických systémov a bol starým otcom moderných a úspešných súčasných systémov AI, ktoré poznáme ako hlboké neurónové siete.

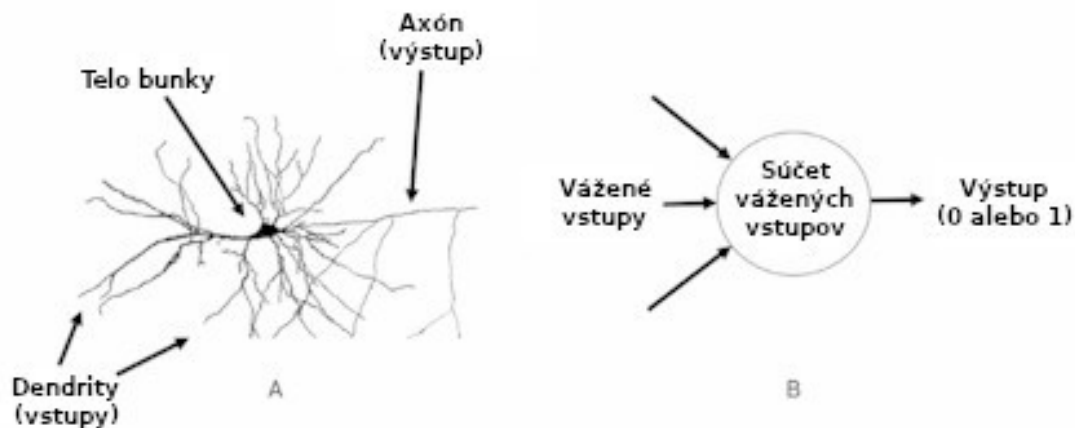
Fungovanie perceptronu bolo principiálne inšpirované spôsobom, ako neuróny spracúvajú informácie: neurón ako mozgová bunka spracúva elektrické, resp. chemické vstupy (vzruchy) z iných neurónov, s ktorými je prepojená. Veľmi zjednodušene povedané, neurón zráta všetky vstupy z ostatných neurónov a ak celková hodnota presiahne určitú hraničnú hodnotu, neurón vyšle impulz. Pritom je dôležité, že jednotlivé spojenia od ostatných neurónov (synapsy) sú rozlične silné. Keď následne neurón ráta výsledný sumár, vstupy zo silnejších synáps majú väčšiu váhu než vstupy zo synáps slabších. Podľa neurovedcov práve úprava sily jednotlivých spojení medzi neurónmi je jedným z kľúčových aspektov mechanizmu učenia sa mozgu.

Perceptron je v podstate počítačovou simuláciou vyššie uvedeného procesu spracovávania informácií v neurónoch. Túto simuláciu vyjadruje obrázok č. 2., pričom časť A zobrazuje neurón s označenými rozvetvenými vláknami, ktoré privádzajú vstupy do bunky (dendrity) a výstupným kanálom (axon). Časť B je schematickým náčrtom

⁹³ Procesy, ktoré u ľudí zahŕňajú zmyslové vnímanie, abstrahovanie pojmov, vytváranie súdov a úsudkov.

⁹⁴ ROSENBLATT, F. *The Perceptron: A Probabilistic Model of Information Storage and Organization in the Brain*. In: *Psychological Review*. 1958, 65, č. 6, s. 386-408.

jednoduchého perceptronu, ktorý sčítava vstupy a v prípade, že výsledný súčet je rovný alebo väčší ako nastavená prahová hodnota, výstup je 1 (analogické vyslaniu impulzu u neurónu), v opačnom prípade je hodnota výstupu 0.



Obr. 2: A – neurón v mozgu, B – jednoduchý perceptron.⁹⁵

V schematickom náčrte perceptronu ako simulácie neurónovej bunky je slovným spojením „vážené vstupy“ vyjadrený podstatný aspekt, bez ktorého nie je možné simulovať mechanizmus učenia sa mozgu – a to vyjadrenie sily jednotlivých vstupných dendritov a ich úprava ako súčasť tohoto mechanizmu učenia sa, t.j. vyjadrenie váhy jednotlivých vstupov perceptronu a proces ich modifikácie ako súčasť učenia sa (adaptácie) systému umelej inteligencie.

Možno si ani neuvedomujeme, ako tento jednoduchý koncept využívame pri niektorých svojich rozhodnutiach. Čo nás napríklad môže viesť k zhliadnutiu konkrétneho filmu v kine? Referencie od priateľov, ktorí tento film už videli, pričom ich názor určite nedávame na rovnakú úroveň – úsudku a vkusu konkrétneho priateľa dôverujeme viac, inému menej. Vytvára sa nám tak množina vážených vstupov, pri ktorých ak pozitívne naladenie pre daný film presiahne určitú úroveň, rozhodneme sa pre jeho zhliadnutie v kine, či nákup v on-line službe. Samozrejme, prichádzajú do úvahy aj ďalšie vplyvy (tzv. hyperparametre), ako napr. reklama na film (*doplnková* konštanta), konkrétne preferencie pre filmový žáner a hercov (variabilita prahovej konštanty, t.j. úrovne, ktorá ak je presiahnutá, rozhodneme sa pre zhliadnutie či nákup filmu) a pod.⁹⁶ A pri výbere ďalšieho

⁹⁵ MITCHELL, *Artificial Intelligence*, s. 25, upravené autorom.

⁹⁶ Inšpiroval som sa príkladom prof. Melanie Mitchell, ktorý som rozviedol o ďalšie aspekty systémov AI.

filmu na zhliadnutie úplne podvedome prichádza k zmenám v dôvere k úsudku jednotlivých našich priateľov – na základe toho, do akej miery bola ich rada pri predchádzajúcom filme relevantná a korešpondujúca s dojmom, ktorý v nás zhliadnutý film zanechal (automatický mechanizmus na zmenu váhy jednotlivých „vstupov“).

Analogicky k uvedenému príkladu bol jednoduchý koncept perceptronu postupne rozvinutý do podoby, v ktorej sú v systémoch AI implementované samoučiace sa mechanizmy na zmenu váhy stupov, aplikované doplnkové konštanty k súčtu vážených vstupov a variability prahovej hodnoty, atď.

Na rozdiel od systémov symbolickej AI, ktorej procesy „inteligencie“ boli výsledkom presných pravidiel a vzťahov medzi symbolmi (pojmy, frázy, slová,...), **u subsymbolických systémov je všetka ich „inteligencia“ zakódovaná v číslach, ktoré reprezentujú váhy a prahové hodnoty.**⁹⁷ Ide teda o stoh rovníc (viď predchádzajúca kapitola), ktoré fungujú na základe správneho nastavenia váh a prahových hodnôt. Tento proces môžeme nazvať učením, pričom systémy sa učia na vzorkách učebných (trénovacích) dát.

Realizácia systémov umelej inteligencie na základe automatického (samostatného) učenia sa z existujúcich vzoriek, dát, prípadne skúseností⁹⁸ sa nazýva strojovým učením (machine learning). Výsledkom je nachádzanie vzorov vo vstupných dátach a také nastavenie váh a prahových hodnôt, ktoré systému umožní poskytovať relevantné výsledky a rozhodnutia, to znamená adekvátne reagovať na rôzne vstupné hodnoty bez toho, aby bol na ne explicitne naprogramovaný, iba na základe informácií, ktoré sa už naučil.^{99 100}

Analogicky k perceptronu, **simulácia viacerých poprepájaných neurónov tvorí**

Por. MITCHELL, *Artificial Intelligence*, s. 25.

97 Základom pre moderné neurónové siete boli tzv. konekcionistické siete (usporiadanie a prepojenie viacerých „neurónov“), v ktorých termín *connectionist* vyjadruje ideu, že „poznatie“ v týchto sieťach sa nachádza vo vážených spojeniach medzi jednotlivými uzlami.

RUMELHART, D. E., MCCLELLAND, J. L. *Parallel Distributed Processing*. Bradford Book, 1986, zv. 1/2.

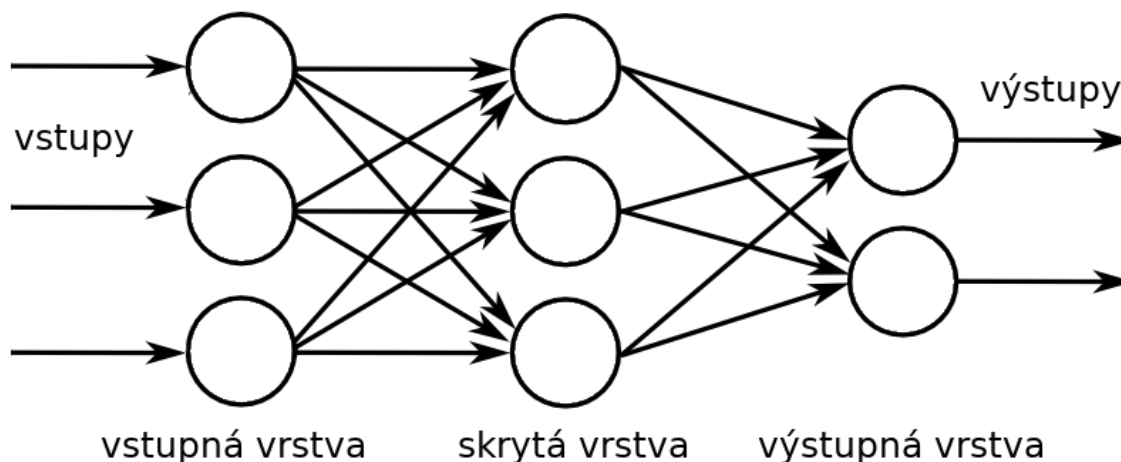
98 Pod skúsenosťou môžeme rozumieť napr. dáta zo senzorov robotického systému, ktorý práve narazil na prekážku a pod.

99 Napr. správna detekcia písaného textu, označenie prekážky na ceste, identifikácia hovoreného slova,...

100 Por. *Algoritmy strojového učenia I.* [on-line]. [cit. 3. januára 2022].

Dostupné na internete: <<https://umelainteligencia.sk/algoritmy-strojoveho-ucenia/>>

neurónová sieť, pričom jednotlivé simulácie neurónov – uzly (node/unit) sú radené do vrstiev (layer). Vstupné dáta sú postupne spracúvané jednotlivými – skrytými vrstvami (hidden layers), ktoré obsahujú skryté uzly. Skryté vrstvy môže predchádzať ešte samostatná vstupná vrstva (input layer). Posledná vrstva, ktorej výstupom je výsledok činnosti neurónovej siete, sa nazýva výstupná (output layer). Uzly medzi jednotlivými vrstvami sú navzájom poprepájané váženými spojeniami.



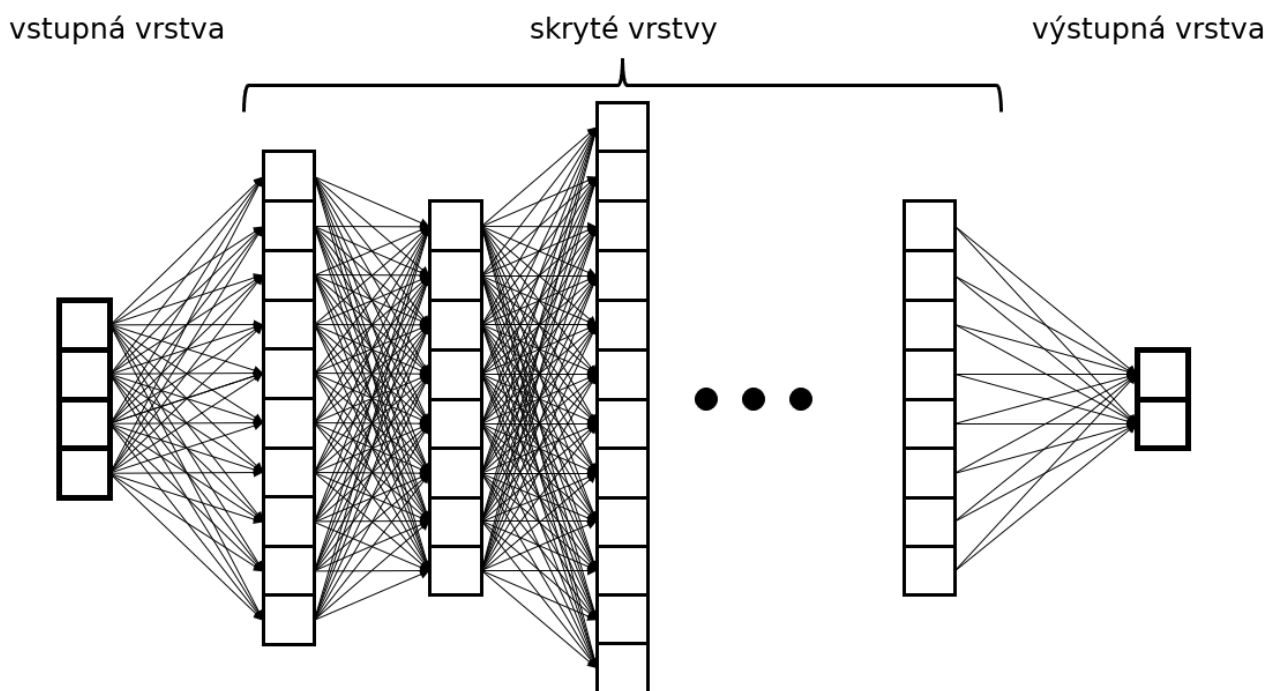
Obr. 3: Viacvrstvová neurónová sieť.¹⁰¹

Ak má neurónová sieť viac než jednu vrstvu, nazýva sa viacvrstvová (multilayered network). Jej príklad je uvedený na obr. č. 3.

Sieť, ktorá obsahuje viac ako jednu skrytú vrstvu, sa nazýva **hlbokou neurónovou sieťou** (deep neural network, prípadne len deep network), ako je uvedené na obr. č. 4.

Súčasný sofistikovaný a najvýkonnejší systém umelej inteligencie využíva typ strojového učenia, ktorý sa nazýva **hlboké učenie** (deep learning). Algoritmy hlbokého učenia sa učia tréningom na obrovskom množstve dát prostredníctvom skrytých vrstiev prepojených uzlov, ktoré – ako sme uviedli – sa označujú ako hlboké neurónové siete.

¹⁰¹ Wikimedia.org, licencia CC, upravené autorom.



Obr. 4: Hlboká neurónová sieť.¹⁰²

1.6. Algoritmy strojového učenia

Podľa toho, ako konkrétne proces učenia prebieha, môžeme algoritmy systémov subsymbolickej AI, resp. strojového učenia rozdeliť do nasledovných skupín:

- učenie s učiteľom (supervised machine learning)
- učenie bez učiteľa (unsupervised machine learning)
- učenie formou odmeňovania (reinforcement learning)

1.6.1. Učenie s učiteľom (supervised machine learning)

Systémy AI využívajúce učenie s učiteľom tvoria v súčasnosti väčšinu systémov strojového učenia. Základom je dostatočne veľká množina správne označených/klasifikovaných tréningových dát (labeled training dataset) s pozitívnymi a negatívnymi príkladmi. Napr. ak chceme naučiť systém rozpoznávať rôznym spôsobom napísanú číslicu 5, musíme mať veľkú kolekciu napísaných 5-iek, pri ktorých je uvedené označenie (label), že ide o číslicu 5, a tiež veľkú kolekciu iných písaných číslic s označením, že to nie je číslica 5. Tieto kolekcie tvoria tzv. tréningové dáta (training set/dataset). Systém v procese učenia

¹⁰² Wikimedia.org, licencia CC, upravené autorom.

vyhodnocuje každú jednu predloženú číslicu, pričom výsledok sa porovná s jej označením (je – nie je to 5). V prípade nezhody, sa generuje signál (supervision signal), označujúci nesprávny výsledok, resp. mieru nezhody. Na základe takýchto signálov si systém mení nastavenia váh a prahových hodnôt, až kým sa nedostane do stavu, že všetky vstupné dáta z tréningovej množiny nie sú správne vyhodnotené (v našom prípade všetky číslice 5 vyhodnotené ako 5 a ostatné vyhodnotené ako niečo iné).¹⁰³

Ide samozrejme o schematický opis fungovania množiny algoritmov pre supervised machine learning, ktorých presný popis nie je cieľom tejto práce.

1.6.2. Učenie bez učiteľa (unsupervised machine learning)

Učenie bez učiteľa sa používa pri dátach, ktoré neboli vopred klasifikované. Chýba nám teda označenie (label) – správna odpoveď, ktorá klasifikuje vstupné dáta (napr. ide o číslicu 5, na vstupe je trojuholník,...).¹⁰⁴

Systémy s učením bez učiteľa sa využívajú v prípade, že na vstupe nemáme klasifikované dáta, alebo ide o dáta, ktoré nepoznáme, takže nevieme natrénovať systém na základe dát, ktoré by sme poznali a vedeli ich klasifikovať, resp. označiť. Základné algoritmy využívajúce učenie bez učiteľa teda nie sú schopné určovať správny výstup – ich cieľom je skôr skúmanie a analýza vstupných dát v snahe objaviť a popísať vzory a štruktúry v týchto dátach, teda dozvedieť sa o dátach, resp. z týchto dát niečo viac.

Učenie bez učiteľa sa využíva na riešenie úloh zhlukovania a asociovania.

Cieľom zhlukovania je spájanie dát do skupín, ktoré majú niečo spoločné, napr. segmentácia zákazníkov do skupín s podobnými preferenciami, niektoré problémy spracovania obrazu a detekcie objektov a pod.

Asociovanie sa využíva pri vyhľadávaní asociačných pravidiel, ktoré popisujú množiny dát

103 Cyklus prijatia informácie na vstupe – jej spracovania – vyhodnotenia výsledku – následná zmena váhových stavov sa pri zložitejších neurónových sieťach nazýva epochou trénovania systému AI.

Pri trénovaní systém prechádza mnohými epochami, ktorých ovocím sú meniace sa váhové stavy i parametre systému a zlepšujúca sa klasifikácia, t.j. lepšie výsledky. Nakoniec systém „konverguje“, t.j. medzi jednotlivými epochami sa váhové stavy takmer vôbec nemenia a neurónová sieť je v princípe v rozsahu trénovacích dát vytrénovaná (naučená).

Por. MITCHELL, *Artificial Intelligence*, s. 80.

104 Por. *Algoritmy strojového učenia II.* [on-line]. [cit. 3. januára 2022].

Dostupné na internete: <<https://umelainteligencia.sk/algoritmy-strojoveho-ucenia-ii-ucenie-bez-ucitela/>>

a vyjadrujú vzťahy medzi nimi. Napr. predajné systémy, ktoré po natrénovaní vedia zákazníkom kupujúcim si určitý produkt ponúknuť relevantné ďalšie produkty (napr. ponuka periférií pri kúpe notebooku), riešenie niektorých problémov pri jazykových prekladoch a pod.

1.6.3. Učenie formou odmeňovania (reinforcement learning)

Ide o osobitnú kategóriu algoritmov strojového učenia, v ktorom sa tréning modelu (agenta) realizuje prostredníctvom interakcie s prostredím metódou pokus-omyl. Základom sú pravidlá, podľa ktorých sa agent môže v danom prostredí správať a odmeňovacia funkcia, prostredníctvom ktorej agent vie vyhodnotiť, či vykonané rozhodnutie bolo pre neho správne alebo nie.¹⁰⁵

Učenie agenta teda neprebieha na žiadnych označených, či neoznačených tréningových dátach, systém v danom prostredí skúša jednotlivé možnosti a učí sa, ktorá možnosť, resp. kombinácia možností je správna a ktorá nie.

Samozrejme ide len o vyjadrenie princípu – jeho realizácie vo forme algoritmov, ako sú napr. Q-learning či Deep Q-learning sú oveľa sofistikovanejšie a využívajú ich napr. šachové systémy, systémy pre pohyb robotických súprav v prostredí a pod.

Pre správne vytrénovanie systému učenie formou odmeňovania treba zvyčajne extrémne veľa iterácií (napr. milióny pokusov).

Tolko aspoň v krátkosti k jednotlivým skupinám algoritmov strojového učenia (detailnejšie rozdelenie je uvedené v prílohách), na ktoré sa teraz môžeme pozrieť trochu z nadhľadu...

1.7. Stačí strojové učenie, alebo hľadáme ďalej?

Prakticky prvým z algoritmov strojového učenia a ich podmnožiny učenia sa s učiteľom bol Rosenblattov perceptron-learning algorithm. Rosenblatt tiež matematicky dokázal, že pre veľmi špecifickú a limitovanú množinu úloh dokáže správne natrénovaný perceptron (resp. podobný systém učenia sa s učiteľom) vykonávať úlohy bez chýb. Avšak pri iných, resp. všeobecnejších úlohách umelej inteligencie je úspešnosť takéhoto systému

105 Por. *Algoritmy strojového učenia III*. [on-line]. [cit. 4. januára 2022].

Dostupné na internete: <<https://umelainteligencia.sk/algoritmy-strojoveho-ucenia-iii-ucenie-formou-odmenovania/>>

Už pri tak jednoduchej realizácii algoritmu strojového učenia je možné si uvedomiť viaceré riziká spojené so systémami umelej inteligencie, napríklad:¹⁰⁹

- nesprávna zmena váh a prahových hodnôt v procese učenia (napr. priveľké zmeny váhových hodnôt medzi jednotlivými krokmi učenia) môže viesť k nesprávnemu natrénovaniu systému (systém bude dávať nesprávne výsledky v prevádzke).
- u niektorých skupín učenia je pre realizáciu konkrétnych modelov treba vykonať nesmierne veľa iterácií v procese učenia. Obmedzenie možností učenia vedie k modelom, ktoré dávajú nesprávne, resp. v určitých situáciách nesprávne výsledky.
- obmedzenie na množinu úloh, ktoré dokáže takýto systém riešiť.
- praktická nemožnosť previesť naučené hodnoty váh a prahových hodnôt do ľudske pochopiteľných pravidiel, systém sa tak stáva „čiernou skrinkou“.

Čím zložitejší systém AI (s miliónmi váhových stavov a prahových hodnôt, s viacerými vrstvami neurónových sietí, atď.), tým sú uvedené problémy vypuklejšie. Vystáva tak dôležitá otázka – **sme potom schopní overovať a nastavovať etické pravidlá, resp. všeobecne obmedzenia a mantinely systémov umelej inteligencie?**¹¹⁰

Obhajcovia subsymbolických systémov umelej inteligencie – argumentujúc „architektúrou“ mozgu na úrovni neurónov – sú presvedčení, že prostriedkami symbolickej AI (symboly

106 Por. MITCHELL, *Artificial Intelligence*, s. 30.

107 MINSKY, M. L., PAPERT, S. L. *Perceptrons: An Introduction to Computational Geometry*. Cambridge, Mass: MIT Press, 1969.

108 Zložitejšie úlohy samozrejme vyžadujú použitie viacvrstvových neurónových sietí, pre ktoré však v tom čase neexistoval učiaci sa algoritmus, t.j. spôsob, ako vo viacerých vrstvách realizovať spätnú väzbu a nastavovanie váhových stavov v procese učenia. Preto Minsky dokonca zavrhol celý koncept subsymbolickej AI. Až neskorší vývoj všeobecného algoritmu učenia sa, tzv. **back-propagation** algorithm znamenal skutočný prelom vo vývoji neurónových sietí, pričom v spojení s hlbokým učením sa tieto stali základom pre aktuálny veľký pokrok v oblasti umelej inteligencie.

NIELSEN, M. *Neural Networks and Deep Learning*. [on-line]. [cit. 19. januára 2022].

Dostupné na internete: <<http://neuralnetworksanddeeplearning.com/>>

109 Komplexnej analýze limitov a rizík súčasných systémov umelej inteligencie sa venujeme v 2. kapitole.

110 Nie je to problém pri systéme na rozpoznávanie číslíc, ale môže to byť skutočný problém pri automatických bojových systémoch, autopilotov vozidiel, prostriedkov na podporu života, systémov na detekciu teroristických hrozieb, podporných systémov pre súdnictvo, atď.

a logické vzťahy medzi nimi) nie je možné dosiahnuť umelú inteligenciu, keďže symboly v mozgu vznikajú na základe neurálnych procesov (na úrovni neurónov).¹¹¹

I napriek všetkým úspechom, ktoré sme v súčasnosti vďaka subsymbolickým systémom, medzi ktoré patria neurónové siete, systémy hlbokého učenia atď., dosiahli, **dovoľíme si s týmto názorom nesúhlasiť**, keďže stále nie sme schopní prekročiť hranice ANI, t.j. hranicu medzi úzko špecializovanými systémami (slabej) umelej inteligencie a všeobecnými systémami AGI.¹¹² Radi by sme uviedli aspoň dve výhrady...

Domnievame sa, že **ak budujeme systémy umelej inteligencie inšpirujúc sa ľudským mozgom, resp. sa chceme priblížiť analógii s mozgom, je treba vziať do úvahy metódy subsymbolickej i symbolickej AI** (napr. ako „vstup“ pre systém symbolickej AI vziať „výstupy“ zo subsymbolického systému a pod.). Na jednej strane by tak išlo o postupy, ktoré sú jednou z možností ďalšieho rozvoja umelej inteligencie,¹¹³ na druhej strane – čo je pre nás podstatné – by išlo o systémy, v ktorých na úrovni symbolickej AI by sme azda boli schopní ľahšie riešiť etické výzvy a implementovať pravidlá aj v prípade komplexných systémov AI.¹¹⁴

Druhou výhradou je fakt, že i napriek veľkej snahe dnešné hlboké neurónové siete a algoritmy hlbokého učenia (teda dnešné subsymbolické systémy) len čiastočne simulujú mozgovú činnosť a prácu neurónov. **Biologické neuróny sa správajú inak ako uzly neurónových sietí a mozgové procesy sú v mnohom komplexnejšie.**

Uzol napodobňujúci činnosť neurónu vysiela impulz, resp. vo všeobecnosti číselné stavy. Výstupom biologického neurónu je však elektrický prúd spôsobený pohybom iónov v bunke. Jednotlivé výstupné prúdy sa prostredníctvom synáps posúvajú do ďalších buniek, kde zvyšujú ich membránový potenciál, ktorý je výsledkom nerovnovážnej koncentrácie iónov. Keď membránový potenciál neurónu prekročí určitú prahovú hodnotu, bunka vyšle impulz (používa sa termín spike), t. j. prúd, ktorý sa má odovzdať do ďalších

111 Por. MITCHELL, *Artificial Intelligence*, s. 31.

112 Informácie o ANI a AGI sú uvedené v kapitole 1.3.

113 V súčasnosti sa pokroky v moderných systémoch AI dosahujú (niekedy až radikálnym) vylepšovaním existujúcich algoritmov a súčasne kombináciou viacerých systémov, metód a postupov na dosiahnutie cieľa (viď poznámky č. 84 a 85).

114 Treba však uviesť, že doterajšie pokusy vytvoriť hybridný systém integrujúci subsymbolické i symbolické metódy AI nemali nejaký väčší úspech.
MITCHELL, *Artificial Intelligence*, s. 41.

buniek.¹¹⁵ V klasických neurónových sieťach je tento proces simulovaný výstupným binárnym impulzom a váhovým stavom „synapsy“ spájajúcej dva uzly.

Biologické neuróny majú navyše vnútornú dynamiku, ktorá spôsobuje ich zmeny v čase. S plynutím času majú tendenciu vybiť sa a znižovať svoj membránový potenciál, čoho dôsledkom je napríklad skutočnosť, že riedke impulzy na vstupe neurónu nespôsobia výstupný impulz!¹¹⁶

Dôležitým rozdielom je i spôsob komunikácie medzi biologickými neurónmi a medzi uzlami. Biologické neuróny komunikujú a spracúvajú informácie asynchrónne, kým typické uzly neurónových sietí synchronne, t.j. v jednom kroku všetky uzly jednej vrstvy neurónovej siete prečítajú vstup, vyrátajú výstup a posúvajú ho ďalej. U biologických neurónov je to iné – v každom okamihu môžu prijať vstupný signál a vytvoriť výstup bez ohľadu na správanie sa ostatných neurónov.¹¹⁷

Tiež treba poznamenať, že biologické systémy majú oveľa prepracovanejší systém spätnej väzby na úrovni neurónov (čo je súčasťou procesov učenia sa a adaptívnej činnosti), než sú dnešné metódy používané v neurónových sieťach (napr. back-propagation algoritmus spomínaný na začiatku tejto kapitoly).

Na jednej strane tak máme súčasné sofistikované neurónové siete, ktoré majú rastúce, ba až neúnosné výpočtové náklady (a tým aj vysokú spotrebu energie, časovú náročnosť, a pod.) potrebné na ich kvalitné vytrénovanie a rozvoj.¹¹⁸ Tieto siete vedia byť veľmi efektívne v konkrétnych scenároch použitia, no ak by mali vedieť riešiť pre človeka

115 *Action potential* [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Action_potential>

116 KORAKOVOUNIS, D. *Spiking Neural Networks: where neuroscience meets artificial intelligence*. [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <<https://theaisummer.com/spiking-neural-networks/>>

117 KORAKOVOUNIS, *Spiking Neural Networks: where neuroscience meets artificial intelligence*. [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <<https://theaisummer.com/spiking-neural-networks/>>

118 „The cost of improvement is becoming unsustainable.“

THOMPSON, N. C., GREENEWALD, K., LEE, K., MANSO, G. F. *Deep learning computational cost*. [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <<https://spectrum.ieee.org/deep-learning-computational-cost>>

jednoduché základné úlohy¹¹⁹, resp. doterajšie zadania zovšeobecňovať a nebodaj aj chápať, sú v koncoch.

Na strane druhej máme prirodzenú inteligenciu s nepatrnou spotrebou energie, schopnú kreativity, riešenia problémov a multitaskingu. Ide o biologické systémy, ktoré si prirodzenou evolúciou osvojili spracovanie informácií a reagovanie na ne.^{120 121}

Ovocie pokračujúceho skúmania mozgových procesov a neurónov, snahy pochopiť, čo ich robí tak efektívnymi a odvaha aplikovať tieto zistenia do oblasti umelej inteligencie viedli k vzniku tretej generácie neurónových sietí, tzv. **spiking neural networks** (SNNs) a ich implementácie v rámci neuromorfnej hardvérovej architektúry.^{122 123}

Vrátiac sa k začiatku tejto diskusie, t.j. rozporovaniu tvrdenia, že súčasné algoritmy strojového učenia sú schopné dosiahnuť skutočnú umelú inteligenciu, vidíme, že je pred nami ešte dlhá cesta a vývoj, ktorý v kontexte doterajších subsymbolických i symbolických prístupov môže priniesť veľa zaujímavých poznatkov a progresu na ceste (nielen) k uvedomelej umelej inteligencii.

1.8. Striedanie ročných období

Skúsme spraviť malú retrospektívu vývoja umelej inteligencie. Roky po Dartmouthskom seminári ubiehali a veľmi optimistické predpovede o pokroku vo vývoji systémov umelej

119 „The easy things are hard.“ – táto problematika je rozoberaná v kapitole 1.9.

120 KORAKOVOUNIS, *Spiking Neural Networks: where neuroscience meets artificial intelligence*. [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <<https://theaisummer.com/spiking-neural-networks/>>

121 V našom kontexte sú tieto systémy i súčasťou etického rámca, ktorý biológii presahuje.

122 PFEIFFER, M., PFEIL, T. *Deep Learning With Spiking Neurons: Opportunities and Challenges*. [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <<https://www.frontiersin.org/articles/10.3389/fnins.2018.00774/full>>

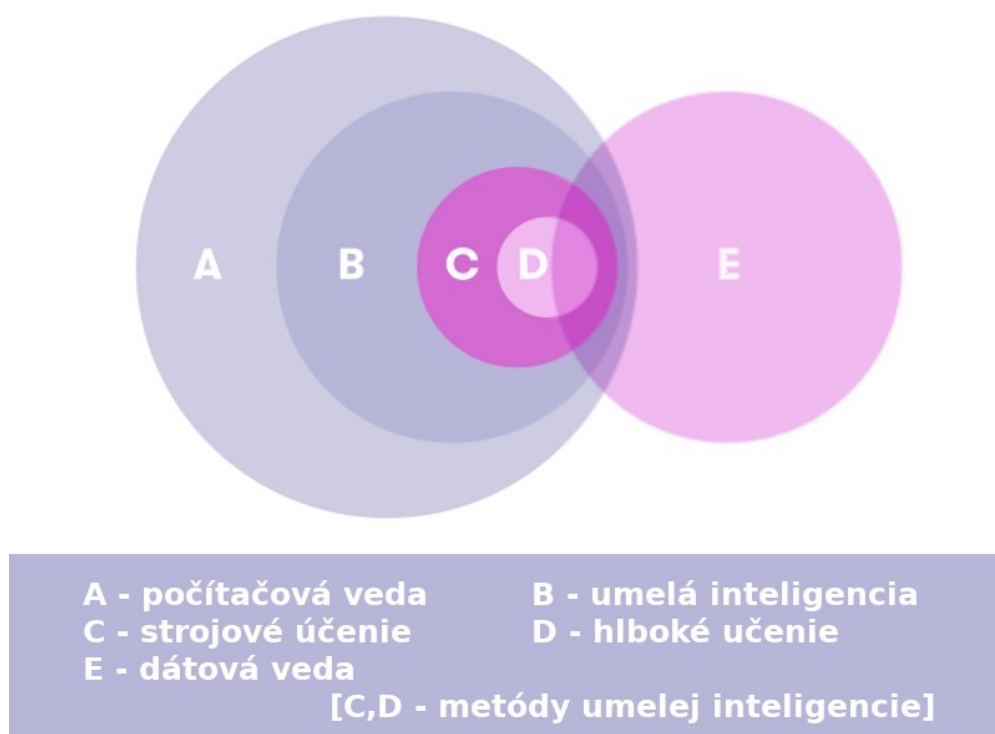
123 Predchádzajúce generácie neurónových sietí sú založené na McCullochových Pittsových neurónoch (t.j. prahových hradlách), resp. sigmoidálnych hradlách, pričom sa javí, že jeden spike neurón dokáže riešiť takú konkrétnu biologickú funkcionálnu, ktorá by si vyžadovala stovky skrytých jednotiek sigmoidálnej neurónovej siete.

MAASS, W. *Networks of spiking neurons: The third generation of neural network models*. [on-line]. [cit. 28. februára 2022].

Dostupné na internete: <<https://www.sciencedirect.com/science/article/abs/pii/S0893608097000117>>

inteligencie sa nenapĺňali.¹²⁴ S vyčerpaním potenciálu konkrétnych dobových poznatkov v oblasti AI a schopností dobovej techniky tak veľmi rýchlo prichádzalo k strate nadšenia i utlmeniu investícií do tejto oblasti.

Vývoj bol udržiavaný v zásade len v rámci základného výskumu viacerých univerzitných centier a pokroku v príbuzných oblastiach (napr. robotika a kybernetika, výpočtová technika, dátová a počítačová veda, v rámci priemyslu a vojenského vývoja), z čoho ťažila aj veľmi vágne ohraničená oblasť umelej inteligencie, ako sme uviedli na konci kapitoly 1.2. Vzťah medzi dátovou vedou, počítačovou vedou a umelou inteligenciou ako vedným odborom vyjadruje schematický obrázok č. 5.



Obr. 5: Vzťah medzi dátovou vedou, počítačovou vedou a umelou inteligenciou.

I napriek útlmu ovocím základného výskumu a pokroku v príbuzných oblastiach bol posun v hľadaní nových metód a prístupov k riešeniu problémov umelej inteligencie. Každý väčší posun sa následne stával štartérom nového entuziazmu a záujmu o umelú inteligenciu, prichádzali nové pozitívne predikcie, ktorých ovocím bol i ďalší prísun investícií a podpory pre túto oblasť. A po vyčerpaní potenciálu daného posunu znovu nastával útlm...

Toto – až do poslednej dekády viac menej pravidelné – **striedanie aktívneho rozvoja**

124 Viac o optimistických predpovediach a predikciách hlavných protagonistov je uvedené v poznámke č.

a útľmu v oblasti umelej inteligencie sa nazýva aj striedanie ročných období:

- jar umelej inteligencie (AI spring) – obdobie prekotného rozvoja, ktoré šartuje novými ideami a pokrokom vo výskume, čoho ovocím sú predikcie prevratného pokroku vo vývoji systémov AI častokrát sprevádzané mediálnym boomom a následným prílevom štátnych dotácií i rizikového kapitálu do akademického bádania i komerčných startupov.
- zima umelej inteligencie (AI winter) – prevratný pokrok neprichádza, výsledky síce sú, no vôbec nespĺňajú extrémne očakávania. Prílev financií a kapitálu ustáva, startupy krachujú a akademický výskum sa spomaľuje a obmedzuje.

Uvedený vzor sa vo vývoji umelej inteligencie v rozličných podobách opakoval v päť až desaťročných cykloch.¹²⁵

S veľkým pokrokom vo vývoji algoritmov vychádzajúcich zo štatistických a pravdepodobnostných teórií, ktoré umožnili rozvoj strojového učenia, v súčasnosti postupne prišlo – až na niekoľko špecifických oblastí – k odmietnutiu symbolických systémov AI¹²⁶ a veľkému nástupu neurónových sietí a strojového učenia. V posledných dvoch desaťročiach tak striedanie ročných období reálne prebieha už len v rámci vlastných cyklov vývoja subsymbolických systémov AI, keďže viackrát sa ukázalo, že na riešenie skutočných problémov reálneho sveta sú aj moderné systémy strojového učenia prikrátke...

V poslednej dekáde možno diskutovať o ďalšom vývoji umelej inteligencie i z pohľadu týchto vývojových cyklov.

Umelá inteligencia v súčasnosti zažíva nebývalý rozmach, dynamický vývoj, extrémny rozsah praktického aplikovania jej systémov, špičkové zabezpečenie výskumu a vývoja, spoločenskú akceptáciu i prehnané mediálne pokrytie a očakávania. Vyzerá to skoro ako definícia jarnej fázy AI.

Na druhej strane však treba povedať, že aktuálne veľké pokroky nie sú dôsledkom

125 Por. MITCHELL, *Artificial Intelligence*, s. 34.

Autorka dokonca spomína, že v čase ukončenia jej vysokoškolského štúdia (1990) sa vývoj AI nachádzal tak hlboko „v zimnom období“, až je radili, aby do profesného životopisu v rámci žiadostí o zamestnanie umelú inteligenciu vôbec neuvádzala.

126 Odmietnutie bolo často sprevádzané dešpektom zhmotneným do označenia symbolickej AI ako GOFAI (podľa Douglasa Hofstadtera ide o skratku pre good old old-fashioned AI).

revolučných zmien v AI, ale len evolučného vývoja existujúcich metód a ich sofistikovaných kombinácií, technologickej podpory z iných oblastí (hlavne rozvoj informačných a komunikačných technológií), dostupnosť extrémne veľkých datasetov, atď. Prevratný pokrok však v podstatných aspektoch neprichádza (vytvorenie AGI). To niekomu skôr pripomenie prvé zimné mrazy...

Keďže miera akceptácie a adaptácie systémov AI v modernej spoločnosti presiahla hranicu, za ktorou je ťažké si predstaviť návrat späť (resp. zmenu priorít rozvoja vedomostnej spoločnosti a Industry 4.0) a keďže trochu plochá krivka základného výskumu a vývoja systémov AI je súčasťou permanentného technologického pokroku, ktorý celkovo v oblasti rozvoja vedy, techniky, objemu dát a znalostí dosahuje takmer exponenciálny rast,¹²⁷ dovoľme si predpokladať, že na aktuálny a budúci vývoj v oblasti umelej inteligencie už nebude možné v plnej miere aplikovať striedanie jarých a zimných období. Výskum a vývoj systémov AI bude neustále pod tlakom prakticky všetkých spoločenských oblastí, v ktorých sa AI v súčasnosti využíva.

Pozorný čitateľ určite vníma, že spochybňujúc pokračovanie vývojových cyklov sme skĺzli do oblasti ANI, t.j. k úzko špecializovaným systémom (slabej) umelej inteligencie, ktorým sa v súčasnosti neskutočne darí.

Inak to však môže byť s AGI, t.j. skutočnou umelou inteligenciou, ktorá sa stáva definitívnou metou súčasného základného výskumu v oblasti AI.¹²⁸ Tu si vývoj nedovoľme predpovedať, keďže na jednej strane principiálny prelom vo vývoji metód umelej inteligencie nenastáva, na strane druhej však vylepšovanie a kombinácia súčasných techník, extrémne veľké datasety, výpočtové systémy s výkonom dizajnovaným priamo pre aplikácie umelej inteligencie i všeobecný vedomostný a výskumný potenciál môžu veľmi ľahko viesť aj k prípadnej skokovej zmene a vývojovému prelomu.

127 ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi*, [online], s. 20, 25, 31, 34. [cit. 7. januára 2022]. Dostupné na internete:

<https://peter.santavy.cloud/docs/Analyza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

128 Filozofi umelej inteligencie, Vincent Müller a Nick Bostrom zverejnili v roku 2013 prieskum vykonaný medzi výskumníkmi AI, v ktorom mnohí vyjadrili 50% šancu na dosiahnutie umelej inteligencie na úrovni človeka do roku 2040.

MÜLLER, V. C., BOSTROM, N. *Future Progress in Artificial Intelligence: A Survey of Expert Opinion*. In: *Fundamental Issues of Artificial Intelligence*. Cham. Switzerland: Springer International, 2016, 555-72.

V rámci viac než šesťdesiatich rokov bádania a snáh o vytvorenie všeobecnej umelej inteligencie mnohí vedci a výskumníci viackrát takmer opustili ideál AGI. V súčasnosti – keď sa debata o umelej inteligencii stala súčasťou hlavného spoločenského prúdu – sa mnohým tento ideál javí ako samozrejмый. Predpokladáme, že len niekoľko rokov bude stačiť, aby sme videli, ktoré tendencie prevažujú: či optimizmus pretrváva, alebo fenomén AGI sa v rámci pokračujúceho výskumu bude javiť ako oveľa sofistikovanejší, komplexnejší a náročnejší nielen na realizáciu, ale hlavne na uchopenie a pochopenie.

1.9. Jednoduché veci sú ťažké

Nadväzujúc na predchádzajúcu kapitolu treba povedať, že v určitej miere nepôjde o nič nové, keďže súčasťou každej zimy umelej inteligencie bolo zistenie, akým problémom pre systémy AI je schopnosť realizácie niektorých vecí. Päťdesiat rokov od konferencie v Dartmouthe to John McCarthy vyjadril slovami: „AI bola ťažšia, ako sme si mysleli“.¹²⁹ A Marvin Minsky už oveľa skôr poznamenáva, že výskum v oblasti umelej inteligencie odkryl **dôležitý paradox: „Jednoduché veci sú ťažké“¹³⁰, keďže riešenie úloh, ktoré sa človeku zdajú byť jednoduché, je pre umelú inteligenciu veľmi náročné. A naopak to, čo je pre človeka extrémne náročné, dokážu systémy AI zvládať veľmi dobre.**¹³¹

Pôvodné ciele vývoja umelej inteligencie – počítače, ktoré dokážu s nami komunikovať ľudskou rečou, vedia popísať, čo „vidia“ kamerou, „chápu“ koncepty na základe len niekoľkých príkladov – to všetko sú veci, ktoré malé dieťa poľahky zvláda, no pre AI sú oveľa ťažšie, ako napr. komplexná diagnostika choroby, víťazstvo v šachu alebo v Go nad najlepšími ľudskými šampiónmi alebo riešenie sofistikovaného algebraického problému.^{132 133}

129 MOEWES, C., NÜRNBERGER, A. *Computational Intelligence in Intelligent Data Analysis*. New York: Springer, 2013, s. 135.

130 MINSKY, M. *Society of Mind*. New York: Simon & Schuster, 1986, s. 29.

131 Kognitívny vedec Steven Pinker to vyjadruje celé: „**The hard things are easy, but the easy things are hard.**“

TROTT, D. *The hard things are easy, but the easy things are hard*. [on-line]. [cit. 8. januára 2022].

Dostupné na internete: <<https://www.campaignlive.com/article/hard-things-easy-easy-things-hard/1498154>>

132 MITCHELL, *Artificial Intelligence*, s. 33-34.

133 Ide o tzv. Moravcov paradox, ktorý bol jedným z výsledkov výskumu Hansa Moravca, Rodney Brooksa a Marvina Minského v osemdesiatych rokoch.

Minský to vyjadril slovami: „Vo všeobecnosti si najmenej uvedomujeme to, čo naša myseľ robí najlepšie.“¹³⁴ A Sparsh Chadha dodáva: „Stroje ... dokážu vykonávať len úlohy, na ktoré boli vytrénované. Vďaka našej schopnosti dynamického myslenia a zdravému rozumu výrazne prevyšujeme stroje. Bez ohľadu na to, ako veľmi natrénujeme naše modely, na koľkých tréningových cykloch trénujeme naše stroje, stále existuje priestor neistoty, ktorý by sa dal človekom veľmi jednoducho vyriešiť, ale umelá inteligencia by zlyhala, pretože jej chýba schopnosť zdravého rozumu.“¹³⁵

Ide o skúsenosť sprevádzajúcu celé desaťročia vývoja umelej inteligencie, ktorá nám pripomína, ako komplexná a subtílna je ľudská myseľ a čo všetko je vo vývoji AGI ešte pred nami.

Jednou z podmienok zvládnutia pre človeka jednoduchých, ale pre AI náročných vecí je prechod od systémov, ktoré sa javia ako inteligentné (napr. systém AI na jazykový preklad, systém na popis objektov zosnímanej scény a pod.) **k systémom, ktoré sú inteligentné** (AI systém, ktorý rozumie prekladanému textu, systém, ktorý chápe obsah zosnímanej scény a pod.).¹³⁶ Avšak rozumieť a chápať súvislosti - to sa znovu dostávame do oblasti všeobecnej a silnej umelej inteligencie (AGI).

Prakticky vo všetkých strategických oblastiach rozvoja umelej inteligencie je AGI túženým cieľom, keďže jej schopnosti by umožnili riešiť problémy na úplne inej úrovni. V kombinácii s technologickými možnosťami súčasnej informatizácie a automatizácie sa možnosti AGI javia ako úžasné.

Napr. analýza cestnej premávky na základe všetkých dostupných senzorov autonómneho vozidla – teda schopnosť ľudského komplexného vnímania, ktorá by bola spojená s vizuálnymi, zvukovými, radarovými, lidarovými, atď. senzormi a zároveň by umožňovala extrémne rýchle reakcie v priamej interakcii s elektronickým ovládaním vozidla.¹³⁷

MORAVEC, H. *Mind Children: The Future of Robot and Human Intelligence*. Cambridge, Mass.: Harvard University Press, 1988.

134 MINSKY, *Society of Mind*, s. 29.

135 CHADHA, S. *“Common Sense” is the Dark Matter of Artificial Intelligence*. [on-line]. [cit. 4. februára 2022].

Dostupné na internete: <<https://hackernoon.com/the-dark-matter-of-ai-common-sense-is-not-so-common>>

136 Pozri poznámku č. 56 a súvisiaci text v kapitole 1.1.

137 Súčasný autonómny systém sú mnohokrát veľmi pomalé, ak sa majú rozhodovať napr. v situácii

Tieto predpokladané úžasné možnosti však idú ruka v ruke s výzvami, ktoré sa doteraz v oblasti umelej inteligencie neriešili – etický rozmer činnosti strojov: ako ho definovať, ako ho realizovať, ako fixovať požadované vlastnosti a ako ich aplikovať do predmetu činnosti daného systému AI.¹³⁸

1.10. Transhumanizmus a umelá inteligencia

Koketovanie s myšlienkou AGI a túžba po vytvorení uvedomelej umelej inteligencie je často konfrontovaná s porovnávaním s inteligenciou človeka. V optike problémov, ktoré vývoj AGI obnáša a v obavách z disproporcie medzi uvedomelou umelou a ľudskou inteligenciou, mnoho vedcov a futuroológov sa zaoberá myšlienkami transhumanizmu, t.j. futurologickým a filozofickým konceptom, ktorý pojednáva o možnostiach transformácie človeka prostredníctvom využitia moderných technológií. Je to pohľad do budúcnosti, ktorý je založený na premise, že ľudia vo svojej súčasnej podobe nereprezentujú koniec vývoja, ale iba jeho ranú fázu.¹³⁹

V zásade ide o dva základné prieniky medzi umelou inteligenciou a transhumanizmom:

- integrácia systémov AI a človeka
- existenciálne riziko vyplývajúce z pokročilej umelej inteligencie

Zavádzanie prvkov umelej inteligencie obnáša aj asistenčné technológie, ktoré zvyšujú

zložitej križovatky.

138 Etika systémov umelej inteligencie je v súčasnosti pomerne zavedená oblasť výskumu AI. U AGI však neide o etiku vývoja, nasadenia, používaných dát, atď., ale priamo o etický rozmer činnosti uvedomelej AI.

139 Transhumanizmus je futurologický a filozofický koncept, ktorý sa zaoberá možnosťami transformácie človeka prostredníctvom využitia moderných technológií. Ide o filozofické a intelektuálne hnutie, ktoré obhajuje zlepšenie ľudského stavu vývojom a sprístupnením sofistikovaných technológií, ktoré dokážu výrazne zvýšiť dĺžku a kvalitu života, kognitívne schopnosti. Transhumanistickí myslitelia skúmajú potenciálne výhody a nebezpečenstvá vznikajúcich technológií, ktoré by mohli prekonať základné ľudské obmedzenia, ako aj etiku používania takýchto technológií.

MERCER, C., TROTHEN, T. J. *Religion and Transhumanism: The Unknown Future of Human Enhancement*. Praeger, 2014.

Niektorí transhumanisti sa domnievajú, že ľudia sa nakoniec budú môcť premeniť na bytosti so schopnosťami tak výrazne rozšírenými oproti súčasnému stavu, že si zaslúžia označenie post-humánne bytosti (pojednáva o tom samostatný komplexný futuristický smer označovaný ako posthumanizmus).

kvalitu života, od tých jednoduchých, napr. implementácia prvkov rozšírenej reality (augmented reality) v rámci rôznych nositeľných zariadení (wearables), až po oveľa sofistikovanejšie, napr. pripojenie robotických končatín na nervovú sústavu hendikepovaného človeka¹⁴⁰. V tomto kontexte sa pre časť výskumu javí atraktívnou myšlienka prepojiť vhodné systémy z umelej inteligencie s telom a nervovou sústavou človeka. Osobitne, ak na jednej strane dosahujeme pomerne veľké úspechy vo vývoji ANI, no súčasne sa len málo posúvame k dosiahnutiu méty AGI.

Keďže jednou z tém a oblastí transhumanistického výskumu je snaha ochrániť ľudstvo pred existenčnými rizikami,¹⁴¹ ako je napríklad jadrová vojna alebo zrážka s asteroidom, smerovanie k všeobecnej a uvedomelej umelej inteligencii spojené s rizikami, ktoré jej potenciál môže obnášať, viedli transhumanistickú scénu aj k vnímaniu existenciálnych rizík vyplývajúcich z pokročilej umelej inteligencie.¹⁴²

Ak by sme len v úzkom kontexte tejto práce chceli uchopiť etické výzvy spojené s transhumanizmom¹⁴³, išlo by buď o problémy súvisiace s vplyvom virtuálneho sveta na život človeka a spoločnosti, ktoré sme popísali v 2. kapitole licenciátskej práce *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi*¹⁴⁴, teda

140 Elon Musk v rámci svojich vizionárskych aktivít (Neuralink Corporation) vytvára rozhranie medzi mozgom a počítačom pre rozšírenie možností človeka, ako napr. integrácia robotických končatín, alebo prepojenie ľudského mozgu a umelej inteligencie pre dosiahnutie symbiózy, ktorú by sme poľahky mohli zaradiť do oblasti transhumanizmu.

Neuralink [on-line]. [cit. 5. augusta 2020].

Dostupné na internete: <<https://en.wikipedia.org/wiki/Neuralink>>

141 „Elon Musk už dlho hovorí, že umelá inteligencia bude musieť ľudské schopnosti skôr rozširovať, než s nimi súťažiť, aby sa vyhnula hrozivej budúcnosti.“

FERRIS, R. *Elon Musk thinks we will have to use AI this way to avoid a catastrophic future*. [on-line]. [cit. 24. januára 2022].

Dostupné na internete: <<https://www.cnn.com/2017/01/31/elon-musk-thinks-we-will-have-to-use-ai-this-way-to-avoid-a-catastrophic-future.html>>

142 BOSTROM, N. *A history of transhumanist thought*. In: *Journal of Evolution and Technology*. [on-line]. 2005, zv. 14, vyd. 1. [cit. 8. januára 2022].

Dostupné na internete: <<http://www.nickbostrom.com/papers/history.pdf>>

143 Teda etické výzvy spojené s umelou inteligenciou, nie celý rozsah problémov týkajúcich sa transhumanizmu.

144 ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi*, [on-line], s. 25-47. [cit. 8. januára 2022]. Dostupné na internete:

problémy týkajúce sa slabej umelej inteligencie, ktoré sú analogické problémom s ostatnými aspektami informačných technológií, kybernetického priestoru a virtuálneho sveta. Alebo by išlo – lepšie povedané pôjde – o problémy oscilujúce medzi výzvami spojenými s etikou uvedomelej umelej inteligencie a psychológiou, morálnymi aspektami i vierou postmoderného človeka.

Pri masívnom prieniku systémov umelej inteligencie do reálneho života treba teda rátať s vplyvom transhumanistickej filozofie, ako napr. s mravným redukcionizmom a relativizmom, falošnou premisou, že technické „vylepšovanie“ človeka v kontexte umelej inteligencie vedie k jeho kultúrnemu i morálnemu zlepšovaniu¹⁴⁵, materialistickému a technokratickému chápaniu šťastia a naplnenia¹⁴⁶, falošným rovnostárstvom s odovzdaním vlády superpočítačom (~umelej inteligencii)¹⁴⁷ a pokriveným vnímaním eschatologickej roviny viery v kontexte transfigurizmu, t.j. náboženského transhumanizmu.

1.11. Masívne zavádzanie prvkov umelej inteligencie

Základ pre masívne využívanie systémov umelej inteligencie tvorí viacero pilierov, ktoré sú ovocím súčasného rozvoja informačnej spoločnosti:

- pokrok vo vývoji metód a systémov umelej inteligencie, osobitne v oblasti neurónových sietí;
- existencia, jednoduché získavanie a bezprecedentná dostupnosť¹⁴⁸ extrémneho

<https://peter.santavy.cloud/docs/Analiza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

145 STRAHOVNIK, V. *Virtues and transhumanist human enhancement*. In: PETROUŠEK R., ŽALEC B., eds. *Transhumanism as a Challenge for Ethics and Religion*. 2021, s. 37 – 44.

146 *Transhumanist Declaration*. [on-line]. [cit. 26. januára 2022].

Dostupné na internete: <https://hpluspedia.org/wiki/Transhumanist_Declaration>

147 LEE, I. *Equalism: Paradise Regained*. In: LEE N. (ed.). *The Transhumanist Handbook*. Springer, 2019, s. 849 – 863.

148 Už v roku 2010 Erich Smid, v tom čase CEO Google, na konferencii Techonomy v kalifornskom Lake Tahoe vyhlásil: „V súčasnosti každé dva dni vyprodukuje toľko informácií, koľko sme vytvorili od úsvitu civilizácie až do roku 2003.“

SIEGLER, M. G. *Eric Schmidt: Every 2 Days We Create As Much Information As We Did Up To 2003* [on-line]. [cit. 12. decembra 2021].

Dostupné na internete: <<http://techcrunch.com/2010/08/04/schmidt-data/>>

množstva štruktúrovaných i neštruktúrovaných dát;¹⁴⁹

- moderné prístupy, ktoré vyžadujú adaptívne procesy schopné spracovávať také kvantum dát;
- potreba automatizácie – rozvoj automatizácie smerujúci k potrebe autonómnych systémov schopných samostatnej adaptívnej činnosti pri spracúvaní veľkého množstva neštruktúrovaných dát v konkrétnej oblasti.

Tieto piliere sú súčasťou tvoriacej sa a (v niektorých častiach sveta) rozvíjajúcej sa informačnej a znalostnej spoločnosti, ktorá je sprevádzaná paradigmatickou zmenou: informačné technológie a dáta sa z roviny prostriedku dostávajú do roviny kontextu spoločnosti až do tej miery, že na základe zmien v technológiách a vzhľadom na prevratný vedecký i technologický pokrok sa mení medziludská komunikácia vo svojej podstate, menia sa vzťahy, mení sa spoločnosť a princípy jej fungovania.¹⁵⁰

Nástup a rozšírenie systémov umelej inteligencie patrí k tejto zmene paradigmy, pričom si dovoľíme tvrdiť, že smerovanie k AGI a dosiahnutie systémov uvedomelej umelej inteligencie by pre mnohých bolo jej naplnením.

149 Senzorické systémy, počítačové siete, digitalizácia, informatizácia vedy i technologických procesov a predovšetkým on-line priestor komunikácie a sociálnych sietí sú zdrojom veľkého množstva dát. „Big data“, ako súčasť technologickej a informačnej revolúcie, umožnili rozvoj viacerých v súčasnosti dominantných metód umelej inteligencie, keďže vďaka veľkému množstvu dát je možné dostatočne a relevantne trénovať systémy AI, aby boli použiteľné v reálnom nasadení.
Por. MITCHELL, *Artificial Intelligence*, s. 80.

150 ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi*, [online], s. 19-20. [cit. 7. decembra 2021].
Dostupné na internete:
<https://peter.santavy.cloud/docs/Analiza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

2. Limity a riziká súčasných systémov umelej inteligencie

*Každá dostatočne pokročilá technológia je na nerozoznanie od mágie.*¹⁵¹

V súčasnosti sme svedkami rozsiahleho nasadzovania systémov umelej inteligencie do praxe, pričom nejde len o širokospektrálne využitie v oblasti vedy a výskumu, resp. v niektorých strategických oblastiach, ale aj o bežné používanie v spotrebnej elektronike a v systémoch, ktoré sú súčasťou nášho každodenného života. Využitie systémov AI sa navyše stáva tak samozrejmým, že ich bežné používanie si ani neuvedomujeme, avšak ich akceptujeme a pomaly – zvykajúc si na benefity ich nasadenia – ich až vyžadujeme, keďže v mnohom nám prinášajú pridanú hodnotu, ktorej sa nechceme vzdať.¹⁵²

Ak použijeme analógiu z oblasti kybernetickej bezpečnosti, zameranie nastupujúcich generácií na komfort a benefity informačných technológií je mnohokrát silnejšie, než opatrnosť a bezpečné správanie sa v rámci ochrany pred kybernetickými hrozbami, únikom a zneužitím osobných údajov a pod. Zámerne spomínam nastupujúcu generáciu, ktorá sa hrdí familiárnosťou vo využívaní informačných technológií a médií, v rámci vzdelávania dostáva aspoň základy gramotnosti v oblasti informačnej bezpečnosti, ale návykovosť a komfort mnohých technológií v reálnom použití víťazí nad zdravou mierou opatrnosti a aplikovania bezpečnostných zásad.

Je takmer isté, že podobná ľahkovážnosť bude sprevádzať spoločnosť aj pri využívaní systémov umelej inteligencie. Súčasná miera akceptácie jednoduchých systémov AI a spôsob ich využívania dáva tušiť, že používanie sofistikovaných systémov AI môže obnášať riziká, ktoré si ani nevieme predstaviť.

V tejto kapitole sa však venujme problému, o ktorých vieme a treba s nimi rátať pri návrhu, tvorbe a využívaní systémov umelej inteligencie. Stále sa pritom pohybujeme v oblasti ANI, teda úzko špecializovaných systémov umelej inteligencie

¹⁵¹ Arthur C. Clarke.

¹⁵² Napr. pre zákazníkov to môže byť využitie jednoduchých systémov AI vo fotoaparátoch chytrých telefónov – modely, ktoré tieto funkcie nemajú, sú v danej kategórii prakticky nepredateľné.

Pre majiteľov elektronických obchodov sa stali nepostrádateľnými systémy, ktoré dokážu vytvárať profily prichádzajúcich návštevníkov e-shopov a zvýšiť šancu predaja rôznych produktov.

Viacere bankové domy si v súčasnosti nevedia predstaviť niektoré svoje on-line služby bez dohľadu natrénovanej umelej inteligencie, ktorá dokáže odhaľovať podvody a falošné transakcie...

(narrow AI), ktoré sú optimalizované na zvládnutie konkrétnej úlohy, resp. množiny úloh. Ide súčasne o systémy slabej umelej inteligencie (weak AI), ktoré vykazujú inteligentné 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.¹⁵³

Keďže v súčasnosti neexistujú systémy AGI, t.j. uvedomelej silnej a všeobecnej umelej inteligencie, zameranie na ANI je samozrejmé. Navyiac – vzhľadom na aktuálnu absolútnu preferenciu neurónových sietí a strojového učenia – sa sústredíme na ich súčasnú prezentáciu hlbokými neurónovými sieťami (deep neural networks) a hlbokým učením (deep learning). Pokúsime sa poukázať na limity a riziká týchto systémov AI, popísať technologické výzvy v oblasti bezpečnosti procesov umelej inteligencie a s tým súvisiacej kybernetickej bezpečnosti, uvedomiť si extrémnu komplexnosť týchto systémov a akcentovať viaceré dôsledky, ktorých zneužitie môže znamenať vážne ohrozenie pre jednotlivcov i spoločnosť.

2.1. Vybrané limity a rizikové faktory systémov umelej inteligencie

Pri vývoji, realizácii a nasadzovaní prvkov AI je treba rátať s viacerými obmedzeniami a rizikami súčasných systémov umelej inteligencie.

Ponajprv malé prirovnanie – ak vnímame neurónové siete a subsymbolickú AI ako systémy hlboko inšpirované činnosťou mozgu na úrovni neurónov¹⁵⁴, pre zdôraznenie **dôsledkov zlyhania prvkov neurónových sietí pre ich celkovú funkčnosť** môžeme použiť veľmi hrubú analógiu s mozgovými poruchami spôsobenými na bunkovej úrovni (neurón ~ uzol neurónovej siete), ktoré môžu viesť k veľmi vážnym psychickým poruchám.

Napríklad serotonín (5-hydroxytryptamin, skratka 5-HT) je biologicky aktívna molekula, ktorá pôsobí ako dôležitý neurotransmitter prenášajúci vzruchy medzi neurónmi. Aby po uvoľnení z jednej bunky mohol prenášať vzruchy na inú, viaže sa na príslušné povrchové receptory, medzi ktoré patrí aj serotonínový receptor 2A (5-HT_{2A})¹⁵⁵. Nesprávna hladina a funkcia receptorov 5-HT_{2A} súvisí s vážnymi psychickými poruchami, vrátane

153 Viac o úzko špecializovaných a slabých systémoch AI pojednáva kapitola 1.3.

154 O systémoch AI inšpirovaných činnosťou mozgu na úrovni neurónov sme pojednávali v kapitole 1.5.

155 5-HT_{2A} receptor [on-line]. [cit. 20. januára 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/5-HT2A_receptor>

bipolárnej poruchy, ťažkých depresí a schizofrénie. Správnu (zdravú) rovnováhu serotonínových receptorov v mozgu môže poškodiť veľa faktorov, napr. drogy, niektoré choroby, neurologické lieky, ale aj nedostatok spánku.¹⁵⁶

Podobne aj u neurónových sietí je veľa faktorov súvisiacich nielen s ich dizajnom a technologickým riešením, ale aj s konkrétnymi zvolenými parametrami, tréovacím datasetom, vyhodnocovaním a pod., ktoré môžu ovplyvniť ich správnu a bezpečnú funkčnosť.

Túto analógiu akcentuje aj skutočnosť, že súčasné systémy hlbokého učenia sú priamo modelované na základe objavov a poznatkov z neurovedy – inak povedané – moderné hlboké neurónové siete napodobňujú zodpovedajúce štruktúry mozgu. Napríklad v súčasnosti najúspešnejšie systémy AI pre spracovanie obrazu sú hlboké siete, ktorých štruktúra napodobňuje časti vizuálneho systému mozgu.^{157 158}

2.1.1. Neurónová sieť ako čierna skrinka

Jedno zo základných rizík vyplýva zo skutočnosti, že **prakticky nevieme, na základe čoho robia hlboké neurónové siete svoje rozhodnutia.**¹⁵⁹ Vieme, ako nadizajnovať neurónovú sieť pre konkrétnu oblasť použitia. Vieme, ako ju natrénovať a v rámci možností aj otestovať. Keďže však neurónová sieť neobsahuje súbor presných softvérových postupov na úrovni logického myslenia, ale je tvorená len stohom rovníc, len húštinou ťažko interpretovateľných operácií s číslami, ktoré fungujú na základe správneho nastavenia váh, konštánt a prahových hodnôt, **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

156 GREGOROVÁ, D. *Nedostatek spánku působí na mozek* [on-line]. [cit. 20. januára 2022].

Dostupné na internete: <<https://www.osel.cz/12127-nedostatek-spanku-pusobi-na-mozek.html>>

157 MITCHELL, *Artificial Intelligence*, s. 71.

158 Ide o konvolučné neurónové siete (convolutional neural networks, skratky ConvNets alebo CNN), ktoré sú dizajnované podľa vizuálneho cortexu v mozgu. Pomenovanie „konvolučné“ sa odvíja od názvu matematickej operácie vykonávanej nad vstupnými hodnotami recepčných polí v rámci jednotlivých aktivačných máp skrytých vrstiev neurónovej siete.

Por. MITCHELL, *Artificial Intelligence*, s. 72-80.

KOUSHIK, J. *Understanding Convolutional Neural Networks*. [on-line]. [cit. 31. januára 2022].

Dostupné na internete: <<https://arxiv.org/abs/1605.09081>>

159 Por. MITCHELL, *Artificial Intelligence*, s. 39.

pri činnosti systému. Táto miera nevedomosti rastie s mierou komplexnosti neurónovej siete, t.j. počtom skrytých vrstiev a uzlov, použitými algoritmi a prítomnosťou špecifických úprav.¹⁶⁰

Už od školských lavíc poznáme otázku – ukáž, ako si to spravil – na ktorú odpovedáme ozrejméním postupu, akým sme prišli k výsledku. Vyjadrenie postupu tak umožňuje nielen jednu z možností verifikácie správnosti výsledku, ale i overenie postupu, ktorým k výsledku prichádzame. Ide tak o jeden zo základných stavebných kameňov dôvery a istôt, na ktorých stavíme naše konanie, spoznávanie sveta, vedecko-technologickú činnosť i celkové fungovanie v reálnom svete. V súčasnosti moderné systémy strojového učenia na výzvu – ukáž, ako si to spravil, resp. predved' spôsob, ako pracuješ – nevedia jednoducho odpovedať (lepšie povedané, ani autori týchto systémov nevedia vo všeobecnosti objasniť, ako presne ich hlboké siete prichádzajú k výsledkom).¹⁶¹

Uvedený problém sa stáva ešte vypuklejším v rámci debát o implementácii bezpečnostných mechanizmov a etických mantinelov do systémov AI. **V tejto rovine sa sofistikovaný systém umelej inteligencie javí ako *black box* – čierna skrinka, ktorá niečo vykonáva, ale ako a prečo tak robí, nie je jasné.**

V tejto súvislosti môžeme spomenúť ešte jeden aspekt strojového učenia – tzv. hyper

160 Ide o vážny problém, preto existujú viaceré odborné aktivity snažiace sa black box otvoriť a jeho obsah vzniesť na svetlo, t.j. chápať, ako systém AI pracuje. Úspechy sú však len parciálne a v úzko špecifických nasadeniach algoritmov umelej inteligencie.

Napr. systém AI ExoMiner vyvinutý v NASA: „Unlike other exoplanet-detecting machine learning programs, ExoMiner isn't a black box – there is no mystery as to why it decides something is a planet or not,” said Jon Jenkins, exoplanet scientist at NASA's Ames Research Center in California's Silicon Valley. „We can easily explain which features in the data lead ExoMiner to reject or confirm a planet.“ *New Deep Learning Method Adds 301 Planets to Kepler's Total Count*. [on-line]. [cit. 28. novembra 2022].

Dostupné na internete: <<https://www.jpl.nasa.gov/news/new-deep-learning-method-adds-301-planets-to-keplers-total-count>>

161 Odborný časopis Technology Review z MIT (Massachusetts Institute of Technology) v tejto súvislosti uvádza, že ide o „temné tajomstvo v srdci umelej inteligencie“ a dodáva: „nikto v skutočnosti nevie, ako najpokročilejšie algoritmy robia to, čo robia. To by mohol byť problém“.

KNIGHT, W. *The Dark Secret at the Heart of AI*. In: *Technology Review*. [on-line]. 2017, 11. 4. [cit. 10. februára 2022].

Dostupné na internete: <<https://www.technologyreview.com/2017/04/11/5113/the-dark-secret-at-the-heart-of-ai/>>

parametre. Súčasné systémy umelej inteligencie prevažne pracujú na základe strojového učenia, ktoré sme vo všeobecnosti definovali ako proces učenia sa (resp. realizácie systémov AI na základe automatického/samostatného učenia sa) z existujúcich vzoriek, dát, prípadne skúseností.¹⁶² Pre úplnosť však treba dodať, že aby sa systémy AI dokázali pomocou algoritmov strojového učenia niečo naučiť (vytrénovať/nastaviť váhy spojení neurónovej siete), treba ich najprv nadizajnovat', resp. prednastaviť pomocou tzv. hyper parametrov. Termín hyper parametre (hyperparameters) je zastrešujúcim vyjadrením pre celú množinu parametrov, ktoré musia byť človekom prednastavené, aby neurónová sieť bola vôbec schopná úspešne sa učiť. Do tejto množiny patrí napríklad počet vrstiev neurónovej siete, veľkosť recepčných polí jednotlivých uzlov konvolučných sietí, parametre rýchlosti učenia sa, použitá aktivačná funkcia a spôsob klasifikácie i veľa ďalších technických detailov.¹⁶³

Takže v rámci voľby použitých algoritmov a rozhodnutia o komplexnom dizajne neurónovej siete je treba vyladiť (používa sa termín tuning) veľa hyper parametrov, ktorých vzájomná interakcia vedie k správne fungovaniu systému AI. Navyše, vo väčšine prípadov ide o unikátne „namiešanú zmes“ techník a parametrov, ktorú treba „namiešať“ nanovo pre každé nové zadanie, na ktoré bude neurónová sieť trébovaná. **Vyladenie hyper parametrov sa tak stáva absolútne kľúčovým pre správnu funkčnosť systémov AI.**¹⁶⁴

V súčasnosti neexistuje univerzálny návod na vyladenie hyper parametrov v rámci návrhov systémov strojového učenia. Ide o schopnosť, ktorá je ovocím dôkladných znalostí, vytrvalej učiteľivosti a mnohokrát ťažko nadobudnutej skúsenosti.¹⁶⁵ Eric Horvitz, riaditeľ výskumných laboratórií Microsoftu, preto na margo ladenia hyper parametrov a úspešného návrhu systému AI hovorí: „**To, čo v súčasnosti robíme, nie je veda, ale určitý druh**

162 Napr. využitie back-propagation algoritmu v systémoch učenia sa s učiteľom (supervised learning) alebo skúseností, resp. interakcie s prostredím v učení formou odmeňovania (reinforcement learning).

Základné rozdelenie algoritmov strojového učenia sme popísali v kapitole 1.6.

163 Por. MITCHELL, *Artificial Intelligence*, s. 97.

164 Por. MITCHELL, *Artificial Intelligence*, s. 98.

165 Z takto namiešanej zmesi sa (nielen) v mojej praxi kybernetickej bezpečnosti občas rodia konkrétne zásahy na zastavenie prebiehajúcich hackerských útokov, ktoré je niekedy ťažko explicitne popísať. Žargónom IT je to len magic – kúzlo, čím sa vyjadruje nie presný postup riešenia, ale skôr odborná intuícia, ktorá viedla k očakávanému výsledku.

alchýmie.¹⁶⁶ Čo sa týka elitného klubu odborníkov, ktorí toto ladenie dokážu realizovať, Demis Hassabis, spoluzakladateľ Google DeepMind, poznamenáva: „Je to takmer ako umenie, ako z týchto systémov dostať to najlepšie... Na svete je len niekoľko stoviek ľudí, ktorí to dokážu robiť naozaj dobre.“¹⁶⁷

Nech už hovoríme o čiernej skrinke, alchýmii alebo umení, **súčasný systém AI sprevádzajú oprávnené a vážne obavy z toho, že ak nechápeme ako tieto systémy pracujú, nemôžeme im reálne dôverovať a ťažko dokážeme predpovedať okolnosti, za ktorých tieto systémy zlyhajú.**

Toto v prísnom slova zmysle môžeme povedať aj o zlyhaniach človeka, v tomto prípade však v kontexte ľudského spoločenstva stavíme na určitej *teórii mysle*, teda na modeli správania sa ľudských bytostí, ktoré je na základe fundamentálnych kognitívnych schopností (rozpoznávanie objektov, chápanie scény, porozumenie reči, atď.) štandardné a vývojov overené. Avšak pri pokročilých systémoch AI nič také, ako „teóriu mysle“ nemáme.¹⁶⁸

2.1.2. Zraniteľnosti systémov strojového učenia

V rámci niekoľkých dekád vývoja neurónových sietí a systémov strojového učenia boli postupne identifikované viaceré rizikové faktory a zraniteľnosti systémov AI. Uvedme aspoň tie podstatné:

- malá množina trénovacích dát
- nesprávne zvolená, či nekvalitná množina trénovacích dát a predsudky
- nadmerné prispôbovanie sa tréningovým údajom
- efekt dlhého chvosta
- klamanie hlbokých sietí a ich zraniteľnosti
- povery

Malá množina trénovacích dát (training dataset)

Úspešnosť väčšiny súčasných systémov umelej inteligencie je extrémne závislá

166 METZ, C. *A New Way for Machines to See, Taking Shape in Toronto*. New York Times, Nov. 28, 2017.

167 TANZ, J. *Soon We Won't Program Computers. We'll Train Them Like Dogs*. Wired, May 17, 2016.

Ide o vyjadrenie z roku 2016 – za ten čas sa v oblasti AI veľa zmenilo, no podstata zostáva...

168 Por. MITCHELL, *Artificial Intelligence*, s. 109.

na rozsiahlych¹⁶⁹ a kvalitných súboroch správne označených tréningových dát, resp. tréningových iterácií.¹⁷⁰ Bez nich sa súčasné systémy strojového učenia nedajú vytrénovať a ich nedostatok vedie v lepšom prípade k nekvalitným, v tom horšom k nesprávnym výsledkom a fatálnym zlyhaniam.¹⁷¹ A ako zistíme pri ďalších typoch zraniteľností a rizikových faktorov systémov AI, problematika tréningových dát je oveľa širšia...

Nesprávne zvolená, či nekvalitná množina tréningových dát a predsudky (biases)

Na základe nesprávne zvolenej alebo nekvalitnej (napr. nesprávne označkovanej) množiny tréningových dát **sa systém AI naučí robiť chybné uzávery alebo podávať výsledky „s predsudkami“**.¹⁷²

Ide o problém nielen technologický, ale aj sociologický – predsudky, či zaujatosť v spoločnosti vedú k voľbe nesprávnemu obsahu tréningových dát a následne k chybným výsledkom systémov AI, pričom v mnohých prípadoch hrozí, že systémy AI tréňované na zaujatých dátach môžu tieto predsudky násobiť a spôsobiť reálne škody.

Ďalším aspektom „zaujatých“ systémov AI je znížená úspešnosť ich činnosti na základe predsudkov, čo môže mať veľmi nepríjemné dôsledky napr. v bezpečnostných systémoch, ochrane ľudských práv, sociálnej spravodlivosti a pod.

Nesprávne fungovanie zaujatého systému AI častokrát nie je problém detekovať v rámci

169 I táto skutočnosť poukazuje na rozdielnosť medzi hlbokým strojovým učeníom a ľudským vzdelávaním sa.

170 Bavíme sa o fenoméne „big data“, čo napr. pri spracovaní obrazu konvolučnými sieťami znamená mať k dispozícii rádovo milióny správne označených tréningových obrázkov.

Por. MITCHELL, *Artificial Intelligence*, s. 99.

„Requiring so much data is a major limitation of [deep learning] today.“

NG A. *Deep Learning in Practice: Speech Recognition and Beyond*. [on-line]. In: *EmTech Digital*. 2016, 23. máj. [cit. 3. februára 2022].

Dostupné na internete: <<https://events.technologyreview.com/video/watch/andrew-ng-deep-learning/>>

171 Viac o fenoméne „big data“ a veľkej množine tréningových dát je uvedené v kapitole 2.5.

172 Viac krát sa stalo, že systémy AI museli byť vypnuté, lebo po uvedení do prevádzky sa na základe nesprávne nastavených hyper parametrov a zvolenej množiny tréningových dát vyvíjali nesprávnym smerom. Napr.:

KRAFT, *Microsoft shuts down AI chatbot after it turned into a Nazi*, [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://www.cbsnews.com/news/microsoft-shuts-down-ai-chatbot-after-it-turned-into-racist-nazi/>>

HAMILTON, *Amazon built an AI tool to hire people but had to shut it down because it was discriminating against women*, [on-line]. [cit. 6. augusta 2020].

ostrého nasadenia, no nie je triviálne to odhaliť v predstihu (napr. v procese učenia).

V niektorých oblastiach nasadenia umelej inteligencie nie je jednoduché natrénovať systém AI bez predsudkov. Vyžaduje si to mnohokrát veľkú erudovanosť a nasadenie tých, ktorí pripravujú tréningové dáta, pričom aj pre nich platí, že pokiaľ sú súčasťou spoločnosti akceptujúcej predsudky, podvedome ich môžu prenášať aj do svojej práce.¹⁷³ A tu sa dostávame na veľmi tenký ľad, **pretože v mnohých oblastiach sa ako spoločnosť nezhodneme na tom, čo je a čo nie je predsudok** (napr. pohlavie vs. rod s útokmi na tzv. rodové stereotypy a pod.).

Niektoré problémy s predsudkami takmer určite v budúcnosti dokážu vyriešiť systémy AGI (ak budú existovať), ktoré by mali chápať základné koncepty a v rámci abstrakcie eliminovať celú škálu zaujatostí v oblasti strojového učenia. Zvyšok, možno ten podstatný, bude výzvou rovnako, ako boj s predsudkami v spoločnosti.

Nadmerné prispôsobovanie sa tréningovým údajom (overfitting to training data)

Ide o problém, v rámci ktorého **sa systém naučí z tréningových dát rozlišovať niečo iné, než to, čo sa mal naučiť**. Napr. ak sa má systém naučiť rozlišovať nakreslený kruh a v rámci učiaceho procesu tréningové dáta budú obsahovať vzorky, na ktorých je nakreslený kruh vždy len na zelenom podklade, je možné, že systém sa naučí ako kruh identifikovať nie to, čo nakreslený kruh skutočne obsahuje, ale skôr všetko to, čo je zelené. Ak potom pri ostrej prevádzke bude na vstupe napríklad trojuholník na zelenom podklade, systém ho identifikuje ako kruh.

V uvedenom prípade je odhalenie problému veľmi ľahké a náprava jednoduchá. Avšak v reálnom svete (napr. pri detekčných systémoch AI v kvantovej fyzike, onkológii, atď.) to vôbec nemusí byť jednoduché a korelácia, koherentnosť i kauzalita v rámci tréningového procesu vôbec nemusí byť zjavná.

Ešte jedna poznámka k doteraz uvedeným problémom s tréningovými dátami – minimálne v niektorých oblastiach pokročilých systémov sa na dôvažok javí, že ani nadmieru veľká a kvalitná množina v súčasnosti používaných tréningových dát nebude dostatočná pre pokrok v rozvoji týchto systémov a k ich priblíženiu sa k všeobecnej a uvedomelej AGI. Napríklad v oblasti počítačového videnia prechod od kvalitnej detekcie a rozpoznania objektov k chápaniu scény nie je možný bez detekcie a vytvárania vzťahov medzi snímanými 3D objektami, čo však nie je realizovateľné bez dostatočnej množiny

173 Por. MITCHELL, *Artificial Intelligence*, s. 106-108.

situačných videí danej scény a inovatívnych algoritmov, ktoré ich budú spracúvať, aby následne systém AI konkrétny rozpoznaný objekt (napr. mačku vo dverách, človeka na prechode) vnímal a „chápal“ v príbehu celej snímanej scény.¹⁷⁴

Efekt dlhého chvosta (long-tail effect)

Týmto termínom sa v oblasti umelej inteligencie rozumie veľký rozsah možných neočakávaných situácií, s ktorými by sa systém AI mohol stretnúť. Termín „long-tail“ pochádza zo štatistiky a vyjadruje určité rozdelenie pravdepodobnosti v tvare pretiahnutého chvosta, ktorý indikuje zoznam veľmi nepravdepodobných, ale možných situácií, ktoré vo výnimočných prípadoch môžu nastať. Tieto situácie nazývame aj hraničnými/okrajovými prípadmi. Efekt dlhého chvosta môžeme ľahko ilustrovať na pravdepodobnosti vybraných situácií, s ktorými sa autonómne vozidlá môžu v reálnej prevádzke stretnúť – viď obr. č. 5, ktorý poukazuje na riziko okrajových situácií, na ktoré autonómny systém v reálnej prevádzke vôbec nebude pripravený a v danej situácii môže vykonať nesprávne rozhodnutie.

V reálnom svete jednoducho nedokážeme všetko popísať a predložiť systémom strojového učenia na vytrénovanie.¹⁷⁵

Ak sa nad efektom dlhého chvosta zamyslíme, vidíme, že trénovaním (učením s učiteľom, supervised learning) nie sme schopní systém AI správne naučiť zvládať všetky hraničné situácie, keďže tieto nedokážeme dostatočne zahrnúť do trénovacích dát. A ak systém AI nie sme schopní na tieto hraničné situácie pripraviť, s veľkou pravdepodobnosťou pri ich výskyte budeme čeliť neočakávaným chybám a zlyhaniam.

V kontexte zavádzania systémov umelej inteligencie do reálneho nasadenia sa tento problém v rámci konzervatívneho prístupu rieši kombináciou špeciálnych množín trénovacích dát obsahujúcich rôzne hraničné situácie a osobitne naprogramovanými obmedzeniami pre hraničné stavy, čo však neprináša dostatočnú robustnosť v reálnej

174 Por. KARPATY, A. *Computer vision research feels a bit stagnating in a local minimum of 2D texture recognition on ImageNet...* [on-line]. Twitter. [cit. 10. februára 2022].

Dostupné na internete: <<https://twitter.com/karpathy/status/1491452689825165314>>

175 „We can't realistically label everything in the world and meticulously explain every last detail to the computer.“

BENGIO, Y. *Machines Dream*. In: BEYER, D. ed. *The Future of Machine Intelligence: Perspectives from Leading Practitioners*. Sebastopol, Calif.: O'Reilly Media, 14.

prevádzke. Moderný prístup zasa kombinuje algoritmy učenia s učiteľom a učenia bez učiteľa (supervised and unsupervised machine learning), t.j. natrénovanie systému na určitej množine označených dát (labeled dataset) a následné učenie sa bez učiteľa, t.j. snaha naučiť systém kategorizovať a zoskupovať vstupné dáta do celkov, pre ktoré systém AI má konkrétne riešenia a reakcie.¹⁷⁶



Obr. 6. Pravdepodobnosť výskytu niektorých situácií, s ktorými sa môže autonómne vozidlo stretnúť v prevádzke

Tento prístup koketuje s myšlienkou abstrakcie a analógie, t.j. spôsobom riešenia, v ktorom človek dokáže excelovať, ale pre strojové učenie je to stále nedosiahnuteľná meta.¹⁷⁷

¹⁷⁶ Por. MITCHELL, *Artificial Intelligence*, s. 103.

¹⁷⁷ V tomto kontexte Yann LeCun, špičkový počítačový vedec v oblasti strojového učenia a počítačového videnia uvádza: „Učenie bez učiteľa je temnou stránkou umelej inteligencie.“ (používa výraz dark matter,

Z doteraz popísaných rizikových faktorov systémov AI vyplýva, že akékoľvek priblíženie sa k všeobecnej a uvedomelej umelej inteligencii (AGI) prekračuje možnosti súčasného strojového učenia prostredníctvom kategorizovaných a označených tréningových dát. V kontexte súčasných algoritmov strojového učenia takmer všetko učenie by muselo prebiehať bez učiteľa (unsupervised), avšak stále nemáme k dispozícii také algoritmy učenia sa bez učiteľa, ktoré by boli, resp. v budúcnosti mohli byť dostatočne úspešné. Človek, napriek tomu, že neustále robí chyby, dokáže konať skutočne inteligentne – jednoducho má to, čo chýba všetkým súčasným systémom umelej inteligencie, a to zdravý rozum (predpoklad pre všeobecnú inteligenciu¹⁷⁸), ktorého súčasťou je schopnosť abstrahovať a na základe analógií a konceptov nachádzať riešenia, myslieť a riešiť na báze nesmierneho rozsahu najrozmanitejších vedomostí, poznatkov a skúseností. Ľudská bytosť využíva zdravý rozum podvedome a spontánne v akejkoľvek oblasti života. Preto **pre mnohých je dôveryhodný systém umelej inteligencie, ktorý môže úspešne fungovať v komplexnom reálnom svete, podmienený schopnosťou mať zdravý rozum tak, ako má človek.**^{179 180}

ktorý sa vo fyzike používa na vyjadrenie temnej hmoty [27%], ktorá spolu s temnou energiou [68%] tvorí 95% vesmíru a je pre nás stále tajomstvom, t.j. nedokážeme ju detekovať a popísať, pozorujeme však jej účinky:-)

LECUN, Y., MISRA, I. *Self-supervised learning: The dark matter of intelligence* [on-line]. [cit. 4. februára 2022].

Dostupné na internete: <<https://ai.facebook.com/blog/self-supervised-learning-the-dark-matter-of-intelligence>>

178 Gary Marcus, zakladateľ spoločnosti Geometric Intelligence (získanej spoločnosťou Uber) a profesor psychológie a neurónových vied na Newyorskej univerzite hovorí: „Predpokladom všeobecnej inteligencie je zdravý rozum; kým ho nedosiahneme, zostane nám úzka umelá inteligencia, ktorá je zriedkavo robustná a nikdy nie tak flexibilná ako ľudský rozum.“

CHADHA, S. *“Common Sense” is the Dark Matter of Artificial Intelligence*. [on-line]. [cit. 4. februára 2022].

Dostupné na internete: <<https://hackernoon.com/the-dark-matter-of-ai-common-sense-is-not-so-common>>

179 Por. MITCHELL, *Artificial Intelligence*, s. 104.

180 Akonáhle rozšírime svoj pohľad od učenia sa bez učiteľa k požiadavke „zdravého rozumu“, môžeme LeCunovo vyjadrenie rozšíriť: „Zdravý rozum je temnou stránkou umelej inteligencie.“

CHADHA, *“Common Sense” is the Dark Matter of Artificial Intelligence*, [on-line]. [cit. 4. februára 2022].

Dostupné na internete: <<https://hackernoon.com/the-dark-matter-of-ai-common-sense-is-not-so-common>>

Klamanie hlbokých sietí a ich zraniteľnosti (fooling deep neural networks and vulnerability to hacking)

Čo má spoločné špeciálny make up, pokreslená cesta, samolepkami oblepený stĺp, či avantgardná potlač na tričku? Dokážu dokonale zmiast' rôzne moderné systémy strojového videnia, detekcie objektov, rozpoznávania tvárí, či autonómnych systémov riadenia vozidiel. Žiaľ, osobitne v poslednej dekáde akcelerovaného vývoja systémov strojového učenia zisťujeme, že **je neuveriteľnej jednoduché priebežne a mnohými spôsobmi oklamať hlboké neurónové siete**.¹⁸¹ Navyiac, oklamanie systémov AI je možné vykonať nielen pre človeka ľahko viditeľnými a rozlíšiteľnými spôsobmi, ale častokrát i technikami, ktoré sú pre ľudskú bytosť nepostrehnuteľné. Obr. č. 7 napríklad ukazuje dvojice príkladov, na ktorých konvolučná neurónová sieť AlexNet, ktorá bola najlepšou sieťou na klasifikáciu obrázkov z databázy ImageNet v roku 2012, totálne pohorela. Nesprávne klasifikovaný obrázok z dvojice obsahuje voľným okom prakticky nepostrehnuteľné úpravy na úrovni pixelov, čo stačí na úplné pomýlenie – a ako vidíme na obrázku – aj možné cielené pomýlenie systému AI konkrétnym smerom. Pritom človek s klasifikáciou všetkých uvedených obrázkov nemá najmenší problém.



Obr. 7. Správne a nesprávne klasifikované obrázky sieťou AlexNet.¹⁸²

Neurónové siete môžu byť však klamané aj inými spôsobmi. Napríklad – ak sa budeme

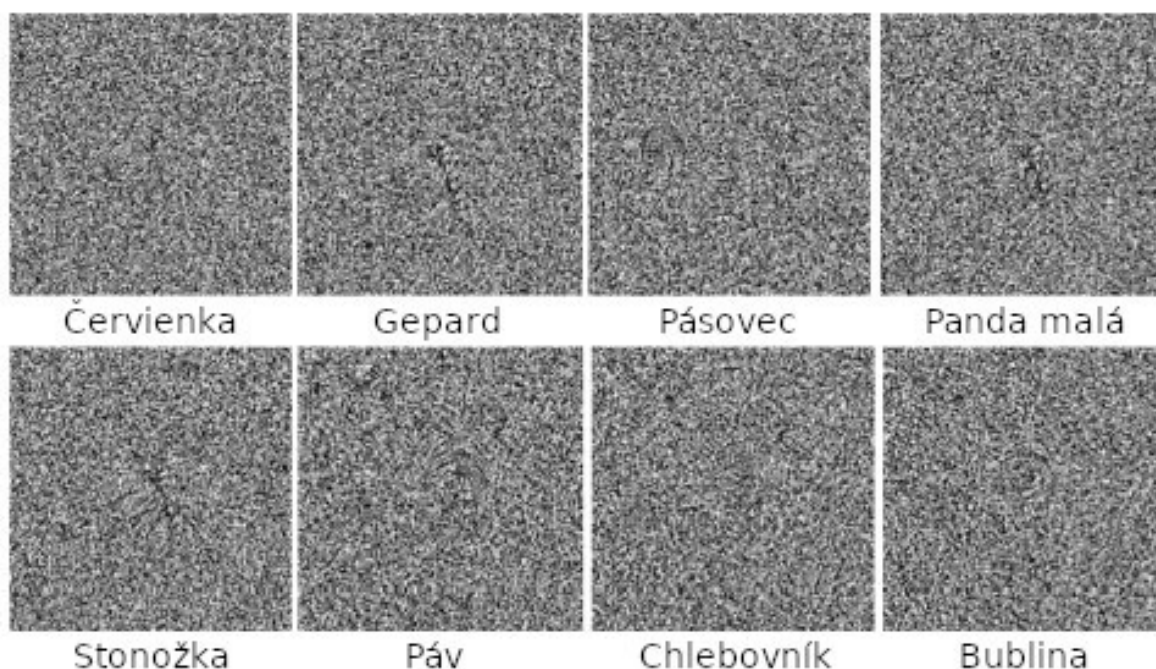
181 Por. MITCHELL, *Artificial Intelligence*, s. 110.

182 MITCHELL, *Artificial Intelligence*, s. 110, upravené autorom.

pohybovať stále v oblasti strojového videnia – pre človeka nedefinovateľné obrázky náhodného šumu môže hlboká sieť vyhodnotiť s veľkou mierou istoty ako konkrétny objekt.

Obrázok č. 8 zobrazuje príklady, ktoré vyzerajú ako náhodný šum, no pre AlexNet a iné konvolučné siete ide s viac než 99% pravdepodobnosťou o konkrétne kategórie objektov.

A aby toho nebolo málo, príklady šumu z obr. č. 8, ktoré zmiatli konvolučné siete, boli vygenerované pomocou výpočtových postupov, ktoré sa nazývajú genetické algoritmy a sú inšpirované procesmi z biologických organizmov.¹⁸³ Ide teda o obrázky, ktoré sa pomocou genetických algoritmov „vyvinuli“¹⁸⁴ do podoby nešpecifického šumu pre človeka, avšak pre systémy AI do podoby jasne identifikovaných kategórií objektov.



Obr. 8. Príklady šumu, ktoré konvolučné siete vyhodnocujú ako kategórie objektov.¹⁸⁵

Rôzne výskumy v tejto oblasti len potvrdili, že v rámci CNN je na zlyhanie náchylná nielen AlexNet, ale zraniteľné sú aj viaceré iné konvolučné siete, a to napriek tomu, že mali

183 Por. NGUYEN, A., YOSINSKI, J., CLUNE, J. *Deep Neural Networks are Easily Fooled: High Confidence Predictions for Unrecognizable Images*. [on-line]. CVPR, 2015. [cit. 12. februára 2022].

Dostupné na internete: <https://cv-foundation.org/openaccess/content_cvpr_2015/papers/Nguyen_Deep_Neural_Networks_2015_CVPR_paper.pdf>

tent_cvpr_2015/papers/Nguyen_Deep_Neural_Networks_2015_CVPR_paper.pdf>

184 MITCHELL, M. *An Introduction to Genetic Algorithms*. Cambridge, Mas.: MIT Press, 1996.

185 NGUYEN, YOSINSKI, CLUNE, *Deep Neural Networks are Easily Fooled: High Confidence Predictions for Unrecognizable Images*, upravené autorom.

rozličné architektúry, hyperparametre a množiny trénovacích dát.¹⁸⁶ I keď to Christian Szegedy a jeho kolegovia vo svojom vedeckom článku nazvali jednou zo „zaujímavých vlastností“ neurónových sietí, ide o reálny problém – jednoducho, konvolučné **neurónové siete sú náchylné na zlyhanie pri záškodníckych dátach** (adversarial examples).

Problém je však širší a netýka sa „len“ CNN, t.j. konvulčných neurónových sietí. Ako potvrdili ďalšie výskumy, pomocou (nielen) genetických algoritmov je možné pripraviť také vstupné dáta a cesty, ktoré podobným spôsobom oklamú hlboké neurónové siete vo všeobecnosti.¹⁸⁷ Tieto útoky v súčasnosti dokážu oklamať systémy identifikácie osôb, autonómnych vozidiel, spracovania lekárskeho dát, rozpoznávania reči, analýzy textu a pod.

Je vážnym zistením, že **mnohé z možných útokov sú prekvapivo robustné** – dokážu účinne oklamať rôzne a diametrálne odlišné sofistikované systémy strojového učenia.¹⁸⁸

Zistené zraniteľnosti pomocou záškodníckych dát tak prinášajú dilemu a poznatok.

Dilemu medzi evidentným úspechom systémov hlbokého učenia v rozličných zadaniach umelej inteligencie a jednoduchosťou, s akou je možné tieto systémy oklamať.

Poznatok, že ani o súčasných pokročilých systémoch strojového učenia nemôžeme povedať, že ich učiaci proces je podobný tomu ľudskému a ich schopnosti nie je možné porovnávať s ľudskými, toľž hovoriť o ich rovnocennosti alebo prekročení tých ľudských.¹⁸⁹ Znovu tak vyvstáva otázka – čo sa v skutočnosti tieto neurónové siete naučili? Ide skutočne o zručnosti korešpondujúce s konceptami, ktoré sa ich snažíme naučiť? Lebo napríklad pri vizuálnych systémoch je esenciálny rozdiel medzi rozpoznávaním objektov a pochopením zobrazovanej scény. A ďalej, **chápanie nielen v kontexte scény, ale vo všeobecnosti (ako sa daný objekt správa v reálnom svete, aké má vlastnosti a pod.) sa tak javí ako nutná podmienka robustného a spoľahlivého systému umelej**

186 SZEGEDY, CH. et al. *Intriguing Properties of Neural Networks*. Proceedings of the International Conference on Learning Representations, 2014.

187 Por. NGUYEN, YOSINSKI, CLUNE, *Deep Neural Networks are Easily Fooled: High Confidence Predictions for Unrecognizable Images*, [on-line]. [cit. 12. februára 2022].
Dostupné na internete: <https://cv-foundation.org/openaccess/content_cvpr_2015/papers/Nguyen_Deep_Neural_Networks_2015_CVPR_paper.pdf>

188 Por. MITCHELL, *Artificial Intelligence*, s. 113.

189 Súčasný najlepší systém CNN síce dosahuje úspešnosť v klasifikácii objektov vyššiu ako človek, ale ako sme uviedli, mýlia sa v situáciách, ktoré ľudská bytosť s prehľadom zvláda.

Z pohľadu zraniteľnosti systémov umelej inteligencie je problém oklamania hlbokých sietí samozrejme riešený, keďže sa vytvárajú tzv. adversarial learning algoritmy a stratégie, ktorých cieľom je ochraňovať systémy AI pred týmto typom zraniteľnosti. Je však pravdou, že i keď sa pochopeniu a obrane voči rôznym potencionálnym útokom venuje veľa z vývoja súčasných systémov AI, na rozdiel od špecifických obranných riešení neexistuje žiadna účinná všeobecná obranná stratégia. Preto v rámci súčasnej úrovne vývoja hlbokých sietí môžeme povedať, že podobne – ako v oblasti klasickej kybernetickej bezpečnosti – ide o nekončiaci zápas medzi hľadaním nových útočných postupov a tvorbou obranných stratégií.¹⁹²

Z pohľadu ďalšieho rozvoja systémov umelej inteligencie ide o jasné indikovanie rozdielov medzi schopnosťami súčasných systémov ANI a riešeniami, ktoré môžu viesť k AGI, t.j. k uvedomelej všeobecnej inteligencii ako takej.

Povery (superstition) – prirýchly Q-learning

Poverou zvykneme nazývať mylnú vieru, že určitá akcia, či úkon môžu pomôcť zapríčiniť dobrý alebo zlý výsledok. V oblasti umelej inteligencie ide o problém prevažne v rámci algoritmov učenia formou odmeňovania (reinforcement learning), pri ktorých sa tréning modelu (agenta) realizuje prostredníctvom interakcie s prostredím metódou pokus-omyl.¹⁹³

V rámci tréningu systému AI (napr. robotického systému) vzniká povera vtedy, ak sa daný systém chybné naučí vykonávať nejaký nepotrebný, ba až možno nebezpečný úkon pre dosiahnutie požadovaného cieľa. Môže tak ísť napríklad o prekvapivý a nebezpečný pohyb robotického ramena alebo nečakaný manéver autonómneho vozidla, ktorý ohrozí účastníkov premávky.

Nielen vyvarovanie sa poverám, ale aj celkový návrh úspešného systému učenia formou odmeňovania je stále ešte určitou alchýmiou či umením, ktoré zvláda relatívne malá skupina expertov s veľkým citom a praxou v ladení hyperparametrov.^{194 195}

190 Por. MITCHELL, *Artificial Intelligence*, s. 112-114.

191 Znovu sa tak dostávame k problematike kognitívneho vnímania a „zdravého rozumu“.

192 Por. MITCHELL, *Artificial Intelligence*, s. 113-114.

193 O učení formou odmeňovania pojednáva kapitola 1.6.3.

194 O ladení hyperparametrov, obraznej alchýmii a umení bola rozprava v kapitole 2.1.1.

195 Por. MITCHELL, *Artificial Intelligence*, s. 138-143.

I napriek neodškriepiteľnému úspechu vývoja a nasadenia súčasných systémov umelej inteligencie v širokom spektre akademického i reálneho prostredia musíme mať neustále na pamäti, že tieto systémy môžu zlyhávať najrozličnejšími a často neočakávanými spôsobmi v dôsledku nemožnosti pripraviť dostatočne veľkú množinu tréovacích dát, prípadne ich nesprávnej voľby bez dostatočnej kvality alebo s predsudkami, nadmernému prispôbovaniu sa tréningovým údajom, efektu dlhého chvosta, rizikám plynúcim z oklamania hlbokých sietí, ich zraniteľností a povier, pod čo sa podpisuje i nedostatok odbornej erudovanosti potrebnej pre dizajn a ladenie hyperparametrov pri príprave funkčného a úspešného riešenia.

Pri hlbšom pohľade na uvedené problémy nás môžu dobiehať aj ich ďalšie dôsledky – nielen riziká priameho zlyhania, ale aj **realita výsledkov, ktoré môže byť ťažké správne interpretovať** (čo sa vlastne sieť naučila, čo výstup z daných dát na vstupe vlastne znamená) a **neschopnosť predvídať, kedy sa jednotlivé zlyhania prejavia** (za akých podmienok, pri akej súhre okolností, s dôsledku akej dynamiky vnútorného vývoja, resp. činnosti systému AI).

V zásade možno povedať, že súčasné systémy strojového učenia budú pri nasadení v reálnom svete stále trpieť neduhmi, ktoré bez dosiahnutia silnej AI nebude možné odstrániť. Totižto **úzke zameranie súčasných systémov strojového učenia je v kontraste s procesom učenia sa človeka, ktoré je vecou všetkých rozmerov ľudskej inteligencie a aktívneho zaradenia sa do sveta.**¹⁹⁶ Pri zavádzaní súčasných systémov AI do nasadenia v reálnom svete si treba uvedomiť, že ich spoľahlivosť je limitovaná a schopnosti sú obmedzené. Analogicky treba s týmto rizikom narábať aj v prípade etických noriem a mantinelov, ktoré dokážeme v rámci systémov umelej inteligencie implementovať.

2.2. Bezpečnosť procesov umelej inteligencie

Ako sme už viackrát uviedli, systémy umelej inteligencie sa stávajú neoddeliteľnou súčasťou fungovania rodiacej sa informačnej spoločnosti a v mnohých oblastiach života sú nasadené a pomerne úspešne využívané. Druhý pohľad na limity a riziká súčasných systémov AI preto zameriame na bezpečnostné aspekty ich nasadenia a útoky, ktorým čelia.

196 Por. MITCHELL, *Artificial Intelligence*, s. 97.

V zásade rozlišujeme tri druhy útokov na systémy AI používané v reálnej prevádzke:¹⁹⁷

- útoky na dôvernosť (confidentiality attacks)
- útoky na zraniteľnosti (evasion attacks)
- útoky s cieľom ovplyvniť model (poisoning attacks)

Útoky na dôvernosť (confidentiality attacks)

Ide o útoky zamerané na dáta uložené v rámci modelov systémov AI, ktorých zámerom je snaha kopírovať model za účelom extrahovania tréningových dát a parametrov z týchto modelov.¹⁹⁸ Ako sme uviedli v kapitole 2.1.2., kvalitné a rozsiahle tréningové dáta sú jedným z podstatných faktorov správne fungujúceho systému AI. Pre reálne nasadenie v konkrétnom sektore, resp. organizácii sa tak súčasťou tréningových dát môžu stať dôverné, osobné, prípadne strategické informácie.

Naviac v kontexte vysokej sofistikovanosti kvalitných systémov strojového učenia je nezanedbateľným dôvodom aj snaha získať informácie o dizajne a hyper parametroch konkrétného systému AI.

Útoky na zraniteľnosti (evasion attacks)

V tejto oblasti sa útočníci zameriavajú na odhaľovanie a zneužitie existujúcich zraniteľností v modeloch za účelom zmanipulovania výsledkov systémov AI. Ide tak – na základe existujúcich zraniteľností a limitov – o ovplyvňovanie činnosti systémov AI priamo počas ich prevádzky.¹⁹⁹

Útoky s cieľom ovplyvniť model (poisoning attacks)

Ide o celú škálu útokov, ktorých cieľom je ovplyvňovanie modelu, tréningového procesu a tým aj výslednej činnosti systému AI. Ide o zámerné ovplyvňovanie tréningového procesu (učenia) modelu za účelom manipulovania jeho následných rozhodnutí v prevádzke.

197 REHÁK, M. *Útoky na systémy umělé inteligence a jejich obrana*. In: *Umělá inteligence 2021*. Praha: TUESDAY Business Network, 2021.

198 Jedným z jednoduchých spôsobov útoku je napr. cielené zadávanie veľkého kvanta vstupných dát do systému konkurencie a sledovanie/analýza výsledkov.

REHÁK, *Útoky na systémy umělé inteligence a jejich obrana*.

199 Oklamanie systému na rozpoznávanie tvárí, podvodné získanie financií, resp. úveru z bankového domu, vyradenie z činnosti systému pre automatické riadenie streľby alebo skupinového riadenia dronov, atď.

V tejto súvislosti je zaujímavým etickým problémom, ktorý sme doteraz nespomínali, aj iný rozmer zámerného ovplyvňovania systému AI – zámerné nastavenie parametrov a ovplyvňovanie tréningového procesu samotnými tvorcami daného systému, či už zo svojej vôle alebo na základe zadania objednávateľa systému. Osobitne v prípade, keď ide o vládny subjekt, nadnárodnú spoločnosť, celosvetovú sociálnu sieť, informačné platformy a pod., sa jedná o vážny a delikátny problém. Avšak problém, ktorý je – žiaľ – reálny.²⁰⁰

Ak sa pozrieme na pomerné zastúpenie výskytu jednotlivých druhov útokov, prax poukazuje na veľkú rôznorodosť výskytu a použitia.

Útokov na dôvernosť (confidentiality attacks) je v súčasnosti minimum, pretože ich cieľ sa dá dosiahnuť inými spôsobmi, ktoré sú mnohokrát jednoduchšie, časovo menej náročné a lacnejšie. Keďže väčšina organizácií útoky tohto druhu neočakáva, ich nasadenie môže byť oveľa väčším rizikom s nedobрым koncom. Preto treba byť opatrný.

Úspešnosť útokov s cieľom ovplyvniť model (poisoning attacks) sa v súčasnosti viaže na pokročilé znalosti z oblasti umelej inteligencie, preto sa tento druh skôr nahrádza jednoduchšími útokmi zacielenými na zraniteľnosti systémov AI.

Potenciál útokov s cieľom ovplyvniť model je však veľký a osobitne s možnosťou nástupu všeobecnej a uvedomelej umelej inteligencie ich atraktivnosť a nutkanie ich zneužiť môže veľmi rásť.

Útoky na zraniteľnosti (evasion attacks) tvoria v súčasnosti takmer 100% útokov. Vzhľadom na neustály vývoj a otvorenú komunikáciu ohľadom zraniteľností systémov AI býva mnohokrát dostatok informácií pre zneužitie konkrétnych zraniteľností (vytvorenie exploitu) a realizácia útokov je pomerne jednoduchá, takže sa dajú ľahko vykonať.²⁰¹

200 V súčasnosti ide prevažne o problém, s ktorým sa potýkajú systémy umelej inteligencie nasadené v oblasti spracovania informácií a informačných tokov.

Viac o možnostiach zneužitia informácií napr. v kapitole *Informácia ako nástroj moci práce*:

ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi*, [online], s. 37-39. [cit. 14. februára 2022].

Dostupné na internete:

<https://peter.santavy.cloud/docs/Analiza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

201 To neznamená, že by bolo vhodnejšie hľadanie, analýzu a riešenie zraniteľností skrývať. Už v rámci klasickej kybernetickej bezpečnosti je overeným faktom, že bezpečnosť prostredníctvom utajenia

V kontexte bezpečnosti procesov umelej inteligencie sa etické výzvy, doteraz rozoberané ako súčasť návrhu a realizácie systémov AI, rozširujú aj o oblasť etiky použitia, resp. riziká zneužitia. V oblasti klasickej informačnej bezpečnosti sú na jednej strane barikády tí, ktorí bezpečnosť systémov strážia (profesionáli v oblasti kybernetickej bezpečnosti), odhaľujú zraniteľnosti s cieľom nápravy a ochrany systémov (napr. etickí hackeri), riešia škody a ochraňujú obeť, na strane druhej je skrytý zástup od jednotlivcov až po organizované skupiny, ktorí sa snažia systémy zneužiť, znefunkčniť, či vykradnúť dôležité dáta z veľmi rôznorodých dôvodov (od ideologických, psychologických, sociologických až po mocenské dôvody a kybernetickú kriminalitu²⁰²). Nie inak je tomu tak aj v oblasti umelej inteligencie, keď prekročenie etického rámca je pre mnohých príliš lákavé pokušenie, ktorému nebudú vedieť a – ak sa etika návrhu a využívania systémov AI nebude zodpovedne riešiť – ani nebudú chcieť a mať prečo odolať.

A keď sme už pri etike, na verejnosť neustále presakujúce informácie o únikoch z vládnych zdrojov a bezpečnostných agentúr dávajú tušiť, že pokušenie je priveľké aj pre tak veľkých „hráčov“ v oblasti nasadzovania systémov AI, ako sú vlády, nadnárodné spoločnosti a mocenské zoskupenia, keď zápasia (*sic!*) s túžbou využitia (zneužitia) umelej inteligencie ako nástroja na dosiahnutie neprimeraného zisku, ideologických a mocenských cieľov.²⁰³

(security by obscurity) je skutočne zlý nápad. Pri otvorenom prístupe k problému zraniteľností sa tieto môžu rýchlejšie identifikovať, komplexnejšie analyzovať a plošne odstraňovať.

202 Informačná kriminalita v súčasnosti patrí spolu s drogami a obchodovaním s ľuďmi medzi tri najvýnosnejšie oblasti kriminality.

ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi*, [on-line], s. 37. [cit. 14. februára 2022].

Dostupné na internete:

<https://peter.santavy.cloud/docs/Analiza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

Ba čo viac – už v rokoch 2015/16 sa kybernetická kriminalita stala oblasťou najvýnosnejšou.

KHIMJI, I. *Cybercrime Is Now More Profitable Than The Drug Trade* [on-line]. [cit. 14. februára 2022].

Dostupné na internete: <<https://www.tripwire.com/state-of-security/regulatory-compliance/pci/cybercrime-is-now-more-profitable-than-the-drug-trade/>>

203 Viac o týchto problémoch pojednávame v kapitole 2.5.

2.3. Kybernetická bezpečnosť

Súčasný systémy umelej inteligencie sú postavené na moderných informačných a komunikačných technológiách. Preto – ako elektronické systémy – zdieľajú aj všetky riziká a nehuhy moderných kybernetických systémov.²⁰⁴ Tretí pohľad na limity a riziká súčasných systémov AI sa preto zameriava na kybernetickú bezpečnosť riešení postavených na umelej inteligencii.

Naviac, **kybernetická bezpečnosť je oblasťou, ktorá určitým spôsobom doteraz uvedené problémy prepája, keďže v reálnom nasadení vo väčšine prípadov nie je možné riziká striktne rozdeľovať podľa vyššie uvedených kapitol. Mnohé riziká napríklad atakujú bezpečnosť procesov umelej inteligencie, zneužívajú nedostatky dizajnu alebo amplifikujú dôsledky rizikových faktorov, pričom však ako vektor útoku používajú zraniteľnosti v oblasti kybernetickej bezpečnosti.**

Napríklad jedna z najpoužívanějších platforiem pre vývoj a prevádzku systémov strojového učenia, TensorFlow od Googlu mala len za rok 2021 oficiálne evidovaných 201 zraniteľností (pridelené CVE²⁰⁵). Je zaujímavé, že z pohľadu kybernetickej bezpečnosti má väčšina týchto zraniteľností nízke rizikové skóre (CVSS Score) – dôsledky ich zneužitia tak nemusia byť veľkým rizikom a viesť ku klasickým kybernetickým kompromitáciám (hacknutiam). Vo viacerých prípadoch môže však byť ovocím ich zneužitia prerušenie prebiehajúcej úlohy alebo jej nesprávne vykonanie. Takže nízke CVSS skóre v oblasti kybernetickej bezpečnosti nemusí znamenať malé riziko pre prevádzku systému AI.²⁰⁶

204 Uvádzame konkrétnejší pojem kybernetický systém, v ktorom sa snúbi skĺbenie výpočtovej techniky a kybernetického informačného priestoru (cyberspace). Pojem informačný systém je teda jeho nadmnožinou, zahŕňajúcou aj oblasti mimo počítačového sveta.

205 Systém CVE (Common Vulnerabilities and Exposures) poskytuje referenčnú metódu pre verejne známe zraniteľnosti a ohrozenia informačnej bezpečnosti. Tento systém spravuje Národné centrum kybernetickej bezpečnosti Spojených štátov amerických (NCF), ktoré prevádzkuje spoločnosť The Mitre Corporation, a financuje ho Národná divízia kybernetickej bezpečnosti amerického ministerstva pre vnútornú bezpečnosť. V oblasti kybernetickej bezpečnosti ide o relevantnú, bežne používanú a celosvetovo akceptovanú metódu klasifikácie zraniteľností.

Common Vulnerabilities and Exposures. [on-line]. [cit. 15. februára 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Common_Vulnerabilities_and_Exposures>

206 Google » *Tensorflow: Security Vulnerabilities* [on-line]. [cit. 11. januára 2022].

Dostupné na internete: <https://www.cvedetails.com/vulnerability-list/vendor_id-1224/product_id-53738/Google-Tensorflow.html>

S postupným rozširovaním využívania tejto platformy a celkovo systémov AI badať aj rapídny medziročný nárast evidovaných zraniteľností,²⁰⁷ čo na jednej strane znamená viac potencionálnych spôsobov kompromitácie, na strane druhej je však indikátorom aj zvýšeného záujmu o bezpečnosť softvérového vybavenia pre masívne zavádzanie AI.

Indikátorom skutočného rizika nie je tak ani počet objavených a zdokumentovaných zraniteľností, ale skôr priemerný čas potrebný na ich odstránenie.

Dôležitým faktorom je tiež spôsob nasadenia technologického vybavenia systémov AI, keďže v mnohých prípadoch ide o systémy, ktoré po nasadení do prevádzky nie sú predmetom údržby a prípadných aktualizácií (opráv) bezpečnostných chýb, takže sa stávajú bezpečnostným rizikom.²⁰⁸ Neaktualizovaný jednoduchý inteligentný asistent (vznešene sa nazývajúci umelou inteligenciou vzhľadom na použité jednoduché prvky z AI) vo fotoaparáte nie je problémom, no systém, ktorý napr. riadi dopravu a nie je ťažké ho zmiast'ť, resp. navodiť dopravný kolaps a nehody – ak sa objavené chyby v jeho dizajne, či implementácii neošetria – môže byť veľmi nebezpečný a možné obete na životoch by boli len otázkou času, kým dané chyby niekto nezneužije, resp. sa nevyskytne súbeh podmienok (race condition), ktoré tento systém AI pomýlia.

Ako každá iná oblasť kybernetickej, resp. vo všeobecnosti informačnej bezpečnosti i oblasť bezpečnosti systémov AI je proces. A ako v každej inej oblasti informačných technológií, ani v oblasti systémov umelej inteligencie neexistuje dokonale bezpečný a spoľahlivý systém.

Okrem snáh o riešenie bezpečnosti systémov AI je jednou z veľkých úloh aj hľadanie spôsobov, ako minimalizovať dôsledky týchto zlyhaní či už pevnými obmedzeniami, dohľadovými riešeniami alebo inými prostriedkami. Zaujímavosťou je, že v súčasnosti takmer všetky moderné nástroje používané v rámci kybernetickej bezpečnosti na detekciu, dohľad a ochranu pred útokmi taktiež obsahujú prvky umelej inteligencie:-)

207 Google » *Tensorflow: Vulnerability Statistics* [on-line]. [cit. 11. januára 2022].

Dostupné na internete: <https://www.cvedetails.com/product/53738/Google-Tensorflow.html?vendor_id=1224>

208 Toto je jedno z vážnych rizík nastupujúceho masívneho využívania internetu vecí (IoT), ktoré s príchodom 5G sietí smerujú k intenzívnej interkonektivite a tým aj k on-line zraniteľnosti. Zatiaľ ani neexistuje pevný rámec, ktorý by tento problém riešil, aj keď viacero iniciatív v tejto oblasti existuje. Doterajšie problémy v oblasti kybernetickej bezpečnosti s IoT poukazujú na veľké výzvy v blízkej budúcnosti.

Tiež si treba uvedomiť, že počet zraniteľností je priamo úmerný komplexnosti systémov – ktorá je v oblasti umelej inteligencie jedným z normatívnych faktorov sofistikovaných a úspešných systémov. Podobne, počet kompromitácií, resp. zneužití týchto zraniteľností priamo rastie s mierou nasadenia a využívania v reálnom svete.

2.4. Technologická komplexnosť a potrebná infraštruktúra

Ako sme uviedli v predchádzajúcej kapitole, súčasné systémy umelej inteligencie sú postavené na moderných informačných a komunikačných technológiách, ktoré sú tvorené nielen softvérovým vybavením, ale i hardvérom, t.j. elektronickým a vo všeobecnosti technologickým zázemím, vďaka ktorému softvér systémov umelej inteligencie môže pracovať. Štvrtý pohľad na limity a riziká súčasných systémov AI sa preto zameriava na hardvérové vybavenie a technologickú infraštruktúru potrebnú pre zabezpečenie spoľahlivej funkčnosti a robustnosti systémov umelej inteligencie.

Výpočtové systémy prevádzkujúce algoritmy AI – vzhľadom na postupy strojového učenia – musia disponovať veľkým výpočtovým výkonom²⁰⁹ a vzhľadom na enormné množstvo spracúvaných dát aj schopnosťou zhromažďovať, ukladať a manipulovať s veľkými štruktúrovanými i neštruktúrovanými dátami (viď príloha č. 2 – štruktúrované a neštruktúrované dáta).

Operácie vykonávané neurónovými sieťami tvoria pomerne špecifickú podmnožinu funkcionality, ktorú vedia moderné počítačové systémy ponúknuť, pričom ich výkon nie je špeciálne optimalizovaný primárne pre vykonávanie operácií napríklad strojového učenia. Kvôli optimalizácii výkonu sa preto v posledných rokoch vyvíja a vyrába špecializovaný hardvér²¹⁰, ktorý ponúka oveľa vyšší výkon v operáciách algoritmov umelej inteligencie než bežné – i keď výkonné – počítačové vybavenie.²¹¹

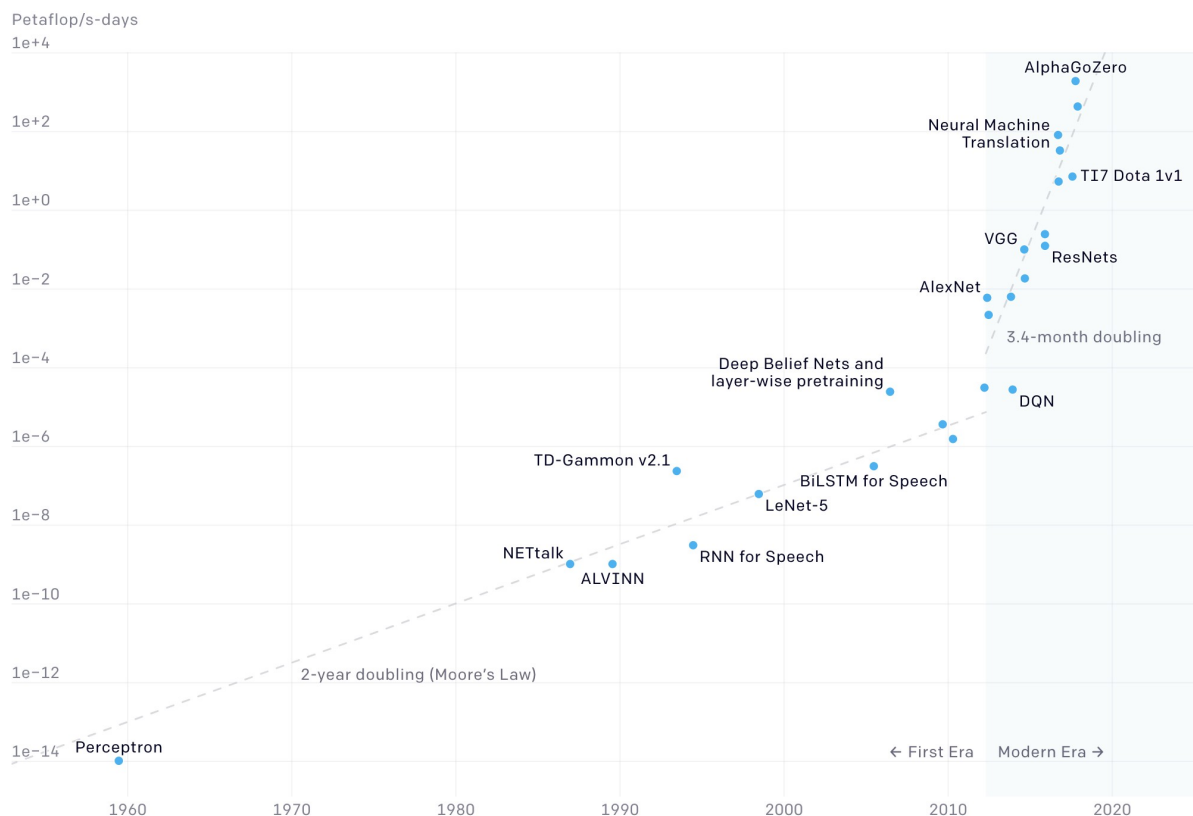
209 **S rozvojom algoritmov strojového učenia sa požadovaný výpočtový výkon v poslednej dekáde zvyšuje exponenciálne!**

AMODEI, D., HERNANDEZ, D. *AI and Compute*. [on-line]. [cit. 19. januára 2022].

Dostupné na internete: <<https://openai.com/blog/ai-and-compute/>>

210 Procesory optimalizované pre AI, hardvérové akcelerátory, atď.: ASIC - Application Specific Integrated Circuit, špecializované integrované obvody navrhnuté na vykonávanie úloh spojených s AI, TPU - Tensor Processing Unit, špecializovaný akcelerátor od Google, využívanie výpočtového výkonu cloudov,...

211 Trh s týmto hardvérom je pomerne rozsiahly – od superpočítačov, cez výkonné systémy autonómnych



Obr. č. 9. Výpočtový výkon systémov AI sa v poslednej dekáde zvyšuje exponenciálne.²¹²

Je to nutné, ak uvažíme, že bez potrebného výkonu nie je možné zabezpečiť nasadenie systémov AI tak, aby pracovali v reálnom čase.

Z apetítu súčasných moderných systémov AI sa javí, že ich ďalší rozvoj – osobitne v túžbe po dosiahnutí uvedomelej AGI – bude pravdepodobne vyžadovať nielen kvantitatívnu, ale predovšetkým kvalitatívnu zmenu v schopnostiach hardvéru výpočtových systémov.^{213 214}

vozidiel, vojenskú techniku a rôznorodé technické vybavenie, špecializované integrované obvody (čipy) pre osobné počítače, tablety a mobily, až po napríklad špičkové a výkonné nástroje pre softvérových vývojárov a vedcov v oblasti AI, ktoré je možné kúpiť za lacný peniac (desiatky Eur) a ktoré sa dajú k počítaču jednoducho pripojiť cez USB, PCIe alebo M2 rozhranie.

²¹² AMODEI, HERNANDEZ, *AI and Compute*. [on-line]. [cit. 19. januára 2022].

Dostupné na internete: <<https://openai.com/blog/ai-and-compute/>>

²¹³ Vývoj sa uberá rôznorodým smerom: či už ide o hľadanie nových polovodičových architektúr, neuromorfné čipy, optické systémy, biotronicke systémy, kvantové počítače,...

²¹⁴ „Kľúčovým prvkom pokroku v oblasti umelej inteligencie sú zlepšenia v oblasti výpočtovej techniky, takže pokiaľ bude tento trend pokračovať, oplatí sa pripraviť na dôsledky systémov, ktoré sú ďaleko

Technológie, potrebné pre úspešné nasadenie systémov AI, netvorí len výkonný hardvér schopný vykonávať extrémne množstvo špecializovaných operácií za sekundu. Viaceré oblasti nasadenia v reálnom svete poukazujú i na ďalšie komponenty, na ktoré nesmieme zabúdať – **špecifické periférne zariadenia a technológie vysoko rýchlostného prepojenia všetkých častí funkčného celku AI v danej oblasti nasadenia.**

Využívanie systémov umelej inteligencie v reálnom svete je pre mnohé scenáre použitia spojené s kvalitnými perifériami, t.j. vstupnými a výstupnými systémami.²¹⁵ Požadovaná kvalita a výkon periférií sa dosahuje ich neustálym zlepšovaním, implementovaním špecializovaných systémov AI (napr. na vylepšenie počítačového videnia z kamery) a orchestráciou rôznych typov periférií do spoločného celku.

Veľký objem zo senzorov generovaných a spracúvaných dát v reálnej prevádzke systémov AI a vysoké dátové toky medzi jednotlivými časťami týchto systémov, ktoré musia byť realizované v zlomkoch sekúnd, vytvárajú veľké nároky na spôsoby vysoko rýchlostného prepojenia (napríklad tzv. zbernice) v rámci jednotlivých systémov AI.²¹⁶

Vo viacerých prípadoch nasadenia však musia systémy AI komunikovať aj s inými systémami, napr. autonómne vozidlá so satelitnými navigačnými systémami, s inteligentnou infraštruktúrou vozoviek, s ostatnými vozidlami, internetom, atď.²¹⁷ A to všetko v reálnom čase.²¹⁸

za dnešnými možnosťami.“

AMODEI, HERNANDEZ, *AI and Compute*. [on-line]. [cit. 19. februára 2022].

Dostupné na internete: <<https://openai.com/blog/ai-and-compute/>>

215 Napr. autonómne vozidlá pre svoju prevádzku vyžadujú detekciu okolia a monitorovanie trasy.

Na detekciu sa v súčasnosti využíva kombinácia radarov, lidarov, satelitnej navigácie, počítačového videnia a zvukový vstupov.

216 Napr. v automobilovom priemysle sa používajú uzavreté vnútrovozidlové zbernice CAN (controller area network), prípadne modernejšie FlexRay, LIN, MOST, ktoré sú však pre veľké objemy dát systémov AI v autonómnych vozidlách nedostatočné. V nich sa zavádzajú automotive ethernet, t.j. nové vysoko rýchlostné zbernice typu ethernet (dobré známe a široko využívané v počítačových sieťach).

217 Ide o tzv. CCAM - cooperative connected automated mobility, t.j. vzájomné prepojenie systémov účastníacich sa dopravnej prevádzky a proaktívna výmena informácií medzi nimi, čo je kľúčové pre nasadenie stupňov 4 a 5 autonómie vozidiel (viď. obr. č. 11 v nasledovnej kapitole).

Ide o tzv. V2X komunikáciu, resp. prepojenie, pričom V2X = V2N + V2V + V2I + V2P (komunikácia vehicle to vehicle, network, infrastructure, pedestrian).

218 Bezdrôtová časť komunikácie vyžaduje vysoké rýchlosti a nízke latencie, čo spĺňajú technológie VANET

Krátky náčrt problematiky technologického vybavenia a nutnej infraštruktúry pre nasadenie systémov AI vo viacerých oblastiach reálneho sveta poukazuje na veľkú komplexnosť týchto systémov. A komplexnosť je problém – je rizikovým faktorom bezpečnosti a stability fungovania systémov.

Nadviažuc na tému predchádzajúcej kapitoly o kybernetickej bezpečnosti systémov AI, v kontexte komplexnosti týchto systémov si musíme uvedomiť aj rapídny nárast tzv. útočnej plochy (attack surface)²¹⁹ na tieto systémy. Ide o rozšírenie možností kybernetických útokov o ďalšie typy vzhľadom na komplexnosť a množstvo použitých technológií i potrebu sofistikovanej infraštruktúry.²²⁰

K faktorom, ktoré asi nikdy nedokážeme naplno ošetriť, patrí aj zlyhanie jednotlivých komponentov a elektronických súčiastok. Ide o malý bonus k doteraz uvedeným problémom, ktorý sa rieši rôznymi metódami redundancie, štatistickými metódami a prvkami kontroly, pričom miera zložitosti a finančnej náročnosti pri systémoch zabezpečených voči zlyhaniam elektroniky mnohokrát neúmerne rastie. **Ekonomický rozmer tak môže veľmi ovplyvňovať aj bezpečnosť systémov AI.**

Keďže komplexnosti sa nedá vyhnúť (v zásade ide o analógiu s biologickými systémami), riešením je vývoj techník zabezpečujúcich stabilitu a bezpečnosť (napr. redundancia, mnoho úrovňové bezpečnostné mechanizmy a pod.).

Čo však v rámci živých organizmov ako mechanizmy ochrany do DNA vložila evolúcia, vo svete umelej inteligencie musia zabezpečiť ľudia – tvorcovia a výrobcovia systémov. A spoľahlivú funkčnosť, ktorú u biologických systémov preverila príroda, musí v oblasti AI zabezpečiť spoločnosť svojimi pravidlami a reguláciami.

V pohľade na technologickú komplexnosť a nutnosť fungujúcej infraštruktúry (kapitola 2.4.), vzhľadom na rôznorodosť kybernetických útokov a reálnu nemožnosť vytvoriť

(Vehicular Ad Hoc Networks – Wifi medzi vozidlami na 5.9 GHz) a 5G siete. Plnú podporu by mali zabezpečiť budúce 6G siete nielen s vyššími rýchlosťami a nižšou latenciou než 5G, no predovšetkým s implementovanou sadou funkcionality využiteľných aj pre autonómne vozidlá (tzv. AI enabled networks).

219 *Attack surface*. [on-line]. [cit. 19. februára 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Attack_surface>

220 Napríklad rádiové rušenie alebo zahltenie (DoS/DDoS útoky), výpadok satelitnej navigácie a internetového pripojenia, špecifické útoky na internet vecí (IoT) a pod.

absolútne bezpečný systém (2.3.), zároveň vo vedomí procesných útokov (2.2.) a rizikových faktorov systémov umelej inteligencie (2.1.), by mala byť **nutnou podmienkou ich prevádzky schopnosť a možnosť človeka prebrať kedykoľvek kontrolu nad týmito systémami, resp. právo a možnosť verifikovať a prehodnotiť výsledky ich činnosti.**²²¹

2.5. Dôsledky pre spoločnosť

Našu rozpravu o limitoch a rizikách súčasných systémov umelej inteligencie môžeme zavrieť pohľadom na oblasť, ktorá sa prakticky dotýka každého jedného človeka modernej spoločnosti – ide o dôsledky a riziká spojené s využívaním prostriedkov AI v našom každodennom živote a interakciu, ktorú vedome či nevedome s nimi neustále podstupujeme. Ide o metamorfózy, v ktorých sa trhové mechanizmy, zábava, vzdelávanie, zdravotníctvo, doprava, vedecko-technický rozvoj, služby a spoločenské procesy, dopované umelou inteligenciou, dotýkajú nášho človečenstva.²²²

V modernej informačnej spoločnosti sa základnou komoditou stávajú informácie.²²³ Systémy umelej inteligencie sú priamo závislé na kvalitných informáciách, pričom o väčšine aktuálne najúspešnejších sofistikovaných systémoch AI sa dá povedať, že sú závislé na extrémnom množstve relevantných dát.

Aby nasadenie systémov umelej inteligencie bolo v spoločnosti úspešné, potrebné dáta musia byť neustále zhromažďované z reálneho sveta a priamo z ľudského prostredia.

Nutnou súčasťou extrémne rýchleho súčasného vývoja v oblasti umelej inteligencie je fungujúca možnosť monetizácie nasadenia systémov AI, čo sa najmarkantnejšie prejavuje

221 V mnohých aplikáciách systémov AI však reálne uskutočniteľná možnosť verifikovať výsledky ich činnosti je **takmer neriešiteľný problém**.

222 V tom lepšom prípade len dotýkajú, v horšom ho možno i menia a pretvárajú.

223 Informácia ako taká sa v informačnej spoločnosti stáva *základnou komoditou jej fungovania*. Je nielen predmetom činnosti znalostnej ekonomiky, ale stáva sa i nástrojom moci, či drogou informačnej závislosti.

ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi*, [online], s. 30. [cit. 7. decembra 2021].

Dostupné na internete:

<https://peter.santavy.cloud/docs/Analýza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

v on-line systémoch a produktoch technologických gigantov (sociálne siete, vyhľadávače,...). U týchto systémov by sa pri povrchnom pohľade mohlo javiť, že základnou komoditou sú informácie, ktoré uvedené spoločnosti bezprecedentným spôsobom zhromažďujú, analyzujú a spracúvajú. Avšak systémy umelej inteligencie so svojimi schopnosťami posúvajú túto časť paradigmatickej zmeny informačnej spoločnosti ešte ďalej: **z informácií a z produktov ich spracovania sa komoditou stávajú priamo ľudia – ba čo viac, ich psychologické profily a vzťahy, konanie a jeho ovplyvňovanie a v konečnom dôsledku i fungovanie celej spoločnosti.**

Ďokonale (povedzme, že len tak kvalitne, ako to umožňujú viaceré súčasné platformy sociálnych sietí) natrénovaná AI sa dokáže zamerať na človeka, konkrétne sociálne skupiny, resp. vo všeobecnosti nás ľudí v spoločnosti a – analogicky využitiu kvalitných psychologických metód – priviesť nás k pozornosti, akceptácii ňou predkladaných informácií (osobitne cez pokročilé metódy, napr. augment reality, podprahové metódy,...) a k **ovplyvňovaniu nášho vnímania, našich rozhodnutí a konania...**²²⁴ A môže to ísť až tak ďaleko, že prichádza k zmene nášho zmýšľania a vnímania, kým vlastne sme. Druhou stranou mince je **schopnosť systémov AI vytvárať modely, ktoré sú schopné predpovedať naše konanie.**²²⁵

Ak dosiahnutý efekt a vplyv môže byť tak závažný a prakticky takmer istý, existuje riziko i pokušenie využiť to, manipulovať a zneužiť (či už vo svete biznisu, ale aj v oblastiach oveľa kritickejších).

Ak efekt má byť takmer istý, treba stavať na extrémne presných predikciách. A extrémne presné predikcie nie je možné vykonávať bez extrémne veľkého množstva dát²²⁶ a techník, ktoré ich dokážu v reálnom čase spracovať.

224 Por. rozhovor s Jaronom Lanierom, počítačovým vedcom a spolu zakladateľom oblasti virtuálnej reality. ORLOWSKI, J. *The Social Dilemma*. [filmový dokument]. Netflix, 2020, 14:25. [cit. 7. decembra 2021]. Dostupné na internete: <<https://www.netflix.com/sk/title/81254224>>

225 Por. rozhovor s Azom Raskinom, pôvodným povolaním matematikom a fyzikom, spoluzakladateľom Centra pre humánne technológie. ORLOWSKI, *The Social Dilemma*, 17:50. [cit. 7. decembra 2021]. Dostupné na internete: <<https://www.netflix.com/sk/title/81254224>>

226 Por. rozhovor s profesorkou Shoshanou Zuboff, autorkou knihy *The Age of Surveillance Capitalism*. ORLOWSKI, *The Social Dilemma*, 15:26. [cit. 7. decembra 2021]. Dostupné na internete: <<https://www.netflix.com/sk/title/81254224>>

S tým súvisí ďalšie riziko - tzv. kapitalizmus dohľadu (surveillance capitalism), ktorý ťaží z nekonečného sledovania - **ľudia i celá spoločnosť sú vystavení neustálemu dohľadu a sledovaniu, ktoré je však len veľmi málo pod kontrolou, ak vôbec**. Osobitne tak môže byť v kontexte systémov umelej inteligencie, u ktorých zhromažďované dáta nie sú takmer vôbec pod ľudským dohľadom (ak taký dohľad je pri veľkých systémoch zhromažďujúcich denno denne extrémne množstvo dát²²⁷ vôbec možné zaručiť).

Ide o začarovaný kruh, čím viac dát systémy AI dokážu získať (a na nich sa učiť), tým presnejšie modely nášho správania ponúkajú a tým lepšie dokážu ovplyvňovať naše vnímanie, postoje a rozhodnutia – rozhodnutia i činnosť, čo následne generuje ďalšie dáta, zberané a využívané systémami AI... Na základe dát a výsledkov systémov AI sa upresňujú algoritmy implementované v týchto systémoch, aby požadované výsledky boli stále presnejšie (a modely nášho správania autentickjšie).

A riziká i možné dôsledky sú v mnohom alarmujúce...

Neuvedomujúc si, ako je naša myseľ a psychika zraniteľná, tak postupne prechádzame od technologického prostredia založeného na systémoch AI k prostrediu založenému na závislosti a manipulácii.

Potenciál a skutočnú silu týchto psychologických zásahov, vyjadruje i skúsenosť ľudí pracujúcich v oblasti implementácie systémov AI v informačnej spoločnosti, osobitne v rámci sociálnych sietí. Mnohí z nich do detailov poznajú spôsob fungovania nasadených systémov AI, poniektorí z nich dokonca tieto systémy vyvíjali a nasadzovali, no sami sa stali obeťami ich fungovania. Otvorene hovoria o obdobiach svojho života, v ktorých zakúsili totálnu neschopnosť sa vymaniť z ich vplyvu a závislosti.²²⁸

227 Už v roku 2010 Erich Smid, v tom čase CEO Google, na konferencii Techonomy v kalifornskom Lake Tahoe vyhlásil: „V súčasnosti každé dva dni vyprodukuje toľko informácií, koľko sme vytvorili od úsvitu civilizácie až do roku 2003.“

SIEGLER, *Eric Schmidt: Every 2 Days We Create As Much Information As We Did Up To 2003* [on-line]. [cit. 12. decembra 2021].

Dostupné na internete: <<http://techcrunch.com/2010/08/04/schmidt-data/>>

228 Napríklad Tim Kendall, bývalý vedúci pracovník Facebooku a následne prezident Pinterestu, hovorí o svojej vlastnej závislosti na technológiách Pinterestu, ktorých vývoj a zavádzanie sám viedol. Vedel o tom, presne poznal celú psychológiu a mechanizmy pôsobenia, no nevedel si v tom čase pomôcť. „Klasická irónia: počas dňa pracujem na niečom a tvorím niečo, čomu napokon sám podľahnem.“ „Je zaujímavé, že i napriek tomu, že som vedel, čo sa deje v zákulisí, nedokázal som svoje používanie sociálnych médií ovládať.“

Dr. Anna Lembke, lekárska riaditeľka liečby závislostí na Stanfordskej univerzite, poukazuje na potenciál závislosti, ktorú môžu vytvárať moderné sociálne siete, postavené na systémoch AI a dostatočne sofistikované, aby virtualizovali a nahrádzali naše reálne vzťahy: „Sociálne médiá sú droga. Máme základnú biologickú potrebu byť v kontakte s inými ľuďmi, čo má priamy vplyv na uvoľňovanie dopamínu v mezilimbickej dráhe. Tento systém, ktorý majú na svedomí milióny rokov evolúcie, sa nás snaží spájať a núti nás žiť v komunitách, aby sme si mohli nájsť partnerov a množiť sa. Takže niet pochyb o tom, že niečo ako sociálne médiá, ktoré umožňujú spojenie medzi ľuďmi, budú mať potenciál vyvolať závislosť.“²²⁹

So závislosťou prichádzajú aj ďalšie súvisiace problémy. Algoritmy sociálnych médií sú postavené na neustálych podnetoch, ktoré súvisia s uvoľňovaním dopamínu a pozitívnou reakciou našej mysle. Vytvára sa tak túžba po neustálom pozitívnom ohodnotení, ktoré už nie je primeranou a hlavne adekvátnou súčasťou vnemov v reálnom svete, ale umelo držanou hladinou, na ktorej sa používatelia stávajú závislí. Znižuje sa tak schopnosť prijímať a adekvátne spracovávať negatívne reakcie a prijať aj svoje slabé stránky. A v konfrontácii s realitou alebo s negatívnou skúsenosťou prichádza k psychickým problémom a kolapsu.²³⁰ Ľudská bytosť je určitým spôsobom konfrontovaná s neľudským perfekcionizmom algoritmov AI, ktorý je na hony vzdialený od empatie, emócií a nedokonalostí človeka, na ktoré však dokáže katastrofálne pôsobiť.²³¹

ORLOWSKI, *The Social Dilemma*, 31:20, 31:55. [cit. 16. februára 2022].

Dostupné na internete: <<https://www.netflix.com/sk/title/81254224>>

229 ORLOWSKI, *The Social Dilemma*, 33:15 [cit. 16. februára 2022].

Dostupné na internete: <<https://www.netflix.com/sk/title/81254224>>

230 Podľa Dr. Jonathana Haidta, sociálneho psychológa z New York University, behom posledných desiatich rokov, ktoré sa prekrývajú s celosvetovým rozšírením využívania sociálnych sietí mladými ľuďmi a deťmi, prišlo k enormnému nárastu prípadov depresíí, úzkosti a samovrážd medzi mladistvými a deťmi.

Napr. u dievčat (ktoré sú viac náchylné na ohodnotenie viditeľnej dokonalosti) nasledovne: u starších dievčat (15-19) stúplo seba poškodzovanie o 62%, u mladších (10-14) o 189%! Počet samovrážd narástol u starších dievčat o 70% a u mladších o 151%.

Ide o generáciu mladých, ktorí sa ako prví stretli so sociálnymi sieťami už na základnej škole, resp. skôr.

ORLOWSKI, *The Social Dilemma*, 40:00 [cit. 17. februára 2022].

Dostupné na internete: <<https://www.netflix.com/sk/title/81254224>>

231 Zaujímavé myšlienky a implikácie na tému dokonalosti AI uvádza článok Umelá inteligencia nesmie nahradiť nedokonalosť ľudskej empatie.

ADIB-MOGHADDAM, A. *Artificial intelligence must not be allowed to replace the imperfection of human*

Naviac, perimeter závislosti od pozitívneho hodnotenia nie je ohraničený reálnymi osobami z môjho okolia, ale rozširuje sa na tisíce až desiatky tisíc virtuálnych osôb, ktorých hodnotenie má na našu psychiku enormný dosah. Ľudský mozog sa nevyvinul tak, aby bol schopný prijímať spoločenské ohodnotenia v minútových intervaloch a od nesmierného množstva iných subjektov. **Pod intenzívnym vplyvom sociálnych sietí a virtuálneho sveta sa ľudia stávajú závislí na svojej viditeľnej dokonalosti a neustálom prísune krátkodobých signálov odmeňovania až do tej miery, že si to spájajú s hodnotami a s pravdou.**²³² Skutočne hodnotné a pravdivé sa tak stáva to, čo prináša najviac pozitívnych hodnotení a čo ľudí udržiava v kontrolovanom pozitívnom stave krátkodobých signálov odmeňovania. No nie je to zadanie ako vyšité pre systémy umelej inteligencie, ktoré bežia na pozadí najrozšírejších sociálnych sietí?

Vzhľadom na bezprecedentnú rozšírenosť týchto sociálnych sietí (a tým aj prakticky najrozšírejšie využívanie systémov AI, ktoré interagujú s človekom a učia sa na reálnych dátach ľudí) ide o zásahy, ktoré ovplyvňujú celú modernú spoločnosť. Častokrát sa objavuje argument, že ide o krátkodobý výkyv, ktorý sprevádzal nástup akejkoľvek pokrokovej technológie v dejinách, resp. príchod ďalšej priemyselnej technológie a veľkej spoločenskej zmeny, a že je len otázkou času, kým si na to ľudstvo privykne. Dovoľme si tvrdiť – spolu s mnohými odborníkmi na etiku sociálnych sietí a informačnú spoločnosť – že v tomto prípade je to iné a to z dvoch dôvodov.

Realitou súčasnej digitálnej éry a ostatnej priemyselnej revolúcie je jej časová neohraničenosť – nachádzame sa v období permanentnej technologickej revolúcie, ktorá nie je uzavretým procesom, ktorého dôsledky spoločnosť vstrebáva. Ľudstvo je skôr vystavené neustálemu dopingu technológií, informácií a vplyvov systémov AI, ktorého ovocím je disproporcia medzi technikou a schopnosťou spoločnosti ju správne a bezpečne využívať.²³³

empathy. [on-line]. [cit. 20. februára 2022].

Dostupné na internete: <<https://theconversation.com/artificial-intelligence-must-not-be-allowed-to-replace-the-imperfection-of-human-empathy-151636>>

232 Por. vyjadrenia Chamatha Palihapatiyu, bývalého viceprezidenta Facebooku pre rast. Svojho času bol priekopníkom novátorských stratégií rastu, vďaka ktorým Facebook neskutočne rástol a ktoré sú v súčasnosti využívané naprieč technologickými firmami celého Silicon Valley.

ORLOWSKI, *The Social Dilemma*, 39:20 [cit. 17. februára 2022].

Dostupné na internete: <<https://www.netflix.com/sk/title/81254224>>

233 Por. ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi*,

Druhým dôvodom je fakt, že informačné technológie stojace za týmito systémami sa nielen permanentne, ale predovšetkým exponenciálne zlepšujú. Na jednej strane interakcie je tak ľudský mozog, ktorý sa ďalej nevyvíja a na strane druhej technológie, ktoré sa za posledných šesťdesiat rokov mnohonásobne zmenili a porástli (len napríklad výkon počítačov sa za ten čas zvýšil v násobkoch miliárd).²³⁴

V konfrontácii uvedeného sa tak dotýkame niečoho, čo sme už skôr spomenuli ako technologickú singularitu²³⁵ a čo je pre niektorých potvrdením obáv a výziev, ktoré obnáša koncept transhumanizmu²³⁶, konkrétne vytvorením nových druhov ľudí genetickými manipuláciami alebo integráciu systémov AI a človeka pre eliminácie dôsledkov činnosti moderných systémov AI v službe sociálnych sietí a schopnosti spoločnosti ich asimilovať a integrovať.

Mnohí odborníci, pohybujúci sa v oblasti umelej inteligencie, podvedome čakajú na moment, keď umelá inteligencia premôže ľudskú silu a inteligenciu. Inak povedané, kedy nastane vek uvedomelej AGI, ktorá nás dokáže nahradiť v práci a bude múdrejšia

[on-line], s. 25, 29-30. [cit. 17. februára 2022].

Dostupné na internete:

<https://peter.santavy.cloud/docs/Analiza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

234 Por. Randima Fernando, spoluzakladateľ Centra pre humánne technológie, ORLOWSKI, *The Social Dilemma*, 45:00 [cit. 17. februára 2022].

Dostupné na internete: <<https://www.netflix.com/sk/title/81254224>>

Por. ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi*, [on-line], s. 25, 29-30. [cit. 17. februára 2022].

Dostupné na internete:

<https://peter.santavy.cloud/docs/Analiza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

235 Pod termínom technologická singularita sa myslí teoretický bod vo vývoji vedeckej civilizácie (znalostnej spoločnosti), v ktorom sa technologický pokrok zrýchli do nekonečna a prevýši všetky predpovede.

Singularita [on-line]. [cit. 3. augusta 2020].

Dostupné na internete: <<https://sk.wikipedia.org/wiki/Singularita>>

Singularitou v umelej inteligencii je myslený stav, ktorý nastane, ak umelá inteligencia so schopnosťou učiť sa, zlepšovať sa a samostatne sa vyvíjať rýchlo dosiahne a následne prevýši inteligenciu človeka.

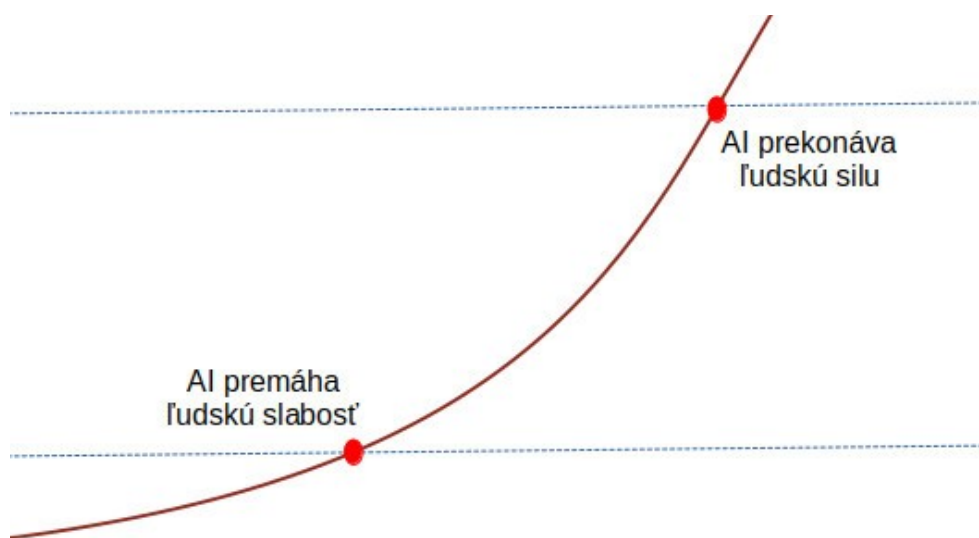
Teda situácia, keď sa počítačové systémy stanú inteligentnejšími než ľudia.

por. MITCHELL, *Artificial Intelligence*, s. 10.

236 O transhumanizme v kontexte umelej inteligencie sme pojednávali v kapitole 1.10.

než my. Stále sa domnievame, že táto otázka nie je na programe dňa a aktuálne je viac súčasťou pracovnej náplne futuroológov.

Skôr sa – na základe uvedených rizík v tejto kapitole – stotožňujeme s pohľadom Tristana Harrisa, bývalého etika dizajnu vo firme Google a spoluzakladateľa Centra pre humánne technológie, pre ktorého je oveľa dôležitejší ten moment, v ktorom technológia prekoná a ovládne ľudské slabosti. Už vtedy prichádza víťazstvo AI a porážka ľudstva, lebo už vtedy prichádza závislosť, polarizácia a radikalizácia spoločnosti, zaslepenosť, strata schopnosti komunikovať a hľadať pravdu,... jednoducho prehráva všetko ľudské v nás...²³⁷



Obr. č. 10. Dva pohľady na singularitu v oblasti umelej inteligencie

Míľníkom, ktorého by sme sa mali obávať, teda 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.

V úvode tejto kapitoly sme spomenuli, že nutnou súčasťou extrémne rýchleho súčasného vývoja v oblasti umelej inteligencie je fungujúca možnosť monetizácie nasadenia systémov AI. Pri riešení dôsledkov činnosti systémov AI je treba stále pamätať na to, aké ciele sú zadané pre činnosť týchto systémov. Implementácia algoritmov s optimalizáciou na maximalizáciu zisku sa na sociálnych sieťach premieta do predlžovania času stráveného na sieti, počtu zhladaných reklám, množstva dát, ktoré používateľ vyprodukuje a ktoré je následne možné nejakým spôsobom premeniť na finančný benefit.

237 ORLOWSKI, *The Social Dilemma*, 53:35 [cit. 17. februára 2022].

Dostupné na internete: <<https://www.netflix.com/sk/title/81254224>>

Samozrejme, takáto optimalizácia prináša aj svoje dôsledky: sociálne bubliny, v rámci ktorých sú podsúvané len tie informácie, ktoré korešpondujú s pohľadom používateľa a ponúkanie kontaktov na osoby rovnakého razenia, uprednostňovanie falošných informácií, pretože tie viac udržia používateľa pripojeného na sieti, ergo viac zarábajú.

Pre zadané ciele dnešné sociálne siete vytvárajú prostredie (t.j. algoritmy AI takto reagujú na zadané požiadavky a vylepšujú ponúkané informácie), ktoré je toxické - odtrhnuté od reality, s eróziou hodnôt, postmodernou²³⁸ rezignáciou na hľadanie pravdy, zámenou dialógu za konfrontáciu,... Dôsledkom sociálnych sietí poháňaných sofistikovanými systémami AI je deštruktívne ovocie pre život a rozvoj človeka, kvalitu jeho života, vzťahy i celú spoločnosť. Jednoducho povedané – **zlyháva etický rozmer nasadenia AI**.

Ak by sme sa na zmenu smerujúcu k etickému nasadeniu neodhodlali, vyhliadky môžu byť veľmi pochmúrne.²³⁹

- extrémne rozdelenie spoločnosti a narušenie vzťahov
- nárast zatvrdilej nevedomosti a ignorancie faktov
- neschopnosť riešiť aktuálne civilizačné výzvy
- premena demokracií na autokratické a disfunkčné celky
- zničenie globálnej ekonomiky a celosvetového spoločenstva

238 Treba si uvedomiť, že **informačný svet, zasadený do postmoderny s jej relativizáciou pravdy, rezignáciou na racio a akcentovaním emócií a zážitkov, je jedným z kľúčov k chápaniu výziev, rizík a potenciálu, ktorý v sebe vyššie uvedené problémy obnášajú.**

Virtuálny svet a sociálne siete sú tak dejiskom i katalyzátorom prerastania modernistickej racionalizácie etiky a morálnych hodnôt do postmodernistickej rezignácie na racio. Stráca sa schopnosť vnímať a hodnotiť veci i udalosti racionálne a objektívne. Čoraz väčší dôraz sa kladie na formu a zážitok, resp. emócie, ktoré sú prostredníctvom informačných technológií v mediálnom svete navodené. Človek následne častokrát hodnotí a rozhoduje sa pudovo alebo podvedome, povrchné, len na základe pozitívnej, či negatívnej emócie.

I keď badať reálny odstup od racionálneho vnímania, rozvíjanie postmoderny práve v rámci virtuálneho sveta a prostredníctvom prostriedkov IKT a systémov AI navodzuje falošnú atmosféru pseudovedeckých a pseudo hodnotových postojov, čo je živnou pôdou pre akékoľvek nekalé pôsobenie a spoločenskú, obchodnú i psychologickú manipuláciu s ľuďmi.

ŠANTAVÝ, P. *Niektoré výzvy informačnej spoločnosti v oblasti morálnej teológie*. In: *Doctorandum dies 2018 : Varia historia et moralia*. Bratislava: RKCMBF UK, 2018, s. 117-135.

239 ORLOWSKI, *The Social Dilemma*, 1:20:25 [cit. 17. februára 2022].

Dostupné na internete: <<https://www.netflix.com/sk/title/81254224>>

- potencionálne smerovanie k civilizačnému kolapsu...

Ako navrhované riešenie sa podáva ponuka nových algoritmov umelej inteligencie a nástrojov, ktoré to vyriešia za nás – ako napr. uviedol Mark Zuckerberg, zakladateľ a CEO Facebooku (v súčasnosti Meta Platforms) pri vypočúvaní v Senáte v kauze ovplyvňovania prezidentských volieb v USA. V odbornej komunite však prevláda voči tomuto názoru skepsa – **vzhľadom na monetizačné zadanie, ktoré sa takmer vôbec v súčasnosti nedá zmeniť a schopnosti súčasných systémov AI (AI sama nevie rozoznať pravdu a mať etické mantinely) to musíme zmeniť my, ľudia.**

Ďalším dôvodom, prečo je táto zmena výzvou a povinnosťou pre človeka a spoločnosť, je skutočnosť, že **technológie samy o sebe nie sú tou priamou hrozbou - dokážu však prebudiť v človeku aj v spoločnosti tie najhoršie stránky ľudského bytia, ktoré sa tak stávajú existenčnou hrozbou.**²⁴⁰

Dôsledky sa teda dotýkajú celej spoločnosti a doliehajú aj na ľudí, ktorí sociálne siete nevyužívajú. Či už ide o znášanie „ovocia“ činov tých, ktorí sú týmito systémami priamo ovplyvnení, alebo sa stávajú obeťami spracúvania a využívania svojich osobných údajov zo sekundárnych informácií, ktoré o nich nazhromaždili automatizované systémy.

Znovu pripomíname, že pokiaľ bude existovať obchodný model, pre ktorý systémy AI v službe človeku a spoločnosti zlyhávajú, samo od seba sa to nezmení. V kontexte slabej umelej inteligencie (ANI) ešte stále konečnú voľbu cieľov nerobia stroje, ale človek, takže stále je v ľudskej moci tieto dôsledky ovplyvniť a zmeniť. Preto **musia existovať etické pravidlá a regulácie, ktoré jednotlivca i celú spoločnosť budú voči rizikám sofistikovanej práce systémov AI v rámci sociálnych sietí chrániť. Ak nechceme smerovať k dystopii, musíme sa preto na všetkých úrovniach angažovať a pracovať na tom spoločne ako ľudstvo.**

V tomto duchu sa javia ako relevantné aj požiadavky:²⁴¹

- aby produkty a ciele, na ktoré sa systémy umelej inteligencie využívajú, boli humánne, t.j. pre dobro človeka

240 ORLOWSKI, *The Social Dilemma*, 1:18:10 [cit. 17. februára 2022].

Dostupné na internete: <<https://www.netflix.com/sk/title/81254224>>

241 ORLOWSKI, *The Social Dilemma*, 1:28:05 [cit. 17. februára 2022].

Dostupné na internete: <<https://www.netflix.com/sk/title/81254224>>

- aby tieto systémy a ich zadávatelia sa k nám nesprávali ako ku zdrojom, ktoré možno využívať a ťažiť z nich

Viackrát sme sa v tejto kapitole okrajovo dotkli aj presahu k systémom uvedomelej AGI. K tomu ešte môžeme uviesť, že **vo vízii uvedomelej umelej inteligencie sa v odbornej komunite diskutuje aj ďalší špecifický rozmer závislosti – špecifické zameranie sa a naviazanie sa na pokročilú umelú inteligenciu ako na osobu.**

Okrem zhromažďovania a spracovania extrémneho množstva reálnych dát o človeku, smerovania sociálnych sietí a závislostí spomeňme ešte niekoľko ďalších dôsledkov, na ktoré treba pamätať...

Jednou z ďalších zaujímavých oblastí aplikácie umelej inteligencie je vytváranie umelých virtuálnych identít (avatarov), využívaných napríklad ako virtuálni hlásatelia v televíznych reláciách, alebo virtualizované identity reálnych hercov použité pri digitálnom nakrúcaní extrémnych scén moderných filmov. Schopnosti systémov AI v tejto oblasti sú obdivuhodné – ved' takto dokážu „rozpohybovať“ nejednu súčasnú osobnosť, či vytvoriť umelé postavy na nerozoznanie od reálnych aktérov filmových dejov.

Nezávisle od tejto schopnosti tiež badať rozmach celej oblasti falošných (fake) obrázkov a videí zobrazujúcich reálne existujúce osoby v situáciách, na ktorých nemali účasť. Osobitne sa tento nešvár šíri v oblasti britkej internetovej zábavy, pornografie a v poslednom čase aj propagandy. Falošné fotografie a videá v tejto oblasti sú tým reálnejšie, čím viac zdrojov o danej osobe je k dispozícii. Napr. z jednej snímky bežného človeka zachytenej bezpečnostnou kamerou je veľmi náročné vytvoriť falošnú video scénu. Ale v prípade verejne známych osôb, napr. hercov alebo politikov, ktorých fotografií a videí sú plné profily sociálnych sietí, nie je problém, aby sa systém AI dostatočne natrénoval a vytvoril vynikajúce falzifikáty.

Ak sa táto schopnosť spojí so skôr uvedenými systémami umelej inteligencie, ktoré dokážu vytvárať modely ľudského správania, tvoriť psychologické profily a predikovať správanie konkrétnych osôb, vo svete, ktorý sa stále viac orientuje na on-line komunikáciu a kybernetický priestor, prichádzajú na scénu mocné nástroje (deep fake systémy) na **vytváranie dôveryhodných virtuálnych identít ľudí a podvrhnutie ich falošnej komunikácie a konania** (samozrejme konania, ktoré sa nikdy reálne neuskutočnilo).²⁴²

242 Napr. obsah prejavu, ktorý by daný politik nikdy nepovedal; falošné videozábery udalostí, v ktorých vystupujú reálne osoby, ktoré však v skutočnosti nikdy aktérmi takej udalosti neboli; a pod.

Deep fake systémy sa podieľajú na stieraní hranice medzi pravdou a lžou, medzi realitou a fiktívnym svetom.

Ďalšou oblasťou dôsledkov nasadenia systémov umelej inteligencie sú právne aspekty súvisiace s ich využívaním.

Už doterajšia rozprava o závislostiach, zbere dát, monitorovaní, deep fake systémoch, atď. v sebe implicitne zahŕňala otázku legálnosti týchto systémov AI (napr. rozsah zhromažďovaných dát, podprahové ovplyvňovanie ľudí, atď.).

Z viacerých súčasných aktivít na tomto poli sa asi najďalej nachádza Európska únia so svojím regulačným rámcom zahŕňajúcim etické a legislatívne aspekty využívania systémov umelej inteligencie. Keďže ide asi o najdôležitejšiu reguláciu AI súčasnosti, osobitne sa jej budeme venovať v tretej kapitole.

Klasickým príkladom právnych otázok spojených s využívaním umelej inteligencie v reálnom živote sú požiadavky kladené na autonómne vozidlá a zodpovednosť v prípade zapríčinenia dopravnej nehody.

Z pohľadu autonómie rozlišujeme päť stupňov – od vozidiel vybavených asistenčnými službami až po plnú autonómiu riadenia, pričom až posledné dva stupne sú považované za tak samostatné, že pri riadení nevyžadujú priamu pozornosť tzv. bezpečnostného operátora (safety operator), t.j. človeka, ktorý dozoruje autonómny systém riadenia (viď obrázok č.11).

V súčasnosti sa diskutujú prvé návrhy na legislatívne požiadavky na bezpečnostného operátora, hľadajúc objektívnu mieru zodpovednosti ľudského faktora pri riadení autonómneho vozidla v rámci jednotlivých stupňov automatizácie.

Rieši sa dilemma trénovania systémov AI plne autonómnych vozidiel pre osobitné kritické situácie, v ktorých sa musí autopilot rozhodnúť medzi potencionálnym ohrozením života posádky alebo ostatných účastníkov v čase blížiacej sa alebo prebiehajúcej dopravnej nehody.

Vzhľadom na extrémnu technologickú komplexnosť autonómnych vozidiel vyvstáva aj otázka osobitných technologických noriem pre ich systémy AI a potrebnej cestnej (elektronickej) infraštruktúry (čo sme diskutovali v kapitole 2.4.) i akceptovanej miery

It's Getting Harder to Spot a Deep Fake Video. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.youtube.com/watch?v=gLoI9hAX9dw>>

neistoty, ktorá sprevádza autonómne systémy AI vzhľadom na limity a rizikové faktory, ktoré sme uviedli v kapitole 2.1.

Pre cestné vozidlá		 Ľudský vodič	 Automatizovaný systém		
		Riadenie a zrýchľovanie/ znižovanie rýchlosti	Monitoro- vanie jazdného prostredia	Záložný systém v prípade zlyhania automatizácie	Automati- zovaný systém má kontrolu
Ľudský vodič monitoruje cestu	0 BEZ AUTOMATIZÁCIE				–
	1 ASISTENCIA VODIČA				NIEKTORÉ REŽIMY JAZDY
	2 ČIASTOČNÁ AUTOMATIZÁCIA				NIEKTORÉ REŽIMY JAZDY
Automatizovaný systém riadenia monitoruje cestu	3 PODMIENENÁ AUTOMATIZÁCIA				NIEKTORÉ REŽIMY JAZDY
	4 VYSOKÁ AUTOMATIZÁCIA				NIEKTORÉ REŽIMY JAZDY
	5 ÚPLNÁ AUTOMATIZÁCIA				

Obr. č. 11. Päť úrovní automatizácie vozidla.²⁴³

K diskusii dôsledkov využívania systémov umelej inteligencie pre spoločnosť treba doplniť oblasť, ktorá síce nie je na výslni, no je o to dôležitejšia – oblasť nasadenia algoritmov a systémov umelej inteligencie v spravodajských službách a v armáde.

²⁴³ AMUNDSEN, M. *The 5 levels of vehicle automation*. Twitter.

Dostupné on-line <<https://twitter.com/mamund/status/1171089135089700865>>, upravené autorom.

2.6. Umelá inteligencia prichádza do spravodajských služieb

Presnejšie povedané – umelá inteligencia v spravodajských službách už dávno nie je žiadny zelenáč a s technologickým rozvojom pracuje na svojom kariérnom postupe.

Umelá inteligencia nachádzala svoje uplatnenie v analytických sekciách spravodajských služieb už v osemdesiatych rokoch minulého storočia. V tom čase išlo primárne o využívanie symbolických systémov AI²⁴⁴, reprezentovaných nástrojmi dolovania údajov (data mining²⁴⁵) a z dnešného pohľadu jednoduchého spracovania spravodajského obrazového materiálu. Primárne išlo o dáta získavané odposluchom z monitorovacích staníc rôzneho typu (od telefónnych operátorov, cez systémy rádio-elektronického boja až po Echelon²⁴⁶) a z fotografií získavaných prieskumným letectvom a družicami. Neskôr – s rozvojom využívania internetu – sa pripojilo i monitorovanie globálnej počítačovej siete (išlo napr. o systém Carnivore²⁴⁷), ktoré postupne naberalo na dôležitosti úmerne tomu, ako sa v rámci rodiacej informačnej spoločnosti internet pretváral na virtuálny kybernetický priestor on-line života.

Preukázateľný profit z týchto systémov, ktorý presahoval vojenské a spravodajské záujmy

244 Len pripomeňme, že symbolické systémy na základe definovaných pravidiel a postupov spracúvajú jednotlivé symboly (pojmy, slová, frázy,...) a vykonávajú priradené úlohy. Veľmi zjednodušene povedané – pomocou matematickej logiky sa snažia emulovať procesy myslenia. Symbolické systémy sme popisovali v kapitole 1.4.

245 Data mining (dolovanie dát) je analytický proces navrhnutý na skúmanie veľkého množstva dát podľa konkrétnych vzorov a väzieb. Výsledkom sú množiny dát, splňujúcich žiadané podmienky, prípadne predikcia ďalšieho vývoja, správania, atď.

Data Mining Techniques [on-line]. [cit. 7. decembra 2015].

Dostupné na internete: <<http://documents.software.dell.com/statistics/textbook/data-mining-techniques>>

246 Ide o globálny systém pre zhromažďovanie a vyhodnocovanie odpočúvaných dát (SIGINT).

Z pôvodného sledovania vojenských a diplomatických cieľov počas studenej vojny sa ku koncu 20. storočia rozvinul do celosvetového systému na zachytávanie i súkromnej i komerčnej komunikácie a stal sa tak nástrojom masívneho dohľadu a komplexnej špionáže.

Echelon [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <<https://en.wikipedia.org/wiki/ECHELON>>

247 Carnivore bolo odpočúvacie a monitorovacie zariadenie FBI zamerané na emailovú a internetovú komunikáciu. Aktívne využívané bolo v rokoch 1997 až 2005, následne bolo nahradené výkonnejšími a globálnymi dohľadovými systémami.

Carnivore [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <[https://en.wikipedia.org/wiki/Carnivore_\(software\)](https://en.wikipedia.org/wiki/Carnivore_(software))>

a prinášal prvé ovocie aj v oblasti priemyselnej špionáže a globálneho dohľadu, dal vyrásť celému ekosystému spravodajských informačných nástrojov a ich pevnému zakoreneniu v spravodajských agentúrach. Súbežne s tým prichádzalo k prekračovaniu limitovaných technologických možností a lokálnych záujmov, čo dalo vznik globálnemu hromadnému dohľadu, t.j. hromadnému dohľadu nad celými populáciami a naprieč širokým svetom.²⁴⁸

Významným medzníkom pre celoplošné a detailné sledovanie boli dôsledky tragického teroristického útoku na Dvojčky v New Yorku z 11. septembra 2001. Podľa riaditeľa NSA, admirála Michaela S. Rogersa „útoky z 11. septembra podnietili 'zásadnú zmenu' v spôsobe, akým americká vláda využíva a zdieľa spravodajské informácie“.²⁴⁹

Bezpečnosť dostala absolútnu prioritu pred právom na súkromie: USA Patriot Act prijatý v októbri 2001, účasť nadnárodných digitálnych korporácií na sledovacích programoch,²⁵⁰ kontroly počítačovej techniky rôznych oficiálnych predstaviteľov štátu pracovníkmi CIA (2014), odhalenie sofistikovaných nástrojov používaných na prienik do smartfónov, počítačov a dokonca aj do „smart“ televíznych prijímačov pripojených na internet,²⁵¹ integrácia systémov umelej inteligencie na rozpoznávanie tvárí (Clearview AI) v roku 2020, atď.²⁵²

Zmena priority zo súkromia na národnú bezpečnosť v USA spustila celý sled udalostí,

248 Okrem Echelonu a Carnivore môžeme hovoriť o takých systémoch, ako sú XKeyscore, PRISM, Dishfire, Stone Ghost, Tempora, Frenchelon, Fairview, MYSTIC, DCSN, Boundless Informant, Bullrun, Pinwale, Stingray, SORM, RAMPART-A, Mastering the Internet a Jindalee Operational Radar Network. *Global surveillance* [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Global_surveillance>

249 GARAMONE, J. *9/11 Drove Change in Intelligence Community, NSA Chief Says*. [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <<https://www.defense.gov/News/News-Stories/Article/Article/945544/911-drove-change-in-intelligence-community-nsa-chief-says/>>

250 Microsoft sa údajne v roku 2007 stal prvou veľkou technologickou spoločnosťou, ktorá spolupracovala na programe elektronického sledovania občanov PRISM. Neskôr sa mali pripojiť i Yahoo, Google a Facebook.

251 *WikiLeaks Releases Trove of Alleged C.I.A. Hacking Documents*. [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <<https://www.nytimes.com/2017/03/07/world/europe/wikileaks-cia-hacking.html>>

252 *How 9/11 Changed the Course of Personal Data Collection and Surveillance*. [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <<https://www.startpage.com/privacy-please/startpage-articles/how-9-11-changed-the-course-of-personal-data-collection-and-surveillance>>

ktoré sa ako domino efekt šírili aj v ostatných krajinách: implementácia moderných sledovacích systémov štátnymi orgánmi, ich využívanie na lokálnej úrovni samospráv i v komerčnej sfére a – čo je najhoršie – bol umožnený neriadený prístup k týmto technológiám aj silovým a spravodajským zložkám autoritárskych režimov a diktatúr.

Súčasný využitie systémov umelej inteligencie v dohľadových systémoch a pri sledovaní je ovplyvnené dvomi skutočnosťami:

- spoločenským tlakom a zmenou politickej klímy po kauze Snowden vs. NSA
- veľkým pokrokom v rozvoji systémov AI ako systémov hlbokého strojového učenia

Edward Snowden je americký whistle-blower²⁵³, ktorý v roku 2013 dal novinárom denníka The Guardian k dispozícii dokumenty o sledovacích aktivitách americkej National Security Agency (NSA) či britskej Government Communications Headquarters (GCHQ). Dovtedy pracoval pre NSA na rôznych pozíciách ako špecialista na informačné technológie a spoznajúc praktiky s nelegálnymi dohľadovými systémami sa rozhodol tieto informácie z morálnych dôvodov zverejniť.²⁵⁴

Využívajúc šifrovanie, prostriedky darknetu²⁵⁵ a profesionálne krytie, Snowden sa tajne stretol s redaktormi The Guardian, ktorí 5. júna 2013 zverejnili prvý z článkov založených na jeho odhalení.²⁵⁶

253 Výraz whistleblowing sa označuje zverejnenie nekalých, resp. nelegitímnych praktík spoločností alebo štátnych inštitúcií interným zamestnancom. Ide buď o zverejnenie mediálne, alebo o upozornenie konkrétnych štátnych orgánov.

Por. *Whistleblowing* [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <<https://cs.wikipedia.org/wiki/Whistleblowing>>

254 „Uvedomil som si, že som súčasťou mašinérie, ktorá spôsobuje viac zlého než dobrého.“

Profile: Edward Snowden. In: BBC. [on-line]. 2013, 24. 6. [cit. 23. februára 2022].

Dostupné na internete: <<http://www.bbc.co.uk/news/world-us-canada-22837100>>

255 Darknet je technologicky izolovaná časť internetu a IKT vo všeobecnosti, na ktorej je kladený dôraz na anonymitu a bezpečnosť. Jej možnosti sú najčastejšie využívané buď na ochranu pred sledovaním, alebo na kriminálne aktivity. Darknet je častokrát spojovaný s jeho obsahom (Darkweb) a skrytými službami (Tor services).

Por. *Darknet* [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <<https://en.wikipedia.org/wiki/Darknet>>

256 GIDDA, M. *Edward Snowden and the NSA files – timeline* [on-line]. The Guardian. [cit. 23. februára 2022].

Dostupné na internete: <<http://www.guardian.co.uk/world/2013/jun/23/edward-snowden-nsa-files->

V priebehu nasledovných mesiacov boli z množstva tajných informácií, ktorými Snowden disponoval, zverejnené nasledovné kauzy:²⁵⁷

- tajný súdny príkaz, ktorým americká vláda donútila telekomunikačnú firmu Verizon Communications sprístupniť milióny záznamov telefonických hovorov²⁵⁸
- tajný národný bezpečnostný program elektronického sledovania občanov PRISM, ktorý NSA prevádzkovala v USA od roku 2007²⁵⁹
- odhalenie tajného programu Tempora z dielne britskej spravodajskej služby GCHQ, ktorý od roku 2012 neustále sleduje a aj fyzicky ukladá telefónne hovory (max. 30 dní), internetové informácie a metadáta ľubovoľných osôb na svete, t.j. priamo ich telefónne hovory, obsahy ich e-mailových správ, ich príspevky na Facebooku, zoznam nimi navštívených stránok atď. GCHQ získané dáta a analýzy poskytuje aj NSA.²⁶⁰
- odhalenie odpočúvania čínskych hovorov a sms agentúrou NSA²⁶¹
- odhalenie odpočúvania a napojenia sa na počítačové siete, ktoré NSA konalo v budovách diplomatického zastúpenia EÚ a OSN vo Washingtone a pri OSN. V dokumentoch z roku 2010 sú Európania explicitne uvedení ako cieľ sledovacieho útoku.²⁶²

timeline>

ŠANTAVÝ, *Niektoré výzvy informačnej spoločnosti v oblasti morálnej teológie*, s. 117-135.

257 ŠANTAVÝ, *Niektoré výzvy informačnej spoločnosti v oblasti morálnej teológie*, s. 117-135.

258 Vzhľadom na ústavu USA a jej dodatky išlo o nezákonné a nelegitímne konanie voči občanom USA.

GIDDA, *Edward Snowden and the NSA files – timeline* [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <<http://www.guardian.co.uk/world/2013/jun/23/edward-snowden-nsa-files-timeline>>

259 Por. *PRISM (NSA)* [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <[https://sk.wikipedia.org/wiki/PRISM_\(NSA\)](https://sk.wikipedia.org/wiki/PRISM_(NSA))>

260 Por. *Tempora* [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <<https://sk.wikipedia.org/wiki/Tempora>>

261 LAM, L., CHEN, S. *US spies on Chinese mobile phone companies, steals SMS data: Edward Snowden* In: *South China Morning Post*. [on-line]. 2013, 23. 3. [cit. 23. februára 2022].

Dostupné na internete: <<https://www.scmp.com/news/china/article/1266821/us-hacks-chinese-mobile-phone-companies-steals-sms-data-edward-snowden>>

262 POITRAS, L., ROSENBACH, M., SCHMID, F., STARK, H. *NSA horcht EU-Vertretungen mit Wanzen aus* In: *Der Spiegel*. [on-line]. 2013, 29. 6. [cit. 23. februára 2022].

Snowdenove odhalenia boli sprevádzané veľkou mediálnou odozvou. To, pred čím niektorí odborníci z oblasti informačnej bezpečnosti pre rokom 2013 varovali a čo väčšina odbornej verejnosti i politikov považovala za prehnané či paranoidné, sa ukázalo ako pravdivé: **prostriedky a možnosti informačných technológií a systémov umelej inteligencie v spoločnosti, ktorá sa mení na informačnú, sú reálne zneužívané tajnými službami, vládnymi agentúrami a korporáciami.**²⁶³

Okrem morálneho a etického hodnotenia zverejnenia prísne tajných dokumentov spravodajských služieb USA sa doba po Snowdenovi snaží pragmaticky vyrovnať s odhalenými problémami a v súčasnosti sa už – technologicky, spoločensky i politicky – úplne inak pristupuje k ochrane súkromia a bezpečnosti dát. Jednak rastie povedomie o práve na súkromie, no súbežne sa vo viacerých krajinách stupňuje tlak na legalizáciu odhalených metód špehovania a monitorovania ľudí, aby konanie, ktoré bolo predtým tajné a v súčasnosti je mnohokrát technologicky blokované²⁶⁴, bolo legalizované a nariadené zákonom.²⁶⁵ Základným argumentom pre legalizáciu dohľadu a špehovania býva potreba národnej bezpečnosti (kybernetický terorizmus a terorizmus vo všeobecnosti, vojenská obrana štátu a boj s dezinformáciami) a boj s kriminalitou (detská pornografia, obchodovanie s ľuďmi a drogová činnosť).²⁶⁶

S rastom povedomia o ochrane súkromia i bezpečnostných hrozbách v spoločnosti a zlepšujúcou sa ochranou informačných prostriedkov i širokou dostupnosťou bezpečnostných nástrojov a návodov by sa mohlo zdať nebezpečenstvo zneužívania

Dostupné na internete: <<http://www.spiegel.de/netzwelt/netzpolitik/nsa-hat-wanzen-in-eu-gebaeuden-installiert-a-908515.html>>

POITRAS, L., ROSENBACH, M., STARK, H. *NSA überwacht 500 Millionen Verbindungen in Deutschland* In: *Der Spiegel*. [on-line]. 2013, 30. 6. [cit. 23. februára 2022].

Dostupné na internete: <<http://www.spiegel.de/netzwelt/netzpolitik/nsa-ueberwacht-500-millionen-verbindungen-in-deutschland-a-908517.html>>

263 ŠANTAVÝ, *Niektoré výzvy informačnej spoločnosti v oblasti morálnej teológie*, s. 117-135.

264 Zlepšujúca sa ochrana systémov voči sledovaniu a široká dostupnosť bezpečnostných nástrojov.

265 Svojho času vyvolal veľký rozruch návrh zákona v Austrálii, podľa ktorého môže byť vlastník zašifrovaného zariadenia odsúdený až na desať rokov väzenia, ak odmietne odšifrovať svoje zariadenie. Už v tom čase podobnú legislatívu mala napr. Veľká Británia. Vo viacerých krajinách je zasa povinnosťou cestujúceho pri vstupe do krajiny odblokovať svoje zariadenie a umožniť jeho preverenie úradom.

266 ŠANTAVÝ, *Niektoré výzvy informačnej spoločnosti v oblasti morálnej teológie*, s. 117-135.

technológií na sledovanie obyvateľov zažehnané (ak si teda odmyslíme v poslednom období rastúci tlak na legalizáciu).²⁶⁷

Opak je však pravdou – na scénu prichádza veľký pokrok v rozvoji systémov umelej inteligencie. Súčasné algoritmy hlbokého učenia totižto vedia sklbiť dôsledky spoločenského tlaku, v mnohých prípadoch pretaveného do legislatívnych obmedzení, s dostupnosťou extrémneho množstva dát, ktoré – ako inak – sú vytvárané ľuďmi, teda potencionálnymi subjektami sledovania. Tieto systémy využívajú dve skutočnosti:

- vďaka svojej schopnosti autonómne fungovať a adaptabilite²⁶⁸ tieto systémy dokážu spracúvať dáta, ku ktorým bez príkazu súdov neakceptujeme prístup osôb. To, že výsledná analýza a predložené závery sú prijímané s vysokou mierou dôveryhodnosti, je už len čerešnička na torte.
- spracúvanie metadát, ktoré netvoria priamo obsah komunikácie. Metadáta sú informácie o komunikácii – kto, kedy, komu, ako častokrát, v akej korelácii s inými udalosťami a osobami, atď.²⁶⁹

V roku 2014 bývalý riaditeľ CIA a NSA na The Johns Hopkins Foreign Affairs Symposium vyhlásil: „na základe metadát zabíjame ľudí“.²⁷⁰

Systémy umelej inteligencie, ktorým sa v súčasnosti analýza metadát zveruje, tak de facto majú rastúci potenciál rozhodovať o bytí a nebytí, „generovať“ dôkazy, ktoré môžu byť použité napríklad pri obvinení z vlastizrady, identifikácii podozrivých osôb, dokazovaní nelegálnej činnosti a pod.

Ak si uvedené spojíme s rizikami a limitmi komunikovanými v kapitolách 2.1. - 2.4. a okoreníme to niektorými spoločenskými dôsledkami z kapitoly 2.5., nezostáva nám nič

267 Radi by sme upozornili čitateľa na postupné rozširovanie záberu tejto kapitoly – od spravodajských služieb k sledovaniu a dohľadu vo všeobecnosti...

268 Vďaka patrí nielen moderným algoritmom a špecialistom na hyper parametre, no predovšetkým nám všetkým za poskytnuté množstvo údajov a dát:-)

269 Napr. pri emailovej komunikácii obsah emailu tvorí dáta. Metadátami sú napr. emailová adresa prijímateľa, čas odoslania emailu, ďalší prijímateľa na kópii, interval medzi prijatím emailu a odpoveďou, korelácia medzi emailovými komunikáciami viacerých sledovaných osôb, GPS súradnice miesta, z ktorého bol email odoslaný,....

270 „We kill people based on metadata.“

The Price of Privacy: Re-Evaluating the NSA [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <<https://www.youtube.com/watch?v=kV2HDM86Xgl&t=18m>>

iné, len **akcentovať veľkú opatrnosť a vyžadovať striktnú zákonnosť i dôsledný dohľad demokraticky zvolených zástupcov pri tomto druhu nasadenia systémov AI.**

V úvode tejto kapitoly sme trochu nadnesene spomenuli „kariérny postup“ umelej inteligencie v spravodajstve... Kombinácia extrémne veľkých objemov dát (big data) s modernými algoritmami umelej inteligencie v kontexte informačnej spoločnosti i napriek mnohým legislatívnym a regulačným obmedzeniam zabezpečuje systémom AI dôležité postavenie medzi prostriedkami spravodajských služieb a svojmu kariérnemu postupu dáva „zelenú“.

Ak by sme sa ponorili do sveta spracovania dát i metadát systémami AI, do sveta technológií priebežne rozpoznávajúcich tváre a identifikujúcich osoby, žasli by sme nad aplikačnými možnosťami, ktoré tento svet ponúka a možno by sme zistili, aké pokušenie to okrem legálnych a prospešných scenárov využitia prináša... A odtiaľ už je len krok k ešte väčšiemu rozšíreniu týchto systémov v spoločnosti, než sme uvádzali ako dôsledok zmien bezpečnostných priorít po páde Dvojčiek – či už v pozitívnom alebo v negatívnom slova zmysle.

V súčasnosti badať neustále rastúci počet štátov a korporácií, ktoré zavádzajú pokročilé nástroje umelej inteligencie umožňujúce monitorovanie, dohľadávanie a sledovanie občanov. Jednou z výpovedných štúdií popisujúcej vývoj v tejto oblasti je *The Global Expansion of AI Surveillance* a *AI Global Surveillance (AIGS) Index* z dielne nadácie Carnegie Endowment for International Peace.²⁷¹

Autor štúdie, v snahe uchopiť a mať možnosť riešiť dôsledky nasadenia pokročilých nástrojov sledovania využívajúcich umelú inteligenciu, sa zameriava na scenáre ich použitia, t.j. ako sa tieto nástroje nasadzujú a ako sa používajú. Výsledkom je Index globálneho dohľadu nad umelou inteligenciou (AIGS) zhromažďujúci empirické údaje o používaní dohľadu systémami AI v 176 krajinách sveta. Index i celá štúdia v zásade nerozlišujú medzi legitímnym a nelegitímnym používaním dohľadu systémami AI – cieľom je skôr ukázať, ako nové možnosti dohľadu menia schopnosť vlád monitorovať a sledovať jednotlivcov alebo systémy. Štúdia uvádza krajiny, ktoré využívajú technológiu dohľadu pomocou umelej inteligencie, sumarizuje konkrétne typy dohľadu, ktoré sú vládami

271 FELDSTEIN, S. *The Global Expansion of AI Surveillance*. [on-line]. [cit. 28. februára 2022].

Dostupné na internete: <<https://carnegieendowment.org/2019/09/17/global-expansion-of-ai-surveillance-pub-79847>>

zavádzané a vymenúva krajiny i spoločnosti dodávajúce túto technológiu.

Ku kľúčovým zisteniam štúdie²⁷² patrí prekvapivá rýchlosť šírenia implementácie dohľadových systémov AI v rámci širokého spektra štátov. Zo 176 vymenovaných krajín minimálne sedemdesiatpäť využíva tieto prostriedky naozaj aktívne a intenzívne.

K najviac používaným technológiám umelej inteligencie v dohľadových systémoch patria platformy tzv. inteligentných, resp. bezpečných miest (päťdesiatšesť krajín), systémy rozpoznávania tváre (šesťdesiatštyri krajín) a inteligentná polícia (päťdesiatdva krajín).

Hlavným aktérom na poli vývoja, ale i zavádzania týchto technológií do praxe je Čína. Jej ponuka je naviac veľmi atraktívna – s veľkými zľavami, či pomocou zvýhodnených pôžičiek umožňuje zavádzanie svojich systémoch i v krajinách a režimoch, ktoré by si inak podobné systémy vôbec nemohli dovoliť.

Ďalšími významnými hráčmi na trhu dohľadových systémov umelej inteligencie sú Japonsko, USA, Francúzsko, Nemecko a Izrael. **Je znepokojujúce, že demokratické štáty, z ktorých väčšina týchto systémov pochádza, neprijímajú primerané opatrenia na monitorovanie a kontrolu šírenia sofistikovaných technológií spojených s celým radom možných porušení ľudských práv a zneužití autokratickými režimami a diktatúrami.**

Demokratické štáty sú hlavnými používateľmi dohľadu pomocou umelej inteligencie. Z AIGS vyplýva, že 51% vyspelých demokracií využíva systémy dohľadu s umelou inteligenciou. Ide o celý rad technológií dohľadu, od platforiem bezpečného mesta až po kamery na rozpoznávanie tváre. Javí sa, že miera zneužitia je nízka, pretože najdôležitejším faktorom, ktorý určuje, či vlády nasadia túto technológiu na represívne účely, je kvalita a transparentnosť ich vládnutia spojené s regulačnými a legislatívnymi štandardmi.

Vlády v autokratických a semi-autokratických krajinách sú náchylnejšie na zneužívanie dohľadových technológií ako vlády v demokraciách. Využívajú systémy AI na účely masového dohľadu, na posilnenie represie i na dosiahnutie určitých politických cieľov.

Existuje silný vzťah medzi vojenskými výdavkami krajiny a tým, ako vláda využíva

272 FELDMSTEIN, *The Global Expansion of AI Surveillance*. [on-line]. [cit. 28. februára 2022].

Dostupné na internete: <<https://carnegieendowment.org/2019/09/17/global-expansion-of-ai-surveillance-pub-79847>>

dohľadové systémy AI: štyridsať z päťdesiatich krajín s najvyššími vojenskými výdavkami na svete (na základe kumulatívnych vojenských výdavkov) využíva aj technológiu dohľadu pomocou AI.

Tol'ko z citovanej správy a indexu globálneho dohľadu.

Môžeme len doplniť, že kým v niektorých krajinách rastie snaha o transparentné a regulované využívanie systémov AI aj v oblasti dohľadu²⁷³, v iných prináša nasadenie algoritmov umelej inteligencie vyšší stupeň monitorovania a dohľadu nad jednotlivcami i nad celou spoločnosťou. Znepokojivým príkladom tohoto trendu je Čína so svojím systémom sociálneho kreditu²⁷⁴, ktorý pomocou sofistikovaného dohľadového systému, analýzy širokého spektra dát a osobných údajov hodnotí a reguluje činnosť osôb a spoločností v štáte. Tento systém intenzívne využíva technológie umelej inteligencie s cieľom získať efektívnejšie nástroje riadenia správania. Ovocím je bezprecedentný dohľad nad akýmikoľvek aktivitami občanov, ich kategorizácia a obmedzovanie, ak nespĺňajú štátom požadované „parametre“. Ide o nebezpečný precedens, ktorý sa v Číne naďalej rozvíja a – žiaľ – stáva sa inšpiratívnym pre predstaviteľov viacerých krajín (aj demokratického) sveta.²⁷⁵

Pozornému čitateľovi iste neuniklo, že sme – vzhľadom na kontroverznosť a dôsledky sledovania – komunikovali nasadenie algoritmov umelej inteligencie prevažne v rámci dohľadových systémov. No niektoré závery, ktoré sme uviedli zo štúdie *The Global Expansion of AI Surveillance* nám pripomínajú aj ďalšie dva dôležité rozmery spravodajského nasadenia – analytiku a predikcie.

Plošný zber dát, sledovacie a dohľadové systémy by bez analytickej nadstavby boli nefunkčné. Nie je totiž v ľudských silách spracúvať extrémne množstvo spravodajských dát a už vôbec nie sme schopní nad týmito dátami vytvárať pokročilé analýzy a predikcie.

Práve preto bola jednou z prvých vo využití ešte symbolických systémov AI oblasť Data Miningu (dolovanie dát), ktorá – ako sme uviedli v úvode tejto kapitoly – predstavuje analytické metódy a procesy navrhnuté na skúmanie veľkého množstva dát podľa konkrétnych vzorov a väzieb. Výsledkom sú množiny dát, splňujúcich žiadané podmienky,

273 Napríklad aktuálne uvedený regulačný rámec EÚ pre systémy umelej inteligencie.

274 Schematické znázornenie systému sociálneho kreditu je uvedený v prílohe.

275 *Social Credit System*. [on-line]. [cit. 28. februára 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Social_Credit_System>

prípadne predikcia ďalšieho vývoja, správania, atď.

S rozvojom algoritmov hlbokého učenia, fenoménu big data, metadát a kybernetického priestoru data miningové systémy dostali nový dych – sú oveľa sofistikovanejšie, výsledky ponúkajú v reálnom čase a zo vstupných dát vedia vydolovať oveľa viac.

Zároveň – a tu by sme odkázali na postrehy z kapitoly 2.5. ohľadom schopnosti systémov AI vytvárať modely, ktoré sú schopné predpovedať naše konanie – kvalitná analýza je nutnou podmienkou pre modelovanie faktorov, na ktoré sa dohľad a analýza zameriavajú.²⁷⁶

Trio technológií umelej inteligencie – sledovanie a dohľadové systémy spojené s analytickými nástrojmi a schopnosťou vytvárať modely predikujúce ľudské konanie či vývoj situácie – tvorí v súčasnosti skutočne silnú technologickú výbavu (nielen) spravodajských služieb.

V kontexte tejto technologickej výbavy a čínskeho systému sociálneho kreditu vráťme sa ešte k zisteniam štúdie *The Global Expansion of AI Surveillance*. Konkrétne k najviac používaným technológiám AI v dohľadových systémoch, medzi ktoré patria platformy tzv. inteligentných, resp. bezpečných miest, systémy rozpoznávania tváre a inteligentná polícia. Tieto platformy sú súčasťou technológií, známych pod názvami ako **algokracia**, vláda podľa algoritmov, algoritmický právny poriadok, algoritmické vládnutie a pod.

Ide o alternatívnu formu vlády, resp. spoločenského usporiadania, pri ktorom sa na reguláciu, presadzovanie práva a všeobecne na akýkoľvek aspekt každodenného života, ako je doprava alebo registrácia pozemkov, používajú počítačové algoritmy, najmä umelá inteligencia a blockchain.

Algoritmické vládnutie na jednej strane vďaka použitým technológiám môže priniesť nezanedbateľné benefity v správe vecí verejných, skvalitnení a zefektívnení služieb občanom i v riadení štátu, na druhej strane však prináša rôznorodé výzvy a riziká. K výzvam sa určite radí potreba harmonizácie algoritmického vládnutia s princípmi riadenia moderného štátu a e-governmentu i zvládnutie sociologických aspektov a praktických dôsledkov algokracie.

Ako sme videli na záveroch správy nadácie Carnegie Endowment for International Peace,

²⁷⁶ Napríklad modelovanie veľkosti rizika a parametrov teroristického útoku v danej oblasti, vývoj kriminality v konkrétnej časti veľkomesta,...

nasadenie systémov AI smerujúce k algokracii sprevádza viacero problémov v oblasti demokracie a ľudských práv. Nesmieme však zabúdať na skutočnosť, že medzi riziká algoritmického vládnutia patrí aj celý komplex problémov prameniacych z limitov a rizikových faktorov súčasných systémov AI, ktoré sme uvádzali v kapitolách 2.1. až 2.5.

Sumarizujúc túto kapitolu vidíme, že technológie využívajúce algoritmy umelej inteligencie dokážu byť veľmi účinným nástrojom pre spravodajské služby a dohľad s takým presahom do ostatných oblastí verejného života a štátu, ktorý môže znamenať veľký posun v procesoch a spôsobe fungovania spoločnosti. Avšak vzhľadom na politický, ideologický i spoločenský kontext v rôznych častiach sveta a riziká systémov AI treba zabezpečiť, aby prostriedky umelej inteligencie neboli zneužitá alebo sa vymkli kontrole.

Preto je treba v celom spektre nasadenia od spravodajských služieb až po algoritmické riadenie akcentovať veľkú rozvážnosť a vyžadovať striktnú zákonnosť i dôsledný dohľad demokraticky zvolených zástupcov, obmedziť dopady na sociálnu spravodlivosť²⁷⁷ a zabezpečiť dodržiavanie ľudských práv a hodnôt.

2.7. Systémy umelej inteligencie v armáde

Na prelome osemdesiatych a deväťdesiatych rokov minulého storočia sme ešte pred nežnou revolúciou ako študenti technickej kybernetiky žartovali, kedy bude u nás prvý 32-bitový procesor. Odpoveď: s priletom prvej riadenej strely zo západu.²⁷⁸

Do akej miery sa situácia zmenila, vyjadruje skutočnosť, že v súčasnosti už vôbec neriešime, akú procesorovú jednotku moderný zbraňový systém obsahuje, ale hodnotíme, akou technológiou umelej inteligencie disponuje. A nielen „na západe“...

V kapitole 1.8. sme uviedli, že v tzv. zimných časoch umelej inteligencie, t.j. v obdobiach veľkého útlmu na poli AI bol vývoj udržiavaný v zásade len v rámci základného výskumu

²⁷⁷ Digital divide – digitálne rozdelenie ako dôsledok informačnej nerovnosti, ktorá prerastá do digitálnej chudoby a má evidentné dôsledky na život ľudí.

Por. ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi*, [on-line], s. 30-32. [cit. 28. februára 2022].

Dostupné na internete:

<https://peter.santavy.cloud/docs/Analiza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

²⁷⁸ Išlo o riadené strely BGM-109 Tomahawk a MGM-31B Pershing II.

viacerých univerzitných centier a pokroku v príbuzných oblastiach (napr. robotika a kybernetika, výpočtová technika, dátová a počítačová veda, v rámci priemyslu a vojenského vývoja²⁷⁹).

Vzhľadom na predikované možnosti umelej inteligencie už od čias Dartmouthského seminára stál rozvoj vojenskej techniky a armádne nasadenie ako jeden z prvých čakaťel'ov v rade na reálne využitie. A akonáhle ostatná jar AI prerástla do prakticky kontinuálneho rozvoja v tejto oblasti,²⁸⁰ umelá inteligencia sa stala jednou z primárnych technológií nielen moderných zbraňových systémov, ale vojenského využitia v celej svojej šírke a komplexnosti.

Využívanie technológií AI v armáde nie je zviazané len s krajinami, u ktorých by sme to vzhľadom na vedecko-technologický, či vojenský potenciál predpokladali. S dostupnosťou moderných systémov založených na umelej inteligencii a ich komercializáciou sú tieto systémy – ako sme uviedli pri dohľadových systémoch v predchádzajúcej kapitole – lákadlom prakticky pre kohokoľvek.²⁸¹

Využívanie systémov AI vo vojenskej oblasti napreduje predovšetkým v týchto krajinách:

- veľmoci (USA, Čína, Rusko,...), ktoré majú i dostatočný vedecko-technologický potenciál i rozsiahle armádne celky
- vysoko technologicky rozvinuté štáty (Izrael, Japonsko, niektoré štáty EÚ,...), ktorých technologické portfólio takmer prirodzene zahŕňa aj armádne využitie

279 Veľkú roľu v tejto oblasti dlhodobo zohráva DARPA (Defense Advanced Research Projects Agency) – agentúra ministerstva obrany USA zodpovedná za výskum a vývoj nových vojenských technológií. Dostupné na internete: <<https://www.darpa.mil/>>

280 Eventuálny postupný prerod striedania ročných období AI do prakticky kontinuálneho rastu a rozvoja sme diskutovali v kapitole 1.8.

281 Už v roku 2019 vtedajší minister obrany USA Mark Esper predložil informácie o čínskych exportných aktivitách, v rámci ktorých sú na Stredný východ vyvážané drony s inzerovanými schopnosťami smrtiacich autonómnych zbraňových systémov (LAWs – budeme o nich ešte pojednávať).

Napr. vrtuľníkový dron Blowfish A3 od čínskej firmy Ziyan, je vybavený ľahkým guľometom a je schopný operovať na komplikovaných bojových misiách, realizovať prieskumné úlohy a tiež vykonávať vysoko presné likvidačné údery.

TUCKER, P. *SecDef: China Is Exporting Killer Robots to the Mideast*. [on-line]. [cit. 6. decembra 2021]. Dostupné na internete: <<https://www.defenseone.com/technology/2019/11/secdef-china-exporting-killer-robots-mideast/161100/>>

- vysoko militarizované krajiny (India, Turecko,...), ktoré budujú moderné armádne celky a investujú do najmodernejších technológií v tejto oblasti
- problematické štáty (Irán, Severná Kórea,...), ktoré pre svoje ideologické a politické ciele kladú dôraz na rozvoj vojenského potenciálu, no v rámci svojich možností sa zameriavajú na pre nich dostupné a zároveň efektívne technológie, ku ktorým patria aj systémy AI

Uvedené krajiny sa snažia uchopiť využívanie umelej inteligencie vo vojenskej oblasti komplexne. V súčasnosti však nielen ony, ale takmer každá krajina, ktorá sa snaží rozvíjať, resp. budovať moderné armádne zložky, využíva systémy AI aspoň v niektorej z oblastí, medzi ktoré patrí:

- vojenské spravodajstvo
- modelovanie technológií, konfliktov a operácií
- podpora pre velenie
- trenažéry, simulátory a výcvik
- autonómne zbraňové systémy
- skupinové riadenie bojových prostriedkov a autonómnych systémov
- vedenie vojny v kybernetickom priestore

2.7.1. Vojenské spravodajstvo

Využitie systémov umelej inteligencie v oblasti vojenského spravodajstva má na rozdiel od využitia technológií AI v spravodajských službách iných zložiek štátu užší záber činnosti, ktorý mimo autokratických a diktátorských režimov nemá až taký dopad na jednotlivcov a spoločnosť. Na druhej strane však z dôvodu bezpečnosti štátu má oveľa menšie legislatívne obmedzenia (a požíva spoločenský konsenzus) než v civilnom spravodajskom využití. Toto sa ešte stupňuje v režime ohrozenia alebo vojny, keď väčšina regulačných mechanizmov padá, resp. sú nahradené inými.

Čo sme o technológiách AI v spravodajských službách uvádzali v minulej kapitole, môžeme vztiahnuť aj na vojenské systémy. I v oblasti vojenského spravodajstva vzhľadom na limity, riziká a možnosti zneužitia súčasných systémov AI platí požiadavka veľkej opatrnosti a striktnej zákonnosti i potreba dôsledného dohľadu na to určených armádnych štruktúr a mechanizmov, aby prostriedky umelej inteligencie neboli zneužitú, resp. sa

vymkli kontrole. Vzhľadom na obmedzenejšie možnosti demokratického dohľadu a kontrolných mechanizmov – osobitne v čase ohrozenia alebo vojny – to nemusí byť jednoduché zabezpečiť. Hrozbou je i využitie v autokratických režimoch a diktatúrach s prienikom armádnej a občianskej roviny spoločnosti.

2.7.2. Modelovanie technológií, konfliktov a operácií

Modelovanie technológií, konfliktov a operácií je zaujímavou oblasťou využitia systémov AI, ktorá – ak sa bavíme o technológiách – osobitne vo vede slávi nebývalé úspechy. Modelovanie zložitých dejov na kvantovej úrovni, štiepných a fúzných reakcií, makrokozmičských astrofyzikálnych dejov,²⁸² látok v rámci materiálovej fyziky a chémie, biologických a biochemických reakcií – to všetko je nielen príkladom pre využitie vo vojenskom sektore, ale v mnohých oblastiach i priamo súčasťou vojenského výskumu.

Na nasadenie systémov AI v rámci modelovania vojenských technológií môžeme mať dva pohľady. V tom prvom ide o pozitívnu vec, keď sa prostredníctvom modelovania znižuje nebezpečenstvo prameniace z iných foriem výskumu. Ide napríklad o vývoj nových jadrových zbraní, ktoré nevybuchujú na jadrových polygónoch, ale len vo vizualizáciách superpočítačov²⁸³, modelovanie nových prvkov riadených striel, stíhacích lietadiel piatej a šiestej generácie²⁸⁴, odolných materiálov a pod.

Druhý pohľad je skôr varovný. Schopnosť realisticky modelovať a simulovať môže byť lákadlom pre vývoj zakázaných vojenských prostriedkov, napríklad biologických

282 Ako sme uvádzali v poznámkovom aparáte k úvodnej kapitole – AI prvý raz simulovala vesmír: rýchlo, presne a nikto nevie ako.

SUMMER, *The first AI universe sim is fast and accurate—and its creators don't know how it works*. [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://phys.org/pdf480780725.pdf>>

Odvtedy vzniklo a stále vzniká viacero ďalších simulácií vesmíru, resp. jeho častí, pri ktorých sa algoritmy AI neustále zlepšujú a simulácie sú stále presnejšie a detailnejšie.

283 Ako príklad môžeme uviesť národné laboratóriá USA pre výskum a vývoj jadrových zbraní v Los Alamos. Aspoň jeden z ich superpočítačov sa už tradične umiestňuje v prvej desiatke najvýkonnejších superpočítačov sveta.

DOE/NNSA/LANL/SNL [on-line]. [cit. 3. marca 2022].

Dostupné na internete: <<https://www.top500.org/site/50334/>>

284 Tieto stroje sú samy o sebe AI enabled, t.j. intenzívne využívajú svoje vlastné lokálne systémy AI a participujú aj na externých systémoch AI, napr. v rámci skupinového riadenia bojovej operácie.

a chemických zbraní...²⁸⁵

Nech sa však na problematiku modelovania a simulovania vojenských technológií pozeráme akokoľvek, vždy treba mať na pamäti limity a riziká uvedené v kapitolách 2.1. až 2.4. Ak by sme ich nebrali do úvahy, akýkoľvek presun technológie alebo zbraňového systému z virtuálnej simulácie do reálneho sveta môže byť ak nie tragédiou, tak minimálne trpkým sklamaním.

Zaujímavým využitím algoritmov umelej inteligencie je modelovanie takticko-operačných stratégií a postupov. Ak je pre tento účel využitý kvalitne natrénovaný systém AI a je k dispozícii dostatok taktických a operačných dát, môže ísť o nezanedbateľný vklad k optimalizácii armádneho nasadenia. Ovocím môže byť nielen zníženie strát na vojakoch a technike v rámci správnej nasimulovanej operácie, ale – ak je pre to záujem a vôľa – i zníženie kolaterálnych strát, t.j. obetí z radov civilného obyvateľstva, životne dôležitej infraštruktúry a pod. Toto ovocie sa však môže stať trpkým, ak by obrovská technologická prevaha vo využití systémov AI bola zneužitá na napadnutie iných štátov alebo upevnenie diktatúry.

Prečo nepristúpiť od modelovania operácií k modelovaniu celých konfliktov? Myslíme si, že v rámci súčasných možností systémov umelej inteligencie a od toho sa odvíjajúcich problémov, ktoré sme diskutovali v závere kapitoly 2.1.2, reálne modelovanie celých konfliktov nie je na programe dňa. Možno parciálne modelovanie jednotlivých krokov prebiehajúceho konfliktu môže byť veleniu vojenských operácií nápomocné, no viac skutočne reálnych a spoľahlivých výsledkov sa za súčasného stavu rozvoja systémov AI asi dosiahnuť nedá. Rizikom by sa teda mohlo stať, že by armádne vedenie, resp. vládni činitelia uverili predikciám, ktoré by potencionálne modelovanie mohlo prinášať. I keď túto dôveru môžeme považovať za nereálnu až iracionálnu, niektoré varovania uvedené v kapitole 2.5. a technologické mýty²⁸⁶, ktoré informačnú spoločnosť sprevádzajú, by mohli

285 Na základe jednoduchých zmien v nastaveniach systém AI MegaSyn navrhol miesto nových liekov niekoľko desiatok vražedných chemických bojových látok.

URBINA, F., LENTZOS, F., INVERNIZZI, C. et al. *Dual use of artificial-intelligence-powered drug discovery*. In: *Nat Mach Intell.* [on-line]. 2022, 4, s. 189–191. [cit. 22. marca 2022].

Dostupné na internete: <<https://doi.org/10.1038/s42256-022-00465-9>>

286 Ide o sociologický fenomén prehnanej dôvery vo funkčnosť, spoľahlivosť a presnosť systémov IKT a AI, ktorá sa radí medzi tzv. mýty informačnej spoločnosti a prejavuje sa ako rizikový faktor napríklad v oblasti kybernetickej bezpečnosti.

k tejto klamlivej dôvere viesť. A bola by to veľká tragédia, ak by ovocím takejto dôvery bolo rozpútanie vojnového konfliktu...

2.7.3. Podpora pre velenie

Rozvíjaním našej rozpravy o modelovaní takticko-operačných stratégií a postupov i simulovaní konfliktov sme priamo prešli do oblasti podpory pre velenie, pričom v niektorých aspektoch sa delenie na modelovanie a podporu prakticky stiera a navyše sa kombinuje s výsledkami vojenského spravodajstva z oblasti analýzy a predikcie.

Podpora pre velenie však prináša aj ďalšie možnosti využitia systémov umelej inteligencie. Ide napríklad o modelovanie logistických potrieb a trás zásobovania, manažment telekomunikačného zabezpečenia velenia, presné meteorologické predpovede a mnoho iných služieb podpory. Niektoré z nich môžeme považovať za natoľko zabehnuté a overené, že s vážnejšími zlyhaniami ani nerátame, resp. zlyhania sú zakomponované do riadenia rizík.²⁸⁷

Iné však v sebe môžu skrývať riziká špecifické pre algoritmy AI. Napríklad zlyhanie zásobovania počas prebiehajúceho konfliktu alebo humanitárnej operácie, ktorého dôvodom je nesprávne natrénovaný systém AI riadiaci logistické procesy. Ak totiž máme systém trénovaný na dátach možno aj veľmi úspešných vojnových operácií, avšak riadených iným veliteľským štýlom, než aký používa aktuálne velenie (iné vedenie operácií môže znamenať inú [s]potrebu materiálu, pohonných hmôt, podporného personálu, ženijnej podpory a pod.), generované požiadavky na logistiku môžu byť nesprávne.

2.7.4. Trenažéry, simulátory a výcvik

Tréningové, operačné i taktické dáta (aj) zo sledovania, analytika i predikcia, modelovanie a simulácie – to všetko sa spája v moderných trenažéroch, simulátoroch a výcvikových zariadeniach využívajúcich systémy AI.

Ide o technológie, ktoré vďaka algoritmom umelej inteligencie dávno presahujú schopnosti bežných tréningových nástrojov a pomôcok, dokážuc hodnoverne simulovať reálne bojové

Konkrétnejšie sa problematike dôvery voči systémom AI venujeme o niekoľko kapitol ďalej.

287 Bez riadneho manažmentu rizík nie je možné hovoriť o spoľahlivom a bezpečnom nasadení prostriedkov IKT a AI v životne dôležitých oblastiach.

Risk management [on-line]. [cit. 3. marca 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Risk_management>

nasadenie a v rámci výcviku aktívne vystupovať ako skutočne schopní protivníci najrozličnejšieho druhu.

Výhodou nasadenia simulátorov vybavených technológiami AI je zabezpečenie kvalitnej odbornej i taktickej prípravy vojenského personálu a možnosť dlhodobo ho udržať v dobrej kondícii. Navyše bez zbytočných rizík a extrémnych nákladov pre zabezpečenie tak kvalitnej prípravy.

Rizikovým faktorom je psychologický aspekt vytvorenia si podvedomého chápania boja ako virtuálnej reality. **Dôsledkom môže byť strata citlivosti pre ohrozenie ľudských životov a reálnych hodnôt pri nasadení v skutočnom boji.** Tento rizikový faktor je o to záľadnejší, že sa „pod kožu“ dostáva postupne a podvedome, mnohokrát bez varovnej reakcie svedomia.

Ako sme spomenuli, trenažéry a simulátory bojových operácií na báze umelej inteligencie vedia vystupovať ako zdatní súper i vo virtuálnych bojových situáciách. Logickým dôsledkom je tak ich možnosť nasadenia v reálnych zbraňových systémoch...

2.7.5. Autonómne zbraňové systémy

Jednou z najdiskutovanejších oblastí využitia systémov umelej inteligencie v armádnej sfére sú autonómne zbraňové systémy. Ide o armádne vybavenie, u ktorého sa snúbi pokročilé technologické a zbraňové vybavenie s adaptabilitou a autonómiou systémov umelej inteligencie.

Autonómne zbraňové systémy by sme v zásade mohli rozdeliť na smrtiace a ostatné, zahŕňajúce široké portfólio bojovej techniky od prieskumných operácií, cez komunikačnú a logistickú podporu až napríklad po ženijné nasadenie.²⁸⁸

Z pohľadu nasadenia systémov AI vo vojenskej oblasti však zásadnú rolu hrajú smrtiace autonómne zbraňové systémy²⁸⁹, schopné na základe činnosti algoritmov umelej inteligencie samostatne vyhľadávať a zasahovať ciele. Ide o systémy pokrývajúce najrozličnejšie zbraňové technológie a schopné operovať vo vzduchu, na zemi, vo vode, pod vodou alebo vo vesmíre.²⁹⁰

288 Bojová technika z portfólia ostatných autonómnych zbraňových systémov sa vo viacerých oblastiach priamo či nepriamo prekrýva s inými oblasťami systémov AI využívaných v armáde.

289 Lethal autonomous weapon – používa sa skratka LAWS.

290 *Lethal autonomous weapon* [on-line]. [cit. 5. marca 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Lethal_autonomous_weapon>

Náš záujem v kontexte etických výziev sa prirodzene zameriava práve na tieto smrtiace autonómne zbraňové systémy...

Podľa niektorých vojenských expertov autonómne zbraňové systémy nielenže prinášajú významné strategické a taktické výhody na bojisku, ale sú aj z morálnych dôvodov vhodnejšie v porovnaní s nasadením ľudského personálu. Kritici naopak zastávajú názor, že tieto zbrane by sa mali obmedziť, ak nie úplne zakázať, a to z rôznych morálnych a právnych dôvodov.²⁹¹

Protagonisti LAWs teda obhajujú ich využitie v dvoch rovinách – z dôvodu vojenských výhod a z dôvodu morálnosti ich nasadenia.

Medzi vojenské výhody radia ponajprv ich multiplikačný efekt, čo znamená, že na vykonanie bojovej misie je potrebný menší počet vojakov, pričom účinnosť každého z nich je vyššia. Ďalším benefitom je rozšírenie bojiska, keďže LAWs sú schopné preniknúť do oblastí, ktoré by boli bez ich nasadenia inak nedostupné. A v neposlednom rade je dôležitou výhodou i nahradenie ľudského faktora v nebezpečných misiách a tým aj minimalizovanie strát na životoch, čo má ďalekosiahle dôsledky – od psychologických v rámci armády, cez spoločenské (verejná mienka a podpora je dôležitým faktorom pre bojové operácie) až po materiálne a odborné.²⁹²

Ďalšie zaujímavé výhody uvádza aj cestovná mapa nasadenia bezpilotných systémov pre roky 2007-2032: robotické systémy sú vhodnejšie ako ľudia pre nasadenie na „nudné, špinavé alebo nebezpečné“ misie. Medzi nudné misie patria dlho-trvajúce (monotónne) operácie. Ako príklad špinavej misie sa uvádza vystavenie personálu potenciálne škodlivému rádiologickému materiálu. A príkladom nebezpečnej misie môže byť

291 ETZIONI, A., ETZIONI, O. *Pros and Cons of Autonomous Weapons Systems* [on-line]. [cit. 5. marca 2022].

Dostupné na internete: <<https://www.armyupress.army.mil/Journals/Military-Review/English-Edition-Archives/May-June-2017/Pros-and-Cons-of-Autonomous-Weapons-Systems/>>

292 Por. MARCHANT, G. E. et al. *International Governance of Autonomous Military Robots*. In: *Columbia Science and Technology Law Review*. [on-line]. 2011, 12, s. 272–76. [cit. 27. marca 2017].

Dostupné na internete: <<http://stlr.org/download/volumes/volume12/marchant.pdf>>

Ohľadom materiálnych a odborných dôsledkov minimalizovania strát na životoch je potrebné si uvedomiť, že odborný výcvik, vzdelanie a tréning pre moderné sofistikované zbraňové systémy je časovo i intelektuálne náročný a drahý. Napr. strata bojového lietadla v hodnote desiatok miliónov dolárov bolí, no oveľa viac velenie mrzí, ak pri tom príde aj o schopného pilota, ktorého nenahradia tak ľahko ako havarovaný, či zostrelený stroj.

zneškodňovanie výbušnín, či boj pod intenzívnou paľbou protivníka.²⁹³

Medzi ďalšie výhody patrí potenciál operovať v rýchlejšom tempe, ako je možné dosiahnuť u ľudí a schopnosť smrteľne zasiahnuť aj vtedy, keď je prerušené komunikačné spojenie.²⁹⁴ V dlhodobom horizonte ide aj o nemalé úspory na personále – jeho vzdelaní, tréningu a službe.²⁹⁵

I keď sme s uvedením všetkých výhod nasadenia LAWs v armáde ešte neskončili, prichádzame k závažnému zisteniu, pre ktoré je a bude obmedzovanie nasadenia smrtiacich systémov veľmi ťažké – **vo viacerých scenároch nasadenia sú zbraňové systémy vybavené technológiami umelej inteligencie lepšie ako ľudia.**

Z pohľadu LAWs to nie je nič nové, ak do kategórie smrtiacich automatických zbraňových systémov radíme nielen systémy AI, ale aj jednoduchšie automatizované zbrane, napr. radarom navádzané systémy CIWS používané na obranu lodí, ktoré sa používajú od 70. rokov 20. storočia. Takéto systémy dokážu autonómne identifikovať a zaútočiť na prichádzajúce rakety, delostreleckú paľbu, lietadlá a hladinové plavidlá podľa kritérií stanovených ľudskou obsluhou.²⁹⁶

Moderné technológie umelej inteligencie však schopnosti LAWs povzniesli na úplne novú úroveň, pri ktorej sa pokročilá automatizácia mení na skutočnú adaptabilitu

293 CLAPPER, J. R. et al. *Unmanned Systems Roadmap: 2007-2032*. [on-line]. Washington, DC:

Department of Defense [DOD], 2007, s. 19. [cit. 5. marca 2022].

Dostupné na internete: <http://www.globalsecurity.org/intell/library/reports/2007/dod-unmanned-systems-roadmap_2007-2032.pdf>

294 THURNHER, J. S. *Legal Implications of Fully Autonomous Targeting*. In: *Joint Force Quarterly*. [on-line]. 2012, 67, 4, s. 83. [cit. 7. marca 2022].

Dostupné na internete: <http://ndupress.ndu.edu/Portals/68/Documents/jfq/jfq-67/JFQ-67_77-84_Thurnher.pdf>

295 David Francis v *The Fiscal Times* v roku 2013 cituje údaje ministerstva obrany, podľa ktorých „každý vojak v Afganistane stojí Pentagon približne 850 000 dolárov ročne“.

FRANCIS, D. *How a New Army of Robots Can Cut the Defense Budget*. In: *Fiscal Times*. [on-line]. 2013, 2. 4. [cit. 5. marca 2022].

Dostupné na internete: <<http://www.thefiscaltimes.com/Articles/2013/04/02/How-a-New-Army-of-Robots-Can-Cut-the-Defense-Budget>>

296 V lodných systémoch ide napr. o americké systémy CIWS Phalanx, pri tankových napr. ruský systém Arena, izraelský Trophy a nemecký AMAP-ADS.

a autonómnosť.²⁹⁷

Ako príklad môžeme uviesť moderné letecké bojové simulátory, pri ktorých – nadviažuc na pasáž o využití systémov AI pri trenažéroch a výcviku – je možné porovnávať schopnosti ľudskej obsluhy sofistikovaných zbraňových systémov a ich umelých náprotivkov. Systém umelej inteligencie ALPHA²⁹⁸, ktorý pre tento účel vyvinuli na Univerzite v Cincinnati, zdrvujúco porazil v simulovaných bojových operáciách jedného z najlepších inštruktorov stíhacích pilotov v USA, plukovníka Geneho Leea.²⁹⁹ ALPHA ho dokonale zaskočila – excelentne sa chránila a dokázala okamžite správne reagovať. Výborne predvídala zámery protivníka a pohotovo odpovedala na akékoľvek zmeny letu a odpaľovanie rakiet. Nielenže sa vynikajúco vyhýbala odpáleným raketám, ale vedela podľa potreby aj plynule prechádzať medzi defenzívnymi a ofenzívnymi fázami boja. ALPHA proti plukovníkovi bojovala v simulácii reálnych, niekoľko hodín trvajúcich misií, po ktorých bol Lee totálne vyčerpaný – fyzicky i mentálne na dne. Asi netreba dodávať, že ALPHA bola neustále na koni a nemala problém bojovať ďalej...³⁰⁰

Zámerne sme použili príklad simulátora z roku 2016, aby sme si v kontexte neustále pripomínaného veľkého pokroku na poli umelej inteligencie uskutočneného v ostatných rokoch uvedomili, o koľko vyspelejšia je asi súčasná technika LAWs.

Zbraňové systémy vybavené technológiami AI sa vo veľkom presúvajú z oblasti simulátorov do reálnych cvičných a bojových strojov.

Ponajprv išlo a ide o cvičné stroje, ktoré umožňujú využívať výhody umelej inteligencie známe z virtuálneho prostredia v reálnom výcviku. Len v USA ide o produkciu viacerých

297 Ide o základné vlastnosti každého systému umelej inteligencie, ktoré sme uviedli v kapitole 1.3.

298 Ide o fuzzy orientovaný systém umelej inteligencie s genetickými algoritmami (genetic-fuzzy systems).

299 Inštruktor, ktorý vyškolil tisíce pilotov US Air Force, to komentoval slovami: skúsený bojový pilot dokáže bez problémov poraziť väčšinu súčasných bojových umelých inteligencií ako lovnú zver... Pri systéme ALPHA sme lovnou zverou my... Ešte nikdy som sa nestretol s tak agresívnou, vnímavou, dynamickou a vierohodne bojujúcou umelou inteligenciou, ako je ALPHA.

300 MIHULKA, S. *Letecká bojová umělá inteligence si natřela na chleba taktické experty*. [on-line]. [cit. 5. marca 2022].

Dostupné na internete: <<https://www.osel.cz/8903-letecka-bojova-umela-inteligence-si-natrelo-na-chleba-takticke-experty.html>>

firiem³⁰¹, čomu sa snažia úspešne sekundovať výrobcovia z iných krajín.³⁰²

Akonáhle sa využitie cvičných systémov osvedčilo, prichádzalo na nasadenie v prieskumných a bojových misiách, pričom sa rozdiely medzi prieskumným a bojovým využitím stierajú, takže moderné LAWS vedia byť veľmi univerzálne pri reálnom nasadení.³⁰³

Na vývoji LAWS v oblasti vojenského letectva je asi najmarkantnejšie vidieť, ako sa moderné automatizované zbraňové systémy približujú špičkovej leteckej technike. Moderné drony sú vybavené podobnými technológiami ako bojové lietadlá generácií 4+ a 5 a sú schopné s nimi aj v bojových operáciách priamo spolupracovať.³⁰⁴

Zaujímavým prvkom je i kombinácia rôznych technológií v rámci LAWS. Napríklad vývoj moderného lacného³⁰⁵ bojového dronu Skyborg v máji 2021 ešte nebol hotový, no jeho riadiaci mozog (Autonomy Core System, ACS) poháňaný umelou inteligenciou už bol zrelý na testovanie v reálnej prevádzke. ACS bol preto nainštalovaný na taktický dron Kratos UTAP-22 a poslaný do vzduchu.³⁰⁶ Zámerne uvádzame aj túto – pre armádny „mainstream“ – nepodstatnú udalosť, keďže ukazuje, do akej miery moderné zbraňové systémy v posledných rokoch pokročili od špeciálne „vypíplaných“ a na mieru vytvorených celkov k univerzálnym riešeniam.

Pokiaľ teda LAWS využívajú algoritmy umelej inteligencie, dokážu byť neuveriteľne technologicky adaptabilné, relatívne lacné a minimálne rovnako nebezpečné ako najmodernejšie zbraňové systémy riadené ľudskou obsluhou.

Čitateľ, pamätajúc na riziká a limity, ktoré sme v kapitolách 2.1. až 2.4. uvádzali, však okamžite spozornie – **ako dokážeme zabezpečiť, že smrtiace autonómne systémy,**

301 Len z najznámejších ide o General Atomics, Boeing alebo Kratos, z tých najnovších so zaujímavými cieľmi vývoja napr. Exosonic alebo Dynetics.

302 V zásade ide o krajiny patriace do štyroch kategórií štátov najviac napredujúcich vo využívaní AI vo vojenskej oblasti, ktoré sme spomínali v úvode tejto kapitoly.

303 Kto by nepoznal klasický dron MQ-9 Reaper od General Atomics – ostrieľaného kežáka, pardon kovboja, ktorý už bol nasadený v nespočetných prieskumných a bojových misiách.

304 Napr. prúdový bojový dron Gambit od General Atomics, Loyal Wingman od Boeingu a pod.

305 Ide o kategóriu tzv. „attritable drones“ – ktoré sú celkom lacné, takže môžu byť nasadené vo vysoko rizikových konfliktoch bez toho, že by vyššie bojové straty ohrozili rozpočet.

306 SZONDY, D. *Skyborg combat drone's "brain" flies for the first time*. [on-line]. [cit. 7. marca 2022]. Dostupné na internete: <<https://newatlas.com/military/skyborg-combat-drones-brain-first-time/>>

ktoré sú zámerne dizajnované tak, aby dokázali efektívne pracovať aj pri strate spojenia, budú vedieť správne použiť svoju smrtiacu silu a rešpektovať požadované mantinely, ak je nám známe, že systémy AI robia mnohokrát iné chyby ako ľudia³⁰⁷ a nevieme ich vytrénovať tak, aby vedeli adekvátne odpovedať na všetky reálne situácie?

Rozoberajúc autonómne zbraňové systémy sme sa viackrát dotkli ich nasadenia v kooperácii s inými automatizovanými alebo ľudskou obsluhou riadenými systémami.

2.7.6. Skupinové riadenie bojových prostriedkov a autonómnych systémov

Pri skupinovom riadení bojových prostriedkov a autonómnych systémov by sme mohli hovoriť o LAWs na stereoidoch, keďže rôzne výhody uvedené pri ich individuálnom nasadení sa znásobujú. Podstatou skupinového riadenia je spojenie schopností a činnosti autonómnych zbraňových systémov do jedného harmonického celku, ktorého časti síce plnia svoje vlastné samostatné úlohy, no spoločne sledujú jeden bojový cieľ.³⁰⁸

Rozlišujeme dva spôsoby skupinového riadenia a nasadenia autonómnych zbraňových systémov: spoločné nasadenie s bojovými prostriedkami ovládanými človekom a koordinované nasadenie LAWs pod spoločným riadením centrálnym systémom umelej inteligencie, prípadne koordinovaným riadením na základe interakcií jednotlivých autonómnych systémov.

Zbraňové systémy AI spolupracujúce s prostriedkami ovládanými človekom prevažne pracujú v submisívnom režime, pri ktorom bojovú misiu riadi človek. LAWs v tomto režime zvyčajne operujú ako sprevádzajúce a spolupracujúce zbraňové systémy pod velením človeka. Vzhľadom na svoju výbavu systémami AI však dokážu vykonávať parciálne podúlohy bojovej misie nielen autonómne, ale prípadne sú schopné aj prevziať iniciatívu a operovať úplne samostatne.³⁰⁹

307 Omyly a chyby systémov AI sú iné, než tie ľudské – viď kapitola 2.1. alebo diskusia v kapitole 1.1. o programe Eugene, ktorý ako prvý zvládol Turingov test.

308 Americký Defense Science Board používa termín *multiagent coordination* pre okolnosti, v ktorých je úloha rozdelená medzi „viacero robotov, softvérových agentov alebo ľudí“.
Defense Science Board. *Task Force Report: The Role of Autonomy in DoD Systems*. Washington, DC: Office of the Under Secretary of Defense for Acquisition, Technology and Logistics, 2012.

309 Napríklad v dňoch písania tejto kapitoly predstavený plnohodnotný bojový prúdový dron Gambit od General Atomics, ktorý by mal operovať ako sprevádzajúci bojový stroj pilotovaného lietadla a po jeho boku a pod jeho vedením sa zúčastňovať vzdušných bojov. Vzhľadom na svoju výbavu systémami AI

Mohlo by sa zdať, že nasadenie LAWs je v tomto režime pod kontrolou, no v prípade umelou inteligenciou riadených zbraní to tak vôbec nemusí byť. Súčasťou autonómneho vykonávania parciálnych úloh tiež môže byť použitie smrtiacej sily na základe vyhodnotenia, ktoré vykoná algoritmus AI. Prevzatie iniciatívy sprevádza v zásade to isté riziko, len môže byť umocnené nielen úplne autonómnym rozhodovaním systému AI, ale i dôsledkami aktuálneho vývoja bojovej situácie.³¹⁰

Koordinované nasadenie LAWs pod vedením centrálného systému AI, resp. spoločného riadenia zapojených systémov AI samozrejme obnáša rovnaké kategórie rizík ako samostatné autonómne zbraňové systémy, len ich dôsledky sú v prípade zlyhania nebezpečne násobené.

Naviac sa vynára viacero nových problémov týkajúcich sa činnosti týchto zbraňových systémov v prípade výpadku spojenia, resp. kolapsu centrálného systému riadenia. Sú či budú parciálne systémy AI (napríklad u miniatúrnych alebo skôr spomínaných „attributable drones“) tak riešené, aby vedeli bezpečne zvládať aj tieto situácie?

Je koordinované riadenie na základe interakcie jednotlivých systémov AI natoľko robustné, aby prepojené LAWs dokázali správne pracovať aj pri výpadku kritického množstva zapojených automatických systémov?

Vieme dostatočne odhadnúť správanie centrálne riadeného (a zároveň parciálne autonómneho) alebo koordinovane riadeného skupinového automatického zbraňového komplexu vo vyhranených situáciách?

2.7.7. Vedenie vojny v kybernetickom priestore

Last but not least, treba spomenúť využitie systémov umelej inteligencie v rámci vedenia vojny v kybernetickom priestore.³¹¹ Kybernetický priestor v moderných doktrínach

však dokáže vykonávať podúlohy bojovej misie autonómne a prípadne aj prevziať iniciatívu a operovať úplne samostatne.

SZONDY, D. *General Atomics' Gambit autonomous combat drone takes the initiative*. [on-line]. [cit. 7. marca 2022].

Dostupné na internete: <<https://newatlas.com/military/general-atomics-gambit-combat-drone-air-dominance/>>

310 Napríklad zostrelenie veliaceho pilotovaného bojového lietadla v tak špecifickej konkrétnej situácii v rámci mestskej aglomerácie, že následné správne rozhodnutie o použití smrtiacej sily autonómneho dronu v zastavanej oblasti je nad natrénovaný rámec jeho systému AI.

311 Kybernetické zbrane – cyberwarfare.

predstavuje ďalší z rozmerov vedenia vojny, ktorý sa v poslednom desaťročí stáva tak významným, že si vo viacerých armádach vydobyl svoje pevné postavenie v rámci samostatnej armádnej zložky. Mnohí experti ho nazývajú bojovým priestorom 21. storočia.³¹² Kybernetický priestor zároveň predstavuje novú operačnú doménu, ktorej charakter a prepojenosť s ostatnými operačnými doménami – zem, voda, vzduch a vesmír – vytvára priestor pre synergiu a efektívne použitie vojenskej sily, ekonomizáciu vynakladaných prostriedkov a dosahovanie zámerov velenia v reálnom čase.³¹³

Fenomén kybernetického vedenia vojny stavia na dvoch pilieroch – využití informačných a komunikačných technológií v spravodajstve a pokroku v oblasti kybernetickej bezpečnosti, pričom treba zdôrazniť, že pomyselné hranice medzi kybernetickým vedením vojny a spravodajstvom, resp. kybernetickou bezpečnosťou sú veľmi vágne a mnohé kybernetické operácie je ťažké presne zaradiť.

Spravodajský pilier sme načrtli v predchádzajúcej kapitole. Pilier kybernetickej bezpečnosti³¹⁴ je vhodné trochu konkretizovať, keďže pre účely vojenského využitia kybernetická bezpečnosť ide ruka v ruke s využitím metód a nástrojov kybernetickej kriminality.

Bezpečnosť v oblasti IKT prešla v posledných desaťročiach búrlivým vývojom – od praktickej neexistencie, či úplného marginalizovania, cez riešenie problémov, ktoré vznikali pri reálnej prevádzke a potrebe chrániť systémy pred dôsledkami rozvíjajúcej sa počítačovej kriminality, až po dnešný stav, v ktorom ide o dosť rozsiahly vedecký obor, jeden z podstatných a nutných rozmerov dizajnu moderných informačných technológií a vytváranie jednotiek kybernetického boja u armád nielen najsilnejších štátov sveta. Zároveň z pohľadu informačnej kriminality ide v súčasnosti už o najvýnosnejší spôsob

312 Ak v moderných dejinách vojny bolo 19. storočie arénou pozemných jednotiek a námorníctva a 20. storočie ovládol strategický význam letectva, virtuálny priestor vytvorený celosvetovou digitálnou interkonektivitou má byť priestorom kybernetického boja.

GIBNEY, A. *Zero Days*. [filmový dokument]. [cit. 10. marca 2022].

Dostupné na internete: <<http://www.zerodayfilm.com/>>

313 Z predkladaného Návrhu Stratégie kybernetickej obrany Slovenskej republiky.

Stratégia kybernetickej obrany Slovenskej republiky. [on-line]. Ministerstvo obrany SR, Vojenské spravodajstvo, s. 3. [cit. 17. marca 2022].

Dostupné na internete: <<https://www.slov-lex.sk/legislativne-procesy/-/SK/dokumenty/LP-2022-128>>

314 Používame termín kybernetická a nie informačná bezpečnosť, keďže kybernetická bezpečnosť je podmnožinou informačnej bezpečnosti zahŕňajúcej aj oblasti, ktoré nie sú predmetom našej diskusie.

kriminality, predstihujúc tak drogovú kriminalitu a obchodovanie s ľuďmi.³¹⁵

Kybernetická kriminalita stavia na realite chýb v prostriedkoch IKT³¹⁶, ktoré umožňujú vniknutie do systémov, únik informácií a dokonca prevzatie vlády nad systémom.

S nástupom masívneho využívania prostriedkov IKT v mnohých oblastiach života a fungovania spoločnosti prichádzalo k postupnej kriminalizácii zneužívania chýb.

V rámci prebiehajúceho procesu prechodu k informačnej a znalostnej spoločnosti, ktorá je bytostne spätá s fungujúcimi prostriedkami IKT, kybernetická bezpečnosť, resp. jej temná kriminálna stránka nadobúdajú strategickú dôležitosť a stávajú sa predmetom útočných i obranných aktivít v čase vojnových konfliktov.

K charakteristickým črtám prvkov kybernetického vedenia vojny tiež patrí ich zaradenie k tzv. prostriedkom vojen štvrtej generácie³¹⁷, ktoré sú popisované stieraním hraníc medzi vojnou a politikou³¹⁸, medzi armádou a civilným obyvateľstvom, decentralizovaným vedením vojny, guerilovou taktikou a prvkami terorizmu, dezinformačným pôsobením a propagandou, útokom na kultúru a psychologickými metódami na oslabenie protivníka.³¹⁹

315 Kybernetická kriminalita sa v rokoch 2015-2016 stala najvýnosnejšou oblasťou kriminality.

Cybercrime Is Now More Profitable Than The Drug Trade [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <<https://www.tripwire.com/state-of-security/regulatory-compliance/pci/cybercrime-is-now-more-profitable-than-the-drug-trade/>>

Pokročilý ransomware je v súčasnosti najvýnosnejším obchodným modelom kybernetickej kriminality.

Why Advanced Ransomware Is Cybercrime's Most Profitable Business Model [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <<https://blog.knowbe4.com/why-advanced-ransomware-is-cybercrimes-most-profitable-business-model>>

ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi*, [on-line], s. 129. [cit. 10. marca 2022].

Dostupné na internete:

<https://peter.santavy.cloud/docs/Analýza_virtualneho_sвета_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

316 Ide o chyby v hardvéri, základnom softvérovom vybavení (firmware, operačné systémy), aplikačnom softvéri, komunikačných technológiách (chyby šifrovania a pod.).

317 *Fourth-generation warfare*. [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Fourth-generation_warfare>

318 Tiež treba podotknúť, že aktérmi nemusia byť len štáty, resp. štátne zoskupenia. Môže ísť nielen o zástupné organizácie niektorých štátov, ale aj o akékoľvek iné mimovládne činitele.

319 Zaujímavým konkrétnym príkladom komplexnosti štvrtej generácie je štúdia prestížneho amerického

S rastúcou účinnosťou kybernetických prostriedkov, potencionálnym efektívnym dopadom na nepriateľa a výhodami v rámci vedenia vojenských operácií sa vo viacerých armádach kybernetické jednotky postupne etablujú ako samostatná armádna zložka.

V oblasti kybernetickej bezpečnosti, v ktorej sa autor profesionálne pohybuje, osobitne v posledných desaťročiach badať nielen uskutočnenie prvých úspešných a sofistikovaných útokov vedených priamo a len prostriedkami kybernetického vedenia vojny, ale aj neustále ataky, ktoré do tejto oblasti patria. Ide napríklad o testovanie štruktúr IKT niektorých štátov (v posledných čase realizované napr. v špecializovaných DDoS útokoch na štátnu infraštruktúru Estónska a Česka), útoky na komplexné siete internetu vecí, ktoré riadia energetické sústavy, tepelné rozvody, diaľkové distribučné siete fosílnych palív³²⁰, a pod.

Priam exemplárny vzhlľad do problematiky kybernetického vedenia vojny ponúka Stuxnet³²¹, ktorý bol v roku 2007³²² nasadený voči jadrovým zariadeniam v Iráne a ktorý predstavoval v zásade prvý, zároveň najväčší a najsofistikovanejší prípad štátom, resp. viacerými štátmi vykonaného kybernetického útoku.

Stuxnet bol po prekročení svojich operačných hraníc a rozšírení do viacerých častí sveta objavený bieloruskou bezpečnostnou službou VirusBlokAda v roku 2010 a následne analyzovaný špecialistami firiem Kaspersky a Symantec, ktoré sa zameriavajú na kybernetickú bezpečnosť.

Stuxnet bol vyvinutý pre útoky na systémy SCADA, prostredníctvom ktorých boli ovládané

think-tanku RAND Corporation o možnostiach destabilizácie a ekonomického vyčerpania Ruska.

DOBBINS, J., COHEN, R. S., CHANDLER, N. et al. *Overextending and Unbalancing Russia: Assessing the Impact of Cost-Imposing Options*. [on-line]. Santa Monica, CA: RAND Corporation, 2019. [cit. 18. marca 2022].

Dostupné na internete: <https://www.rand.org/pubs/research_briefs/RB10014.html>

320 Tvrdý útok na Colonial Pipeline, najväčšiu distribučnú sieť fosílnych palív v USA, za ktorým stála hackerská skupina DarkSide a ktorý prebehol začiatkom mája 2021.

Colonial Pipeline offline due to ransomware attack. [on-line]. [cit. 13. marca 2022].

Dostupné na internete: <<https://www.fortiguard.com/outbreak-alert/darkside>>

321 *Stuxnet*. [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <<https://en.wikipedia.org/wiki/Stuxnet>>

322 Stuxnet bol vyvinutý pravdepodobne v roku 2005, jeho nasadenie s infikovaním iránskych SCADA systémov sa uskutočnilo v roku 2007. V roku 2009 bol Stuxnet zaznamenaný aj v iných zariadeniach a v roku 2010 sa začal šíriť vo väčšom meradle.

NSA – minimálne po odhalení – Stuxnet uvádzala pod názvom Olympijské hry (Olympic Games).

centrifúgy (odstredivky) používané na separáciu jadrového materiálu.

Po stránke technologickej Stuxnet spôsobil v bezpečnostnej komunite šok, keďže využíval ukradnuté originálne bezpečnostné certifikáty, dokázal atakovať veľmi špecifický hardvér³²³, využíval štyri softvérové chyby nultého dňa (to najlepšie, o čom drvivá väčšina kybernetickej scény ani nevedela), softvérový útok vykonával veľmi cielene a nezaútočil na čokoľvek a akokoľvek, dokázal rozlíšiť dôležitosť cieľa a podľa toho útočiť a i keď infikoval zariadenia prakticky po celom svete, vedel sa aktivovať len v konkrétnych destináciách v Iráne. Stuxnet využíval aj vynikajúce krytie, či už na tú dobu pokročilé techniky skrývania v systéme, alebo priamo počas útoku, keď centrifúgy privádzal k seba deštrukcii, obsluhu totálne zmiatol podávaním klamlivých informácií na ovládacích paneloch. Nikto nevedel, čo sa vlastne deje.

Naviac dizajn a architektúra (i použitý vektor útoku) Stuxnetu nie sú špecifické len pre špecializované verzie PLC v centrifúgach, ale môžu byť prispôsobené ako platforma pre útoky na široké spektrum moderných systémov SCADA a PLC (napr. v továrenských montážnych linkách alebo elektrárnach, rozvodných sieťach, plynovodoch a ropovodoch, atď.), z ktorých väčšina sa nachádza v Európe, Japonsku a Spojených štátoch...

Stuxnet vznikol v čase rozvíjajúcich sa nukleárných ambícií Iránu s potvrdeným spravodajským materiálom o obohacovaní uránu vo veľkom. Podľa zdroja z NSA išlo o politické rozhodnutie na medzinárodnej a medzirezortnej úrovni so zámerom vyvinúť a nasadiť nový druh zbraní. Aktívne sa na vývoji podieľali USA (CIA, NSA) a Izrael (Mosad), pričom spravodajsky celú operáciu podporila aj britská GCHQ.³²⁴

NSA síce nemá právomoc zasahovať, no už skôr bol vytvorený US Cyber Commands³²⁵ – vojenské oddelenie, ktoré bolo zodpovedné za národnú bezpečnosť v kybernetickom

323 V rámci SCADA išlo o PLC (programovateľné logické radiče) od Siemensu, konkrétne špecifické dve verzie frekvenčných meničov, použitých v centrifúgach jadrových zariadení v Iráne.

324 Koniec vlády prezidenta Busha bol sprevádzaný **zvyšujúcim sa tlakom na financovanie útočných kybernetických zbraní**. V rámci ministerstva obrany, ktoré v tom čase viedol Robert Gates, prevládala neistota, kam zaradiť kybernetické zbrane, preto bol Stuxnet priradený nie armáde, ale spravodajským službám. To bol dôvod pre účasť CIA, NSA, Mosadu a GCHQ. Podľa niektorých zdrojov bol do zbierania informácií a dodania USB kľúča so Stuxnetom zaangažovaná aj holandská spravodajská agentúra AIVD.

325 US Cyber Commands vzniklo v rámci operácie Buckshot Yankee, t.j. objavenia a riešenia prieniku do tajných sietí Ministerstva obrany USA (CENTCOM) v r. 2008.

priestore a majúce spoločné velenie s NSA.³²⁶

Samotný Mosad mal už v tom čase niečo podobné – kybernetickú jednotku vojenskej rozvedky 8200, špecializujúcu sa na elektronický a kybernetický boj.

Celú operáciu viedla CIA, pričom Mosad bol ten, kto prakticky určoval tempo.

US Cyber Commands spolu s jednotkou 8200 boli zodpovedné za vytvorenie kódu, testovanie prebehlo v testovacom stredisku NSA Oak Ridge na odstredivkách, ktoré odovzdala Líbya. Výsledok testovania bol predložený do Bieleho domu na schválenie operácie.³²⁷

Ovocie je známe, keďže úspešnému útoku Stuxnetu sa pripisuje niekoľkoročné zdržanie iránskeho jadrového programu spôsobením fyzickej deštrukcie približne pätiny iránskych centrifúg (nehovoriac o následnej anabáze s odstavením, analýzou a preverovaním, odvírením a postupnou obnovou funkčnosti ostatných zariadení).

Na začiatku nášho predstavenia Stuxnetu sme uviedli, že predstavoval v zásade prvý, zároveň najväčší a najsofistikovanejší prípad štátom, resp. viacerými štátmi vykonaného kybernetického útoku. Išlo o dôležitý míľnik, ktorý ukázal vysokú efektívnosť využitia kybernetických nástrojov pre spravodajské služby a armádu.³²⁸

Išlo o štátom riadený útok, na ktorom participovali viaceré štáty a ich rôzne inštitúcie. Stuxnetom tak nastal na tejto úrovni posun od koordinovaného získavania kvalitných

326 Na operácii sa podieľalo i oddelenie TAO-S321, hackerské oddelenie, ktoré ako jediné malo právomoc kompromitovať (hackovať) a vstupovať do cudzích systémov.

327 Vzhľadom na schválenie úradujúceho prezidenta Georga W. Busha bolo do kódu Stuxnetu implementované časové obmedzenie jeho aktívneho pôsobenia (niekoľko dní pred inauguráciou nového prezidenta). Po nástupe Baracka Obamu bolo treba nové schválenie, ktoré prišlo hneď v prvom roku jeho vlády (2009).

328 Zo známych jednotiek kybernetického boja môžeme okrem uvedených v kauze Stuxnet spomenúť aspoň niektoré ďalšie:

Napr. čínska jednotka č. 61419 spadá pod špeciálnu zložku Sily strategickej podpory, ktorá má okrem iného na starosti kybernetické vedenie vojny.

The People's Liberation Army Strategic Support Force. [on-line]. [cit. 13. marca 2022].

Dostupné na internete: <<https://jamestown.org/program/the-peoples-liberation-army-strategic-support-force-update-2019/>>

Alebo hackerská a špiónážna skupina APT28, známa aj pod označením Fancy Bear, ktoré je podľa nórskej tajnej služby (PST) napojená na ruskú vojenskú rozvedku GRU, konkrétne na jej 85. centrum špeciálnych operácií (GTsSS).

spravodajských informácií a obranných operácií k odhodlaniu vykonať útočné aktivity proti inému štátu.

V rámci vojnového použitia ide o ďalšiu z revolučných zmien, ktoré menia pohľad na bojové operácie – tie sú vykonávané elektronicky, pritom však s reálnym a ťažkým dopadom na protivníka, môžu byť realizované na veľké vzdialenosti a prakticky v neobmedzenom rozsahu, môžu mať extrémne rýchly či naopak nebadaný priebeh a minimálne znaky identifikujúce útočiace či bojujúce strany.

Stuxnet bol zároveň súčasťou zmeny v činnosti spravodajských služieb s dopadom na armádne zložky, keď miesto obranných aktivít v kybernetickom priestore sa začal dávať dôraz na útočné aktivity voči iným štátom, resp. aktérom konfliktu. Ruka v ruke s tým, podobne ako v iných oblastiach vojenského nasadenia elektronických systémov, badať stieranie rozdielov medzi obrannými a útočnými aktivitami – tie isté kybernetické zbrane môžu sledovať či brániť, ale i útočiť.

V poznámkach pod čiarou sme vyjadrili, že už koniec vlády prezidenta Georga W. Busha bol sprevádzaný zvyšujúcim sa tlakom na financovanie útočných kybernetických zbraní. Vzhľadom na intenzívne využívanie prostriedkov IKT v tzv. kritickej infraštruktúre USA prinieslo následné obdobie Obamovej administratívy oficiálny dôraz na kybernetickú obranu, pričom sa zdalo, že útočné kybernetické zbrane boli akoby tabu. Z vtedajších rozpočtových údajov však jasne vyplýva, že väčšina kybernetických vojenských výdavkov predsa len išla na útočné zbrane.³²⁹ Toto tvrdenie ostatne podporuje už spomenuté opätovné schválenie Stuxnetu v roku 2009 a napríklad aj neskôr zverejnené štatistiky využívania dronov a všeobecne LAWS v konfliktoch na Blízkom východe.

Stuxnet sa tak stal priekopníkom vývoja nového druhu zbraní a spôsobu ich nasadenia. Bol nie evolúciou, ale revolúciou vo svete kybernetických hrozieb. A i keď priamo v kóde Stuxnetu by sme márne hľadali súčasné technológie umelej inteligencie, išlo o samostatne pracujúci kód, vykonávajúci útočné operácie bez priameho riadenia vzdialeným operátorom.

Jedným z dôsledkov úspešného nasadenia Stuxnetu boli i prvé počiny na poli zaangažovania umelej inteligencie v kybernetických zbraňových systémoch, a to v oblasti obranných systémov. Išlo totiž o hrozby, ktoré neboli dopredu presne definované –

329 Výdavky boli skryté pod rôzne kódové označenia, napr. 10 CN0 (10: armádne, CN0: computer network operations, a pod.).

nebezpečné zraniteľnosti nultého dňa (t.j. dovtedy neznáme), sofistikovaný spôsob šírenia a maskovania kódu, novátorský prístup útoku a pod. Pre moderné systémy AI a algoritmy hlbokého učenia to bola ponuka, ktorá sa neodmieta a preto logicky prišlo k adaptácii týchto technológií aj v oblasti vojenských kybernetických systémov.³³⁰

V súčasnosti sú technológie umelej inteligencie plne integrované v obranných vojenských systémoch, pričom ich význam rastie ruka v ruke s ich technologickým rozvojom. Postupne sa prichádza aj k výraznejšiemu využitiu v scenároch útočného nasadenia.

Tu však môže byť problém. Už Stuxnet ukázal, že **ani tak sofistikovaná zbraň nie je úplne pod kontrolou, keďže sa rozšíril aj do iných systémov, než na ktoré bol zameraný a zaútočil aj na iné služby a organizácie.**

Aby toho nebolo málo, do hry vstupovali aj neštátni aktéri, napr. Equation Group³³¹ (napojená na už skôr spomínanú jednotku TAO z NSA), ktorí stáli za niektorými zneužitými zraniteľnosťami a spolupodielali sa na vývoj ďalších kybernetických zbraní (špiónážny Flame, súrodenec Stuxnetu Duqu,...) a pod. Z ich produkcie pochádza aj šírenie ďalších verzií Stuxnetu.

Celú kauzu nakoniec zavŕšila hackerská skupina The Shadow Brokers,³³² keď v roku 2017 zverejnila obrovský súbor nástrojov patriacich skupine Equation Group. Spolu so zdrojovými kódmi, ktoré už skôr The Shadow Brokers zverejnili, sa tak prakticky ku každému potencionálnemu aktérovi na poli kybernetických zbraní dostala dostatočná zbrojárska výbava.³³³

V tejto konštelácii však ťažko hovoriť o bezpečnej integrácii algoritmov umelej inteligencie a kybernetických zbraní. Veď ako dokáže ktorýkoľvek aktér napr. zabezpečiť dostatočne veľkú, relevantnú a kvalitnú vzorku trénovacích dát na spoľahlivé vytrénovanie systému? Ako dokáže aktér, ktorý útočný kód len prevzal, prípadne pozmenil, domyslieť, ako sa bude daný systém AI správať v neštandardných situáciách? Podobne môžeme klásť

330 Tieto technológie sú v súčasnosti implementované aj v komerčných bezpečnostných produktoch.

331 *Equation Group*. [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Equation_Group>

332 *The Shadow Brokers*. [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/The_Shadow_Brokers>

333 Táto skutočnosť plne zapadá do nášho zaradenia kybernetických zbraní medzi prostriedky vojen štvrtej generácie.

otázky prakticky v každej oblasti rizík, ktoré sme v kapitolách 2.1. až 2.4. predstavili.

A ak technológie umelej inteligencie spojíme s klasickými hackerskými technológiami, ako sú napr. stealth technológie, obfuskovanie kódu, rootkity či polymorfný kód, vieme garantovať bezpečnú funkcionálnosť takéhoto systému?

Aj z vlastnej praxe na poli kybernetickej bezpečnosti môžeme povedať, že o kvalite a robustnosti obranných kybernetických systémov využívajúcich algoritmy umelej inteligencie sme sa presvedčili. Avšak **útočné kybernetické systémy využívajúce súčasné technológie umelej inteligencie považujeme za veľmi rizikové a nevhodné pre vojenské nasadenie**. Ich dôsledky by totiž mohli byť katastrofálne.

Kybernetické zbrane sa v súčasnosti naplno integrujú do armádneho portfólia, či už ako priama súčasť a podpora vojenských aktivít³³⁴ alebo ako samostatné zbraňové systémy.

2.7.8. Ďalšie etické a právne aspekty

Či už ide o využívanie LAWs, dôsledky kybernetického boja alebo zneužitie sily spôsobené nasadením akýchkoľvek iných vojenských systémov AI – k otázkam a výhradám, ktoré sme v rámci ich predstavenia priebežne uvádzali, môžeme pripojiť i ďalšie etické a právne aspekty, zosumarizované v už skôr citovanom dokumente *Pros and Cons of Autonomous Weapons Systems*.³³⁵

Viaceré z výhrad boli spísané v rámci otvorených listov a výziev požadujúcich či už zakázanie alebo zmysluplné technické a legislatívne obmedzenie autonómnych zbraňových systémov.

334 Napríklad začiatok aktuálneho vojnového konfliktu na Ukrajine, ktorý v závere nasledovnej kapitoly tiež spomíname, bol sprevádzaný podpornými útočnými kybernetickými operáciami: po DDoS útoku, ktorý dočasne odstavil mnohé ukrajinské webové stránky, nasledovala kompromitácia škodlivými kódmi HermeticWiper, IsaacWiper a HermeticWizard.

HermeticWiper: New data-wiping malware hits Ukraine. [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <<https://www.welivesecurity.com/2022/02/24/hermeticwiper-new-data-wiping-malware-hits-ukraine/>>

IsaacWiper and HermeticWizard: New wiper and worm targeting Ukraine. [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <<https://www.welivesecurity.com/2022/03/01/isaacwiper-hermeticwizard-wiper-worm-targeting-ukraine/>>

335 ETZIONI, ETZIONI, *Pros and Cons of Autonomous Weapons Systems* [on-line]. [cit. 5. marca 2022].

Dostupné na internete: <<https://www.armyupress.army.mil/Journals/Military-Review/English-Edition-Archives/May-June-2017/Pros-and-Cons-of-Autonomous-Weapons-Systems/>>

Už v júli roku 2015 bol na jednej z medzinárodných konferencií o umelej inteligencii zverejnený otvorený list, ktorý vyzýval na zákaz autonómnych zbraní. V liste sa priamo píše: „technológia umelej inteligencie dosiahla bod, v ktorom je nasadenie týchto systémov prakticky, ak nie právne, uskutočniteľné v priebehu nie desaťročí, ale rokov a v stávke je veľa: autonómne zbrane sa uvádzajú ako tretia revolúcia v zbraňových systémoch, po strelnom prachu a jadrových zbraniach“. Tento **otvorený list poukazuje na pokrok i výhody spojené s vývojom systémov AI, zvažuje riziká a dopady zneužitia LAWs na celú oblasť umelej inteligencie a graduje k výzve na „zákaz útočných autonómnych zbraní, ktoré sú mimo zmysluplnej ľudskej kontroly“**.³³⁶

List má pôsobivý zoznam signatárov, medzi ktorými sú okrem iných Elon Musk (vynálezca a zakladateľ spoločnosti Tesla), Steve Wozniak (spoluzakladateľ spoločnosti Apple), fyzik Stephen Hawking³³⁷ (Univerzita v Cambridge) a Noam Chomsky (Massachusettský technologický inštitút). List podpísalo aj viac ako tritisíc výskumníkov v oblasti umelej inteligencie a robotiky...

Spolu s autormi článku *Pros and Cons of Autonomous Weapons Systems* len pripomíname to, čo sme pri opise LAWs už skôr uviedli: často nie je ľahké rozlíšiť, či ide o zbraň útočnú alebo obrannú (prieskumnú, atď.). V mnohých prípadoch aj neútočné smrtiace autonómne zbraňové systémy môžu pracovať v režime, ktorý by podľa vyššie uvedeného listu mohol byť eticky neprijateľný.

Etické výzvy idú v ruku v ruku s právnymi, s čím súvisí i problém zodpovednosti pri nasadení autonómnych zbraňových systémov. Etik Robert Sparrow upozorňuje na tento eticko-právny problém tým, že základná podmienka medzinárodného

336 *Autonomous Weapons: An Open Letter from AI [Artificial Intelligence] & Robotics Researchers*. Future of Life Institute website, 28 July 2015. [on-line]. [cit. 8. marca 2022].

Dostupné na internete: <<http://futureoflife.org/open-letter-autonomous-weapons/>>

Okrem špičiek výskumu a vývoja umelej inteligencie mohli byť signatármi tohto listu aj ľudia z iných oblastí, a tak autor tejto práce mal možnosť byť tiež jedným zo signatárov.

337 Nielen v tejto oblasti AI, ale aj vo viacerých ďalších Stephen Hawking varuje: „It will either be the best thing that's ever happened to us, or it will be the worst thing. If we're not careful, it very well may be the last thing.“

HIGGINS, A. *Stephen Hawking's final warning for humanity: AI is coming for us*. [on-line]. [cit. 13. marca 2022].

Dostupné na internete: <<https://www.vox.com/future-perfect/2018/10/16/17978596/stephen-hawking-ai-climate-change-robots-future-universe-earth>>

humanitárneho práva alebo *jus in bello* vyžaduje, aby za úmrtia civilistov niesla zodpovednosť nejaká osoba.³³⁸ Akákoľvek zbraň alebo iný vojnový prostriedok, ktorý znemožňuje určiť zodpovednosť za obeť, túto požiadavku nespĺňa, a preto by sa nemal používať vo vojnovom konflikte.³³⁹

Keďže zbraňové systémy vybavené umelou inteligenciou sa rozhodujú samy, je ťažké určiť, či chybné rozhodnutie je spôsobené priamo ľudským faktorom, alebo chybami v algoritmoch AI či v autonómnom uvažovaní týchto tzv. inteligentných strojov. Predstavme si napríklad situáciu, v ktorej autonómne vozidlo porušilo rýchlostné limity tým, že sa pohybovalo príliš pomaly na diaľnici, pričom následne nebolo jasné, komu by mala byť udelená pokuta...³⁴⁰ V situáciách, keď je rozhodnutie o použití smrtiacej sily na človeku, existuje jasný reťazec zodpovednosti, ktorý sa ťahne od toho, kto skutočne „stlačil spúšť“, až po veliteľa, ktorý vydal rozkaz. V prípade autonómnych zbraňových systémov takáto istota a jednoznačnosť neexistuje. **Nie je jasné, kto alebo čo má niesť vinu alebo zodpovednosť.**³⁴¹

Už v roku 2013 skupina inžinierov, odborníkov na umelú inteligenciu a robotiku a ďalších vedcov a výskumníkov z tridsiatich siedmich krajín sveta vydala „Výzvu vedcov na zákaz autonómnych smrtiacich robotov“. V tejto výzve sa uvádza, že chýbajú vedecké dôkazy o tom, že by roboty mohli v budúcnosti disponovať „funkciami potrebnými na presnú identifikáciu cieľa, situačné povedomie alebo rozhodnutia týkajúce sa primeraného použitia sily“. LAWs tak môžu spôsobiť vysokú mieru vedľajších škôd. Vyhlásenie sa končí zdôraznením, že „**rozhodnutia o použití hrubej sily sa nesmú delegovať na stroje**“.³⁴²

338 Niektorí protagonisti hypotetickej všeobecnej a uvedomelej AI smelo hovoria o nej ako o osobe.

Akékoľvek vnímanie AGI ako osoby by prinieslo celé spektrum neľahkých etických a právnych výziev.

Por. THURZO, V. *The Influence of Existentialism and Subjectivism on the Concept of the Human Person*. In *Spiritual and Social Experience in the Context of Modernism and Postmodernism*. Morrisville, 2021, s. 7-34.

339 SPARROW, R. *Killer Robots*. In: *Journal of Applied Philosophy*. 2007, 24, č. 1, s. 62–77.

340 Por. ETZIONI, A., ETZIONI, O. *Keeping AI Legal*. In: *Vanderbilt Journal of Entertainment & Technology Law* [on-line]. 2016, 19, č. 1, s. 133–146. [on-line]. [cit. 8. marca 2017].

Dostupné na internete: <http://www.jetlaw.org/wp-content/uploads/2016/12/Etzioni_Final.pdf>

341 ETZIONI, ETZIONI, *Pros and Cons of Autonomous Weapons Systems* [on-line]. [cit. 5. marca 2022].

Dostupné na internete: <<https://www.armyupress.army.mil/Journals/Military-Review/English-Edition-Archives/May-June-2017/Pros-and-Cons-of-Autonomous-Weapons-Systems/>>

342 *Scientists' Call to Ban Autonomous Lethal Robots*. ICRC, October 2013. [on-line]. [cit. 8. marca 2022].

Obava z delegovania rozhodovania o živote a smrti na neľudské subjekty sa v rámci autonómnych zbraňových systémov osobitne týka **LAWs, ktoré sú schopné vybrať si vlastné ciele**. Z tohoto dôvodu i rešpektovaný počítačový vedec Noel Sharkey vyzval na zákaz „smrtiaceho autonómneho zameriavania/cielenia“, pretože porušuje zásadu rozlišovania. Zásada rozlišovania sa považuje za jedno z najdôležitejších pravidiel ozbrojených konfliktov – nie je jednoduché pre autonómne zbraňové systémy rozlíšiť, kto je civilista a kto aktívne bojuje (napr. pri mestských bojoch je to ťažké určiť aj pre ľudí).³⁴³

Umožnenie systémom AI rozhodovať o zameraní cieľa s najväčšou pravdepodobnosťou povedie k civilným obetiam a neprijateľným vedľajším škodám.

V zásade ide o tie isté argumenty i závery, ktoré uvádzame naprieč celou 2. kapitolou o limitoch a rizikách súčasných systémov AI:

- „nevyspytateľnosť“ čiernej skrinky a rizikové faktory systémov umelej inteligencie (kapitola 2.1.)
- riziká spojené s procesnými útokmi na technológie umelej inteligencie (2.2.)
- realita rôznorodosti kybernetických útokov a reálna nemožnosť vytvoriť absolútne bezpečný systém (2.3.)
- technologická komplexnosť a nutnosť fungujúcej infraštruktúry (2.4.)
- toxické psychologické a sociologické dôsledky dopadajúce od jednotlivcov až po celú spoločnosť (2.5.)
- riziká špecifického zneužitia v nasadení, ktoré dávno presahuje pôsobnosť spravodajských služieb (2.6.)
- hrozby, ktoré sú spojené so smrtiacimi automatizovanými zbraňovými systémami, ktoré vďaka algoritmom AI dokážu prekonávať schopnosti človeka, no pravdepodobne bez možnosti reálne do nich implementovať etické pravidlá a legislatívne obmedzenia (2.6.)

A prichádzame k tomu istému záveru (už skôr sme ho formulovali ako požiadavku, ktorou sme končili kapitolu 2.4.): **nutnou podmienkou prevádzky ľubovoľného systému AI, ktorý môže predstavovať riziko pre akúkoľvek ľudskú osobu, je schopnosť**

Dostupné na internete: <<http://www.icrac.net/>>

343 SHARKEY, N. *Saying 'No!' to Lethal Autonomous Targeting*. In: *Journal of Military Ethics*. 2010, 9, č. 4, s. 369–383.

a možnosť človeka prebrať kedykoľvek kontrolu nad týmto systémom, resp. právo a možnosť verifikovať a prehodnotiť výsledky jeho činnosti.

Naše závery sú kompatibilné s myšlienkovým rámcom, ktorý zastávajú Sharkey, Sparrow. signatári vyššie uvedeného otvoreného listu a ďalší. Ide o návrh na stanovenie limitov pre technologický vývoj autonómnych zbraňových systémov a vytýčenie červených čiar a mantinelov, ktoré by budúci technologický vývoj nemal prekročiť. Len pripomíname, že tieto **limity, regulácia a obmedzenia LAWs by mali predstavovať etický rámec stanovený na základe morálnych hodnôt ľudskej spoločnosti.**

Viacerí odporcovia tohoto prístupu (prevažne z armádnych a politických kruhov) vidia v tomto postoji zbytočné a neúmerné obmedzovanie vývoja a nasadenia autonómnych smrtiacich systémov. Dávajú skôr prednosť „následnej regulácii“, pri ktorej by sa mal uplatňovať vyčkávací prístup a regulácia by sa mala odvíjať od toho, ako sa objavujú nové pokroky vo vývoji a možnostiach nasadenia LAWs. V zásade špekulujú o vývoji etiky a hodnotového rámca, prispôsobiac sa tak stavu vývoja zbraňových systémov na báze umelej inteligencie. Právni vedci, ako napríklad Kenneth Anderson a Matthew Waxman, ktorí tento prístup obhajujú, tvrdia, že regulácia bude musieť vznikáť priebežne spolu s technológiou a domnievajú sa, že etika a mantinely morálnej obhájitelnosti sa budú vyvíjať spolu s technologickým rozvojom. **Podporovatelia „následnej regulácie“ sa tak ľahko môžu dostať do oblasti hodnotového a morálneho relativizmu, snažiac sa prehodnotiť argumenty o nenahraditeľnosti ľudského svedomia a morálneho úsudku.**³⁴⁴

Zástancovia „následnej regulácie“ predpokladajú dosiahnutie tohoto cieľa prirodzeným vývojom, v rámci ktorého si ľudia budú postupne zvykať na stroje vykonávajúce funkcie s dôsledkami ohrozujúcimi život až po prípadnú smrť (napríklad riadenie automobilov alebo vykonávanie chirurgických operácií a pod.). Spoločnosť tak postupne začne akceptovať niečo podobné pri začleňovaní technológií umelej inteligencie do výzbroje.³⁴⁵

V tomto kontexte Anderson a Waxman odmietajú nami zastávaný etický rámec, regulácie a pevne stanovené podmienky pre vývoj a prevádzku LAWs. Prikláňajú sa skôr

344 Por. ANDERSON, K., WAXMAN, M. C. *Law and Ethics for Autonomous Weapon Systems: Why a Ban Won't Work and How the Laws of War Can*. Stanford University, Hoover Institution Press, Jean Perkins Task Force on National Security and Law Essay Series, 2013.

345 Podobnosť s inými oblasťami erózie hodnôt je úplne náhodná...

k vytvoreniu nejakého „komunitného riešenia“ s usmerneniami a zásadami, vyjadrujúcimi očakávania spoločnosti ohľadom legálneho a eticky vhodného správania a fungovania týchto systémov.³⁴⁶

I keď sme zástancami jasných etických pravidiel a noriem vyplývajúcich z morálnych hodnôt, niektoré postrehy podporovateľov „**následnej regulácie**“ nám môžu byť nápomocné pri uchopení etických výziev a problémov, ktoré sú a budú súčasťou extrémne rýchleho vývoja systémov umelej inteligencie, a to nielen v oblasti vojenského využitia.

Na margo eticko-právneho diskurzu, z ktorého veľká časť prebieha v armádných kruhoch USA, treba povedať, že súčasné nastavenie je priaznivé realizácii jasného etického rámca a pravidiel. Vládny Defense Innovation Board³⁴⁷ v roku 2019 v dokumente *AI Principles: Recommendations on the Ethical Use of Artificial Intelligence by the Department of Defense*³⁴⁸ navrhol princípy využívania systémov umelej inteligencie v americkej armáde, pričom je pre tieto princípy podstatné, že rozhodovacia právomoc zostáva na človeku, osobitne ak ide i misie s použitím smrtiacej sily.³⁴⁹

Medzi podstatné návrhy uvedené v dokumente môžeme zaradiť:

- už uvedená **zodpovednosť** (rozhodovacia právomoc) pri nasadení, v rámci ktorej by vojenský personál mal uplatňovať primeranú úroveň úsudku a zostať zodpovedný za vývoj, nasadenie, používanie a dôsledky využívania systémov AI.
- požiadavka na extrémnu **opatrnosť pri príprave tréningových dát**, aby sa systémy AI vyhli neúmyselnej zaujatosti/predsudkom (unintended bias), ktorých dôsledkom by boli nesprávne rozhodnutia v operačnom nasadení.³⁵⁰

346 Por. ANDERSON, K., WAXMAN, M. C. *Law and Ethics for Robot Soldiers*. In: *Policy Review*. 2012, 176, s. 46.

347 Dostupné na internete: <<https://innovation.defense.gov/>>

348 *AI Principles: Recommendations on the Ethical Use of Artificial Intelligence by the Department of Defense*. [on-line]. [cit. 9. marca 2022].

Dostupné na internete: <[https://admin.govexec.com/media/dib_ai_principles_-_supporting_document_-_embargoed_copy_\(oct_2019\).pdf](https://admin.govexec.com/media/dib_ai_principles_-_supporting_document_-_embargoed_copy_(oct_2019).pdf)>

349 TUCKER, P. The Pentagon's AI Ethics Draft Is Actually Pretty Good. [on-line]. [cit. 9. marca 2022].

Dostupné na internete: <<https://www.defenseone.com/technology/2019/10/pentagons-ai-ethics-draft-actually-pretty-good/161005/>>

350 Tzv. predsudky systémov AI ako dôsledok nesprávne zvolenej, či nekvalitnej množiny tréningových dát sme rozoberali v kapitole 2.1.2.

- **dosledovateľnosť**, v rámci ktorej musí existovať možnosť vrátiť sa späť k akémukoľvek dielčiemu výstupu alebo procesu a umožniť tak analýzu rozhodovacieho procesu algoritmov AI. K dosledovateľnosti prináleží požiadavka na kontrolovateľné metodiky, zdroje údajov, postupy a dokumentácie dizajnov systémov.³⁵¹
- **spoľahlivosť** ako výsledok presne definovanej oblasti využitia, zodpovedajúcej dizajnu a exaktnému vývoju systému AI, v kombinácii s celou paletou prísnych testov verifikujúcich spoľahlivosť v tejto definovanej oblasti.
- **ovládateľnosť**, t.j. schopnosť odhaľovať a brániť neúmyselnému poškodeniu alebo narušeniu činnosti systému AI kombinovaná s reálnou možnosťou zastavenia či prerušenia činnosti v prípade detekovaných problémov.

I keď sme to priamo zatiaľ neuviedli, zo všetkých oficiálnych aktivít na poli nasadenia vojenských systémov umelej inteligencie vyplýva, že **jednotlivé krajiny sa nedokážu vzdať tak lákavej technologickej výhody a sú pevne rozhodnuté technológie umelej inteligencie implementovať v celej šírke možného zmysluplného využitia.**

Navrhnuté pravidlá pre využívanie vojenských systémov umelej inteligencie v armáde USA, ktorá je prakticky najďalej s vývojom a nasadením technológií AI, poukazujú tiež na **jasný zámer definovať a implementovať pevný etický rámec**, bez ktorého by sa rozmachom armádných systémov AI otvorila Pandorina skrinka.

Vzhľadom na spektrum štátov, ktoré pracujú na aktívnom nasadení technológií umelej inteligencie vo vojenskej oblasti, **je však diskutabilné, či analogický záujem o striktný etický rámec bude existovať aj u ostatných.**

Na základe doterajších skúseností v oblasti limitov a rizík súčasných systémov AI **je však otázne, či v celom spektre technológií a algoritmov umelej inteligencie bude možné pevný etický rámec dodržať.** Vieme si to predstaviť pri cielenom – čo môže znamenať aj limitujúcom a tým aj znevýhodňujúcom – dizajne systémov AI s dôrazom na dodržiavanie navrhnutého rámca (ethics by design), avšak pri súčasnej snahe o čo najširšie adoptovanie umelej inteligencie v armáde a získanie konkurenčnej výhody sme v tomto smere skeptickí.

³⁵¹ Ide o našu známu transparentnosť, aj keď treba povedať, že vzhľadom na fenomén black boxu ťažko uveriť, že pri skutočne zložitých systémoch AI by toho vojenská špecialisti boli schopní.

Ako dovetok tejto kapitoly môžeme uviesť vyjadrenie ruského ministra obrany Sergeja Šojgu v jednom z interview na medzinárodnom vojenskom a technologickom fóre Army-2020: „Už to nie sú veci, ktoré vyzerajú ako hračky z detského sveta, je to niečo, čo predstavuje vážne nebezpečenstvo a seriózny zbraňový systém, pokiaľ je to reálne riadené neurónovými sieťami a umelou inteligenciou...“³⁵²

Stat' o systémoch umelej inteligencie v armáde je písaná v prvých týždňoch vojnového konfliktu na Ukrajine. Zo správ, ktoré prichádzajú z brífingov velení priamo či nepriamo zúčastnených strán tohoto konfliktu i z monitoringu na rôznych platformách zverejňovanej komunikácie členov armád a bojujúcich skupín³⁵³ badať veľmi nešťastnú kalkuláciu, ktorá tento prevažne nesymetrický konflikt, doplnený okrem bežnej vojenskej techniky aj niektorými skutočne modernými zbraňovými systémami, sprevádza: civilné obyvateľstvo a bezbranní ľudia sa používajú ako súčasť obrannej i útočnej taktiky a sú vystavení kalkuláciám, v ktorých na prvom mieste je čokoľvek iné, než dobro človeka.

Tento fenomén je známy i z rôznych konfliktov na blízkom východe. Ak ho navyše skombinujeme s nasadením smrtiacich automatizovaných zbraňových systémov, ohrozenie ľudských bytostí môže smerovať ku katastrofe.

2.8. Môžeme umelej inteligencii dôverovať?

Podľa štúdie zverejnenej tímom odborníkov na stránkach Scientific Reports algoritmy strojového učenia začali zasahovať do úloh, ktoré boli tradične vyhradené pre ľudský úsudok, a sú čoraz schopnejšie podávať dobré výkony v nových a náročných úlohách.

Na troch rôznych experimentoch tento tím ukázal, že ľudia sa s rastúcou náročnosťou úloh viac spoliehajú na algoritmické rady než na interakciu okolia. Ľudské subjekty mali tiež tendenciu výraznejšie ignorovať nepresné rady označené ako algoritmické v porovnaní s rovnako nepresnými radami označenými ako rady pochádzajúce od skupiny kolegov. Výrazne sa prejavil i vplyv sociálnych médií, online recenzií alebo osobných sociálnych sietí – tento vplyv je jednou z najdôležitejších síl, ktoré ovplyvňujú rozhodovanie

352 От робототехники до беспилотников: Шойгу рассказал о новинках ОПК на форуме «Армия-2020». [on-line]. [cit. 1. marca 2022].

Dostupné na internete: <<https://youtu.be/U8V4OZyNSac?t=152>>

353 Sociálna sieť Telegram si v tomto priamo žiada analytické nasadenie systému AI – medzi propagandou a dezinformáciami najrozličnejšieho druhu sa totiž nájde aj veľa faktických, operačných i taktických informácií z frontovej línie.

jednotlivcov.³⁵⁴

Javí sa, že pokiaľ prejavy a výstupy systémov umelej inteligencie spadajú do rámca exaktnosti, logiky a matematickej presnosti, t.j. do intelektuálneho rámca, ktorý sme si podvedome vytvorili, resp. prijali pre vedu a techniku, sme náchylní týmto systémom a priori dôverovať.³⁵⁵

Nerozporujúc úžasný potenciál a pokrok vo vývoji i úspechy v nasadení systémov ANI sme sa však v rámci tejto kapitoly potýkali s rôznymi chybami, limitmi a rizikami jej rôznych druhov a viacerých scenárov nasadenia.

Preto v kontexte rastúcej dôvery v systémy umelej inteligencie a vnímajúc problémy, s ktorými sa nasadenie týchto systémov potýka i oblastí, ktorých sa ich využitie dotýka, môžeme a musíme stanoviť podmienky, bez splnenia ktorých by nasadenie systémov AI do reálneho sveta, v ktorom interagujú s človekom a vplývajú na spoločnosť, nemalo byť umožnené.

Systémy umelej inteligencie vo všeobecnosti musia byť:

- **legálne** – vyhovovať požadovaným normám, zákonom a reguláciám
- **etické** – spĺňať požadované etické kritériá
- **bezpečné** – dosahovať potrebné štandardy bezpečnosti a robustnosti

Len takto sa môžeme priblížiť k cieľu, ktorým je **dôveryhodná umelá inteligencia**. Hovoríme tak o **systémoch AI, ktoré sú zamerané na človeka** a podporujú zodpovedné riešenia dbajúce na ľudské potreby, bezpečnosť a súkromie.³⁵⁶

354 BOGERT, E., SCHECTER, A., WATSON, R. T. *Humans rely more on algorithms than social influence as a task becomes more difficult*. In: *Sci Rep*. [on-line]. 2021, 11, 8028. [cit. 20. februára 2022].

Dostupné na internete: <<https://doi.org/10.1038/s41598-021-87480-9>>

355 Naviac, táto dôvera rastie úmerne náročnosti úloh, ktoré s pomocou týchto systémov máme realizovať, a to v kontexte postmodernej rezignácie na rúco neveští nič dobrého.

356 *Trustworthy AI is human-centered*. [on-line]. [cit. 20. februára 2022].

Dostupné na internete: <<https://www.ibm.com/watson/trustworthy-ai>>

3. Umelá inteligencia a etika

Existujú len dva druhy priemyslu, ktoré svojim zákazníkom hovoria používatelia.

Nelegálne drogy a softvér.³⁵⁷

Pohybujúc sa už viac než dve desaťročia v oblasti kybernetickej bezpečnosti si uvedomujeme, že prakticky ešte nikdy nebolo badať toľké obavy z nasadenia novej technológie a tak seriózný záujem³⁵⁸ o etické otázky ako v prípade systémov umelej inteligencie. Kľúčové slovo či značka (hashtag) #AIEthics v komunite odborníkov neoznačuje okrajovú záležitosť, ktorá je mimo zorného uhľa pohľadu tých, čo umelej inteligencii skutočne rozumejú, ale stáva sa súčasťou hlavného prúdu vývoja, implementácie a používania systémov AI v rámci reálneho sveta.

Aby sme si rozumeli, s futurologickými obavami takmer na úrovni sci-fi sa spoločnosť stretáva už od pionierskych čias formovania konceptu umelej inteligencie a prvých krokov na poli vývoja týchto systémov. V našom (a nielen našom) prípade však ide o seriózný záujem celej spoločnosti, ktorá je konfrontovaná s technológiami, ktoré majú potenciál hlboko zasahovať a ovplyvňovať jednotlivcov i celé spoločenské celky.

Tento seriózný záujem sa neobmedzuje len na oblasť umelej inteligencie, ale má oveľa širší záber, ktorý koreluje a má i kauzálnu súvislosť s treťou priemyselnou revolúciou, digitálnym vekom a zmenou paradigmy nastupujúcej informačnej spoločnosti, v rámci ktorej sa na základe zmien v technológiách a vzhľadom na prevratný vedecký i technologický pokrok mení medziľudská komunikácia vo svojej podstate, menia sa vzťahy, mení sa spoločnosť a princípy jej fungovania.³⁵⁹

Prvý zviditeľnením tohoto záujmu bol na verejnosť „prevalený“ zápas o ochranu osobných

357 Prof. Edward R. Tufte.

358 Už skôr boli komunikované témy a vybojované zápasy, bez ktorých by sme si moderný svet informačných technológií nevedeli predstaviť, napr. zápas open-source vs. closed software, legitimita slobodných licencií typu Creative Common a pod. Vždy však išlo z pohľadu spoločnosti len o okrajovú záležitosť.

359 Por. ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi*, [on-line], s. 20. [cit. 20. februára 2022].

Dostupné na internete:

<https://peter.santavy.cloud/docs/Analýza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

údajov a jasné pravidlá pre sledovanie a monitoring občanov i boj o prístup k informáciám s cieľom angažovať sa v oblasti ľudských práv.³⁶⁰ Išlo a ide o zápas, ktorý prináša svoje ovocie, či už ide o rast povedomia a verejnú diskusiu o ochrane osobných údajov a práve na súkromie, alebo o legislatívne rámce a regulácie, z ktorých najznámejšou a pravdepodobne i najúčinnějšíou je Nariadenie Európskeho parlamentu a Rady (EÚ) 2016/679 (General Data Protection Regulation – GDPR) o ochrane fyzických osôb pri spracúvaní osobných údajov a o voľnom pohybe takýchto údajov.³⁶¹

Ďalším krokom na ceste upevnenia záujmu celej spoločnosti o problémy, ktoré ju bytostne ovplyvňujú, je akcentovanie etického rozmeru širokej oblasti umelej inteligencie. Skúsme sa tomu venovať aj v tejto kapitole.

3.1. Aké etické výzvy prináša umelá inteligencia

Rozoberajúc limity a riziká súčasných systémov umelej inteligencie v druhej kapitole, sme nie raz vyvodzovali závery, ktoré možno naplno zaradiť do agendy etiky AI.

Pripomeňme si aspoň tie podstatné:

- ³⁶²**Súčasný systém AI sprevádzajú oprávnené a vážne obavy z toho, že ak nechápeme ako tieto systémy pracujú, nemôžeme im reálne dôverovať a ťažko dokážeme predpovedať okolnosti, za ktorých tieto systémy zlyhajú.**

Keďže v zásade nevieme, čo presne sa neurónová sieť naučila a ako spoľahlivo to dokáže aplikovať, prakticky každý sofistikovaný systém umelej inteligencie sa javí ako *black box* – čierna skrinka, ktorá niečo vykonáva, ale ako a prečo tak robí, nemusí byť jasné.

Naviac celkový dizajn, voľba algoritmov a nastavenie hyperparametrov, prípadne

³⁶⁰ Náhľad do týchto problémov prináša napríklad analýza viacerých káuz z oblasti úniku a zverejňovania tajných informácií (whistleblowing) – Cablegate z roku 2010 a Snowden vs. NSA začínajúca sa v roku 2013 s dosahom na celé posledné desaťročie, prípadne viaceré škandály Facebooku ohľadom zneužitia a monetizácie osobných údajov používateľov.

Por. ŠANTAVÝ, *Niektoré výzvy informačnej spoločnosti v oblasti morálnej teológie*, s. 117-135.

³⁶¹ Nariadenie Európskeho parlamentu a Rady (EÚ) 2016/679 z 27. apríla 2016 o ochrane fyzických osôb pri spracúvaní osobných údajov a o voľnom pohybe takýchto údajov, ktorým sa zrušuje smernica 95/46/ES (všeobecné nariadenie o ochrane údajov). [on-line]. [cit. 19. februára 2022].

Dostupné na internete: <<https://eur-lex.europa.eu/legal-content/SK/TXT/?uri=celex%3A32016R0679>>

³⁶² Kapitola 2.1.1.

niektoré ďalšie aspekty návrhu tých najlepších systémov AI sú postavené viac na odbornej intuícii a metóde pokus-omyl, než na exaktných postupoch, ktoré by mohli byť univerzálne aplikovateľné a tým ja vo väčšej miere verifikovateľné.

Problém s verifikovateľnosťou postupov a výsledkov činnosti systémov umelej inteligencie sa tak premieta i do oblasti implementácie etických pravidiel, či už ide o spôsob, ako hodnotiť a kontrolovať činnosť systémov AI, alebo ide priamo o spôsob implementácie regulácií technológiami AI.³⁶³

- ³⁶⁴Súčasný systémy umelej inteligencie trpia rôznymi zraniteľnosťami. **Musíme mať neustále na pamäti, že tieto systémy môžu zlyhávať najrozličnejšími a často neočakávanými spôsobmi** v dôsledku nemožnosti pripraviť dostatočne veľkú množinu tréningových dát, prípadne ich nesprávnej voľby bez dostatočnej kvality alebo s predsudkami, nadmernému prispôbovaniu sa tréningovým údajom, efektu dlhého chvosta, rizikám plynúcim z oklamania hlbokých sietí, ich zraniteľností a tzv. povier, pod čo sa môže podpísať i nedostatok odbornej erudovanosti potrebnej pre dizajn a prípravu funkčného a úspešného riešenia i ladenia hyperparametrov.

Nielen riziká priameho zlyhania, ale aj realita výsledkov, ktoré môže byť ťažké správne interpretovať (čo sa vlastne sieť naučila, čo výstup z daných dát na vstupe v skutočnosti znamená) a neschopnosť predvídať, kedy sa jednotlivé zlyhania prejavia (za akých podmienok, pri akej súhre okolností, s dôsledku akej dynamiky vnútorného vývoja, resp. činnosti systému AI).

S týmito zlyhaniami a obmedzeniami treba rátať aj v oblasti snahy o technologickú implementáciu etických noriem a mantinelov. Nemusí to byť vôbec triviálne a v niektorých scenároch nasadenia ani uskutočniteľné.

- ³⁶⁵**Systémy umelej inteligencie robia chyby, ktoré sú diametrálne odlišné od ľudských chýb a zlyhaní.** Preto môžu byť prekvapivé, neočakávané a nebezpečné.

Naviac, **ak by sme priamo algoritmami AI implementovali regulačné**

³⁶³ Vo viacerých oblastiach systémov AI ani nie je možné technologicky tieto dva rozmery implementácie od seba oddeliť – v praxi môže ísť o súčasť funkcionality, algoritmického návrhu a procesu učenia sa celého systému.

³⁶⁴ Kapitola 2.1.2.

³⁶⁵ Kapitola 2.1.2.

mechanizmy, musíme rátať s rovnakým rizikom nečakaných pochybení týchto mechanizmov počas prevádzky.

- ³⁶⁶Vzhľadom na rozličné typy útokov zameraných na procesy umelej inteligencie sa **etické výzvy neviažu len na návrh a realizáciu systémov AI, no rozširujú sa aj o oblasť etiky použitia a z toho prameniace riziká zneužitia.**

Ide teda o pohľad na **etický rozmer zodpovedného vývoja systémov a procesu učenia**, pri ktorom je jednou zo základných požiadaviek princíp „**etika už od návrhu**“ (**ethics by design**).

Minimálne v rovnakej miere ide aj o pohľad na **etický spôsob využívania technológií AI**, t.j. bez akejkoľvek snahy o ich kompromitáciu, resp. zneužitie proti iným, čo sa v nemalej miere premieta aj do problematiky kapitol 2.5. až 2.7.

- ³⁶⁷**Systémy umelej inteligencie sú vystavené aj problémom v oblasti kybernetickej bezpečnosti**, keďže mnohé riziká atakujúce bezpečnosť procesov umelej inteligencie, zneužívajúce nedostatky dizajnu alebo amplifikujúce dôsledky rizikových faktorov, ako vektor útoku používajú zraniteľnosti v oblasti kybernetickej bezpečnosti. Kybernetická bezpečnosť je proces, ktorý má svoju dynamiku a vyjadruje skutočnosť, že **ani v oblasti systémov umelej inteligencie neexistuje dokonale bezpečný a spoľahlivý systém.**

S rozvojom zložitých systémov AI treba mať neustále na pamäti, že počet zraniteľností je priamo úmerný komplexnosti systémov. Podobne i množstvo kompromitácií, resp. zneužití týchto zraniteľností priamo rastie s mierou nasadenia a využívania v reálnom svete.

Okrem opakujúcej sa požiadavky na etický rozmer zodpovedného vývoja a etického využívania systémov AI treba pamätať aj na relatívnu bezpečnosť a spoľahlivosť týchto systémov. Akákoľvek ich absolutizácia môže viesť nielen k ekonomickým škodám, ale v niektorých oblastiach aj k ohrozeniu človeka a spoločnosti, ktoré môže byť – v pohľade na možné spektrum nasadenia od autonómnych vozidiel až po algokratiu – fyzické, psychologické, morálne, právne,...

- ³⁶⁸Krátky náčrt problematiky technologického vybavenia a nutnej infraštruktúry

366 Kapitola 2.2.

367 Kapitola 2.3.

pre nasadenie systémov AI vo viacerých oblastiach reálneho sveta, ktorý sme uviedli v kapitole 2.4., poukazuje na veľkú komplexnosť týchto systémov. **A komplexnosť je problém – je nielen rizikovým faktorom bezpečnosti a stability fungovania systémov, ale mnohokrát i neľahkou výzvou pre úspešnú a jasnú implementáciu regulácií a etických pravidiel.**

- ³⁶⁹Úspešné nasadenie mnohých moderných systémov AI vyžaduje prísun extrémneho množstva dát z reálneho sveta a priamo z ľudského prostredia. Ľudia i celá spoločnosť sú preto vystavení neustálemu dohľadu a sledovaniu, ktoré je však len veľmi málo pod kontrolou, ak vôbec. **Pri veľkých systémoch zhromažďujúcich denno denne extrémne množstvo dát je veľmi ťažké zaručiť etiku spracúvania týchto údajov a vylúčiť riziko ich zneužitia.**

Systémy umelej inteligencie v oblasti médií a sociálnych sietí posúvajú paradigmatickú zmenu informačnej spoločnosti na vyššiu úroveň: miesto informácií sa základnou komoditou stávajú priamo ľudia, keďže prichádza k ovplyvňovaniu ľudského vnímania, rozhodnutí i konania a k neustálemu zlepšovaniu schopností systémov AI vytvárať modely predikujúce konanie konkrétneho človeka. **Systémy umelej inteligencie sa s veľkou mierou istoty stávajú schopnými vytvárať veľmi presné psychologické profily, odhaľovať akékoľvek väzby a mnohé osobné informácie, predpovedať a ovplyvňovať naše konanie. Tieto systémy však mnohokrát nie sú predmetom regulácie, resp. verejnej kontroly.**

Ak dosiahnutý efekt a vplyv môže byť tak závažný a prakticky takmer istý, existuje riziko i pokušenie využiť to, presnejšie – manipulovať a zneužiť (či už vo svete biznisu, ale aj v oblastiach oveľa kritickejších).

Neuvedomujúc si, ako je naša myseľ a psychika zraniteľná, tak **postupne prechádzame od technologického prostredia založeného na systémoch AI k prostrediu založenému na závislosti a manipulácii.**

Systémy AI sa podieľajú na stieraní hranice medzi pravdou a lžou, medzi realitou a fiktívnym svetom. Prichádza tak k extrémnemu rozdeleniu spoločnosti a narušeniu vzťahov, k erózii hodnôt a ideologickej divergencii, k nárastu zatvrdilej nevedomosti, ignorancii faktov a relativizácii pravdy, k neschopnosti riešiť aktuálne

368 Kapitola 2.4.

369 Kapitola 2.5.

civilizačné výzvy, ku kríze demokracií a spoločenských mechanizmov, k ekonomickým dôsledkom a pre mnohých i k potencionálnemu smerovaniu k civilizačnému kolapsu...

Míľnikom, ktorého by sme sa mali obávať, teda 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.

- ³⁷⁰Technológie samy o sebe nie sú tou priamou hrozbou - dokážu však prebudiť v človeku aj v spoločnosti tie najhoršie stránky ľudského bytia, ktoré sa tak stávajú existenčnou hrozbou.

V kontexte slabej umelej inteligencie (ANI) ešte stále konečnú voľbu cieľov nerobia stroje, ale človek, takže stále je v ľudskej moci tieto dôsledky ovplyvniť a zmeniť. Preto musia existovať etické pravidlá a regulácie, ktoré by dokázali chrániť jednotlivca i celú spoločnosť voči rizikám sofistikovanej práce systémov AI v rámci sociálnych sietí, systémov riadenia spoločnosti, resp. analogických platforiem.

- ³⁷¹**Viacere aplikácie systémov umelej inteligencie v reálnom svete budú musieť v zlomkoch sekúnd riešiť rôzne kritické situácie – a vedieť ich vyriešiť eticky.** Napríklad rozhodovanie autopilota plne autonómneho vozidla medzi potencionálnym ohrozením života posádky alebo ostatných účastníkov v čase blížiacej sa alebo prebiehajúcej dopravnej nehody.

Mnohé z týchto systémov sú a budú využívané v oblastiach, kde môže byť ohrozené zdravie a život človeka. Základné etické normy pre ich prevádzku musia byť legislatívne ustavené. Otázkou je, **do akej miery tieto eticko-právne regulácie budú aj technicky úspešne implementované.**

- ³⁷²**Vo viacerých oblastiach nasadenia (napr. autonómne vozidlá) bude treba vyriešiť rozdelenie zodpovednosti medzi človekom a riadiacim systémom AI i s tým súvisiace legislatívne požiadavky pre reálnu prevádzku týchto systémov.**³⁷³

370 Kapitola 2.5.

371 Kapitola 2.5.

372 Kapitola 2.5.

373 Napríklad v oblasti nasadenia autonómnych vozidiel do bežnej prevádzky sa diskutujú legislatívne

- ³⁷⁴Ako stanoviť **akceptovanú mieru neistoty v oblasti spoľahlivosti a robustnosti** systémov umelej inteligencie?
- ³⁷⁵**Fenomén digitálneho rozdelenia (digital divide)**, ktorý sa s neustálym rozvojom informačných technológií a prechodom k informačnej spoločnosti stáva reálnym problémom, môže byť vďaka necitlivému nasadeniu technológií umelej inteligencie oproti súčasnosti ešte znásobený. Či už ide o dopady v oblasti života spoločnosti (2.5.) alebo – a to je vážnejšie – v oblastiach súvisiacich s algoritmickým riadením spoločnosti (2.6.), napr. v dostupnosti elektronických služieb štátu, zdravotnej starostlivosti, vzdelávaní, finančných službách a pod.
- ³⁷⁶**Existuje reálne riziko zneužitia systémov AI v rámci algokracie na efektívnu kategorizáciu a obmedzovanie ľudských práv občanov.** Napr. čínsky systém sociálneho kreditu.

Technológie umelej inteligencie umožňujú aj za sprísnených legislatívnych obmedzení reálne zneužitie tajnými službami, vládnymi agentúrami a korporáciami v dohľadových, analytických a predikčných systémoch. **Aké sú možnosti nastavenia etických mantinelov a reálne dodržiavaných právnych rámcov v oblasti riadenia štátu, spravodajstva a dohľadu?**

- ³⁷⁷Trio technológií umelej inteligencie – sledovanie a dohľadové systémy spojené s analytickými nástrojmi a schopnosťou vytvárať modely predikujúce ľudské konanie či vývoj situácie – tvorí v súčasnosti skutočne silnú technologickú výbavu (nielen) spravodajských služieb.

Systémy umelej inteligencie, ktorým sa v súčasnosti analýza metadát, predikcia a tvorba záverov zveruje, tak de facto majú rastúci potenciál rozhodovať o bytí a nebytí, „generovať“ dôkazy, ktoré môžu byť použité napríklad pri obvinení z vlastizrady, identifikácii podozrivých osôb, dokazovaní nelegálnej činnosti a pod.

požiadavky na tzv. bezpečnostného operátora, hľadajúc objektívnu mieru zodpovednosti ľudského faktora pri riadení autonómneho vozidla v rámci jednotlivých stupňov automatizácie, ktoré sme uviedli v kapitole 2.5.

374 Kapitola 2.5.

375 Kapitola 2.5. a 2.6.

376 Kapitola 2.6.

377 Kapitola 2.6. a 2.7.1.

Ako realizovať striktnú zákonnosť i dôsledný dohľad demokraticky zvolených zástupcov a spoločnosti pri nasadení systémov AI v rámci spravodajských služieb, dohľadových systémov a všetkých foriem i stupňov algoritmického riadenia spoločnosti?

Uvedomujeme si, že technológie využívajúce algoritmy umelej inteligencie dokážu byť veľmi účinným nástrojom pre spravodajské služby a dohľad s takým presahom do ostatných oblastí verejného života a štátu, ktorý môže znamenať veľký posun v procesoch a spôsobe fungovania spoločnosti. Avšak vzhľadom na politický, ideologický i spoločenský kontext v rôznych častiach sveta a riziká systémov AI treba zabezpečiť, aby prostriedky umelej inteligencie neboli zneužitá alebo sa vymkli kontrole.

Preto je treba v celom spektre nasadenia od spravodajských služieb až po algoritmické riadenie akcentovať veľkú rozvážnosť a vyžadovať striktnú zákonnosť i dôsledný dohľad demokraticky zvolených zástupcov, obmedziť dopady na sociálnu spravodlivosť a zabezpečiť dodržiavanie ľudských práv a hodnôt.

- ³⁷⁸Pokročilé modelovanie technológií, simulácia priebehu vojenských operácií i priebehu častí konfliktu a silný virtualizačný efekt môžu viesť k morálnemu znecitliveniu obsluhujúceho personálu a velenia. Ovocím môže byť modelovanie nových druhov zbraní hromadného ničenia, návrh vojenských operácií bez vnímania strát na životoch a utrpenia civilného obyvateľstva a pod.

Rizikovým faktorom je teda psychologický aspekt vytvorenia si podvedomého chápania boja ako virtuálnej reality. **Dôsledkom môže byť strata citlivosti pre ohrozenie ľudských životov a reálnych hodnôt pri nasadení v skutočnom boji.** Tento rizikový faktor je o to záľudnejší, že sa „pod kožu“ dostáva postupne a podvedome, mnohokrát bez varovnej reakcie svedomia.

- ³⁷⁹Vďaka technológiám umelej inteligencie sú vo viacerých scenároch nasadenia **autonómne zbraňové systémy lepšie ako ľudia**, pričom treba rátať s multiplikačným efektom účinnosti ich nasadenia. Preto bude obmedzovanie nasadenia týchto systémov veľmi ťažké.

378 Kapitola 2.7.2. - 2.7.4.

379 Kapitola 2.7.5.

LAWs (autonómne zbraňové systémy so smrtiacim účinkom) využívajúce algoritmy AI dokážu byť neuveriteľne technologicky adaptabilné, relatívne lacné a minimálne rovnako nebezpečné ako najmodernejšie zbraňové systémy riadené ľudskou obsluhou. **Ako dokážeme zabezpečiť, že smrtiace autonómne systémy, ktoré sú zámerne dizajnované tak, aby dokázali efektívne pracovať aj pri strate spojenia, budú vedieť správne použiť svoju smrtiacu silu a rešpektovať požadované mantinely, ak je nám známe, že systémy AI robia mnohokrát iné chyby ako ľudia a nevieme ich vytrénovať tak, aby vedeli adekvátne odpovedať na všetky reálne situácie?**

- ³⁸⁰Pri skupinovom riadení autonómnych zbraňových systémov sa pýtame – je koordinované riadenie na základe interakcie jednotlivých systémov AI natoľko robustné, aby prepojené LAWs dokázali správne pracovať aj pri výpadku kritického množstva zapojených automatických systémov?

Vieme dostatočne odhadnúť správanie centrálne riadeného (a zároveň parciálne autonómneho) alebo koordinovane riadeného skupinového automatického zbraňového komplexu vo vyhranených situáciách?

Nehrozí, že vzhľadom na znásobenú a komplexnejšiu účinnosť celej skupiny LAWs riadených umelou inteligenciou pri zlyhaní tohto riadiaceho systému AI budú jeho jednotlivé časti konať bez obmedzení a životy ohrozujúcim spôsobom?

- ³⁸¹Vďaka nebyválnemu rozmachu informačných technológií, ich prieniku prakticky do všetkých oblastí života rodí sa informačnej spoločnosti a strategickej dôležitosti pre kritickú infraštruktúru štátov i vďaka pokročilým technológiám na poli kybernetickej bezpečnosti a kriminality, ktoré úspešne implementujú algoritmy umelej inteligencie, sa vedenie vojenských operácií v kybernetickom priestore stalo neoddeliteľnou súčasťou portfólia moderných armád.

Technológie umelej inteligencie nasadené v obranných kybernetických systémoch sú účinnou – a treba zdôrazniť aj legitímnu – pomocou pri ochrane nielen vojenských, ale akýchkoľvek dôležitých aktív spoločnosti.

380 Kapitola 2.7.6.

381 Kapitola 2.7.7.

Avšak útočné kybernetické zbraňové systémy, ktoré svojimi doterajšími úspechmi v reálnom nasadení dokázali svoju strategickú opodstatnenosť, obnášajú riziká, ktoré treba eliminovať. Už Stuxnet ukázal, že ani tak sofistikovaná zbraň nie je úplne pod kontrolou, keďže sa rozšíril aj do iných systémov, než na ktoré bol zameraný a zaútočil aj na iné služby a organizácie. Kybernetické zbraňové systémy využívajúce algoritmy AI principiálne nie sú imúnne voči žiadnemu z rizík, ktoré sme uviedli v kapitolách 2.1. až 2.4., pritom však ich dôsledky sa prejavujú v reálnom svete – od ohrozenia konkrétnych životov cez zlyhanie dôležitej infraštruktúry až po havárie kritických systémov spoločnosti.

Preto útočné kybernetické systémy využívajúce súčasné technológie umelej inteligencie považujeme za veľmi rizikové a nevhodné (nielen) pre vojenské nasadenie.

- ³⁸²Vzhľadom na riziká a dopady zneužitia LAWs na celú oblasť vývoja a využitia systémov umelej inteligencie sa v poslednom desaťročí stupňuje apel mnohých odborníkov z tejto oblasti **proti akémukoľvek nasadeniu útočných autonómnych zbraní, ak sú mimo zmyslupnej ľudskej kontroly.**

V súčasnosti **nemáme vedecké dôkazy o schopnosti automatizovaných systémov disponovať funkciami potrebnými na presnú identifikáciu cieľa, situačné povedomie alebo rozhodnutia týkajúce sa primeraného použitia sily.** Umožnenie systémom AI rozhodovať o zameraní cieľa s najväčšou pravdepodobnosťou povedie k civilným obetiam a neprijateľným vedľajším škodám. **LAWs tak môžu spôsobiť vysokú mieru vedľajších škôd a preto sa rozhodnutia o použití hrubej sily nesmú delegovať na stroje.**

Obava z delegovania rozhodovania o živote a smrti na neľudské subjekty sa v rámci autonómnych zbraňových systémov osobitne týka LAWs, ktoré sú schopné vybrať si vlastné ciele. Schopnosť strojového smrtiaceho autonómneho zameriavania a cielenia porušuje zásadu rozlišovania, ktorá sa považuje za jedno z najdôležitejších pravidiel ozbrojených konfliktov – napríklad nie je jednoduché pre autonómne zbraňové systémy rozlíšiť, kto je civilista a kto aktívne bojuje.

- ³⁸³**Schopnosti umelou inteligenciou poháňaných autonómnych zbraňových**

382 Kapitola 2.7.8.

383 Kapitola 2.7.8.

systémov a kybernetických zbraní i napriek všetkým rizikám vedú k zvyšujúcim sa tlakom na financovanie a zavádzanie útočných kybernetických zbraní. Jednotlivé krajiny sa nedokážu vzdať tak lákavej technologickej výhody a sú pevne rozhodnuté technológie umelej inteligencie implementovať v celej šírke možného zmysluplného využitia.

Problematika obmedzenia týchto útočných systémov je skomplikovaná aj reálnym stieraním hraníc medzi obranným a útočným nasadením takmer vo všetkých oblastiach vojenského využitia technológií umelej inteligencie.

- ³⁸⁴**Nutnou podmienkou prevádzky ľubovoľného systému AI, ktorý môže predstavovať riziko pre akúkoľvek ľudskú osobu, je schopnosť a možnosť človeka prebrať kedykoľvek kontrolu nad týmto systémom, resp. právo a možnosť verifikovať a prehodnotiť výsledky jeho činnosti.**

Limity, regulácia a obmedzenia LAWS by mali predstavovať etický rámec stanovený na základe morálnych hodnôt ľudskej spoločnosti, nie na základe relativistickej tzv. „následnej regulácie“.

Na základe uvedomenia si problémov spojených so smrtiacim využitím algoritmov AI a angažovanosti odbornej verejnosti si vlády a armády viacerých krajín uvedomujú potrebu regulačného rámca. Napríklad navrhnuté pravidlá pre využívanie vojenských systémov umelej inteligencie v armáde USA, ktorá je prakticky najďalej s vývojom a nasadením technológií umelej inteligencie, poukazujú i na jasný zámer definovať a implementovať pevný etický rámec, bez ktorého by sa rozmachom armádných systémov AI otvorila Pandorina skrinka.

Na základe doterajších skúseností v oblasti limitov a rizík súčasných systémov AI je však otázne, či v celom spektre technológií a algoritmov umelej inteligencie bude možné pevný etický rámec dodržať. Vieme si to predstaviť pri cielenom – čo môže znamenať aj limitujúcom a tým aj znevýhodňujúcom – dizajne systémov AI s dôrazom na dodržiavanie navrhnutého rámca (ethics by design), avšak pri súčasnej snahe o čo najširšie adoptovanie umelej inteligencie v armáde a získanie konkurenčnej výhody sme v tomto smere pomerne skeptickí.

384 Kapitola 2.7.8.

- ³⁸⁵Podľa štúdie *Humans rely more on algorithms than social influence as a task becomes more difficult*, pri ktorej sme sa záverom druhej kapitoly pristavili, sme náchylní týmto systémom umelej inteligencie a priori dôverovať za podmienky, že ich prejavy a výstupy spadajú do rámca exaktnosti, logiky a matematickej presnosti, t.j. do intelektuálneho rámca, ktorý sme si podvedome vytvorili, resp. prijali pre vedu a techniku.

A priori dôvera – nedbajúc na limity a riziká technológií AI – môže byť veľmi nebezpečným trendom a v určitých rozmeroch (psychologický a sociologický vplyv, algokracia, vojenské nasadenie) aj ohrozením fungovania modernej spoločnosti.

Preto v kontexte rastúcej dôvery v systémy umelej inteligencie a vnímajúc problémy, s ktorými sa nasadenie týchto systémov potýka i oblastí, ktorých sa ich využitie dotýka, môžeme a musíme stanoviť podmienky, bez splnenia ktorých by nasadenie systémov AI do reálneho sveta, v ktorom interagujú s človekom a vplývajú na spoločnosť, nemalo byť umožnené.

Uvedený sumár problémov technológií umelej inteligencie s priamym či nepriamym dopadom na človeka a spoločnosť nás núti **vytvoriť si všeobecné zásady prístupu k fenoménu umelej inteligencie:**

- nutnou podmienkou vývoja a reálneho nasadenia akéhokoľvek systému AI musí byť **schopnosť a možnosť človeka prebrať kedykoľvek kontrolu nad týmto systémom**, resp. právo a možnosť verifikovať a prehodnotiť výsledky jeho činnosti.
- dizajn a funkcionálnosť každého systému AI, ktorý má byť reálne nasadený a využívaný, musí byť dostatočne známa, pochopiteľná a verifikovateľná. Jednoducho **musíme vedieť, ako daný systém AI funguje.**
- musíme byť schopní **analyzovať a chápať dôsledky nasadenia konkrétnej technológie AI.**
- **potrebujeme mať jasný manažment rizík**, ktorý zahŕňa riešenie zlyhania a limitných situácií činnosti algoritmov AI.
- musíme **nastaviť primerané požiadavky a podmienky (regulácie) pre dizajn, fungovanie a využívanie systémov AI, ktoré vychádzajú z etických princípov, právnych noriem a morálnych hodnôt.**

385 Kapitola 2.8.

- nasadenie akéhokoľvek systému AI, ktorý má potenciál ovplyvniť fungovanie, bezpečnosť i budúcnosť celej spoločnosti, **musí byť pod verejným dohľadom a kontrolou.**
- akákoľvek technológia umelej inteligencie **musí byť dôveryhodná, spĺňajúc požadované požiadavky legislatívne, etické i bezpečnostné a zameraná na človeka.**

V prehľade etických výziev, prameniach z limitov a rizík súčasných systémov umelej inteligencie, ktoré sme uviedli v 2. kapitole i diskutabilných spôsobov a oblastí využitia, prípadne zneužitia, tak nachádzame **tri oblasti nutnej implementácie etických noriem, eticko-právnych regulácií a morálnych zásad:**

- etické normy, zákonné regulácie a morálne zásady **tvorcov** systémov AI
- etické normy, zákonné regulácie a morálne zásady **poskytovateľov a používateľov** týchto systémov
- implementované eticko-právne požiadavky a obmedzenia priamo **v systémoch AI**

Uvedené úvahy primárne vzťahujeme na systémy ANI³⁸⁶, u ktorých z podstaty veci musíme rátať so zaangažovaním ľudského faktora vo všetkých oblastiach tvorby, používania i realizácie prostriedkov umelej inteligencie. Preto **aplikovanie etických princípov a regulácií v ANI – akokoľvek náročným až neuskutočniteľným by sa to v praxi mohlo ukázať – má vďaka prítomnosti ľudského faktora jasné zásady, vyhranené oblasti a viac menej definované kritériá.**

Takmer všetkými odborníkmi v oblasti kybernetickej bezpečnosti je za najslabší článok bezpečnostného reťazca považovaný človek, ktorý s informačnými systémami interaguje. V oblasti aplikácie eticko-právnych regulácií v systémoch umelej inteligencie sme si osobitne v rámci 2. kapitoly ukázali, že pri týchto technológiách to tak nemusí byť. Keďže nie je prakticky možné pokryť všetky etické výzvy spojené so systémami AI pomocou technologických riešení – **treba investovať nemalé úsilie do vzdelania, osvety i prevencie a formovať morálne postoje vývojárov, poskytovateľov i používateľov**

386 Len pripomeňme, že hovoríme o úzko špecializovaných systémoch umelej inteligencie (narrow AI), ktoré sú optimalizované na zvládnutie konkrétnej úlohy, resp. množiny úloh. Ide súčasne o systémy slabšej umelej inteligencie (weak AI), ktoré vykazujú inteligentné správanie na základe modelov a aplikovaných metód i tréningových dát. Sú to teda systémy zamerané na riešenie konkrétnych úloh a sú závislé na ľudskom vstupe a konfigurácii.

týchto technológií.

Edukácia a osвета sa tak stávajú nutnou podmienkou pre rast spoločenskej citlivosti a zodpovednosti v oblasti umelej inteligencie, a to vo viacerých oblastiach:

- spoločenská osвета a vzdelávanie za účelom rastu všeobecného povedomia o limitoch, rizikách a reálnych možnostiach technológií umelej inteligencie.³⁸⁷
- trh s umelou inteligenciou v súčasnosti zažíva extrémny boom a vyžaduje veľa špecialistov v tejto oblasti. Utešene nám rastú počty robotníkov AI (riešia nasadenie a prevádzku bežných systémov AI)³⁸⁸, no nie tak počty odborníkov na AI, ktorí „vidia do hĺbky“ týchto systémov a dokážu riešiť nielen technologické, ale aj eticko-právne výzvy pre zavádzanie dôveryhodných technológií umelej inteligencie.
- tímy odborníkov na AI by mali mať interdisciplinárny rozmer, bez ktorého ťažko uchopiť problematiku dôveryhodnosti umelej inteligencie a jej zamerania na človeka.

Smerovanie k AGI³⁸⁹, teda k všeobecnej a silnej umelej inteligencii, nás však aj v oblasti etiky stavia pred nové výzvy a takmer určite ešte neznáme možnosti, riziká i obmedzenia.

Ponajprv nie je jasné, kde sa nachádza, resp. bude nachádzať hranica medzi ANI a AGI. Tiež nevieme, do akej miery bude táto hranica jasná a presne definovaná. Je možné, že

387 Jednou z výborných aktivít, ktorá sa rozšírila prakticky v celej Európskej únii, je vytvorenie kurzu *Elements of AI* z dielne Helsinskej univerzity a organizácie Reaktor Education. Ide o kurz základov umelej inteligencie, ktorý prináša elementárny vzhľad do problematiky AI a poukazuje na viaceré etické výzvy, ktoré sú s umelou inteligenciou spojené. Doteraz kurz absolvovalo cez 750 tisíc občanov EÚ. *Elements of AI*. [on-line]. [cit. 30. marca 2022].
Dostupné na internete: <<https://www.elementsofai.com/>>

388 V zásade sú nahraditeľní prostredníctvom budúcich sofistikovanejších systémov AI:-)
Problematiku robotníkov vs. odborníkov AI diskutoval prof. Peter Sinčák z Technickej univerzity v Košiciach na konferencii ITAPA: Umelá inteligencia v bežnom živote.
Umelá inteligencia v bežnom živote. [on-line]. Bratislava: ITAPA, 2021. [cit. 24. júna 2021].
Dostupné na internete: <<https://www.itapa.sk/12826-sk/program/>>

389 Skutočná umelá inteligencia, ktorá je všeobecná (general) a silná (strong). Všeobecná, keďže dokáže zvládnuť akúkoľvek intelektuálnu úlohu a má schopnosť generalizovať, t. j. zovšeobecňovať a prenášať, či adaptovať naučené schopnosti na iné úlohy. Silná, pretože aj skutočne rozumie tomu, či rieši a vykonáva.

nás prekvapí **existencia čiastočne uvedomelých systémov AI**, ktorá si bude vyžadovať hľadanie svojich vlastných riešení etiky a regulácie.

Ak prezumujeme úspešný vývoj skutočne uvedomelej umelej inteligencie a na základe súčasných skúseností z vývoja systémov hlbokého učenia predpokladáme, že táto uvedomelá inteligencia bude inšpirovaná činnosťou ľudského mozgu, otvára sa nám úplne nová perspektíva spôsobu riešenia etických výziev a regulácií smerujúca viac k psychológii, kognitívnym vedám a právu než k informatike, kybernetike a dátovej vede.

V prípade uvedomelej umelej inteligencie inšpirovanej ľudským mozgom sa v oblasti etiky a regulácie zmysluplne nepohneme ďalej, pokiaľ nebudeme mať vyriešený jej model správania sa na spôsob teórie mysle a schopnosti zdravého rozumu u človeka.³⁹⁰

V mnohých diskusiách o možnostiach a schopnostiach systémov umelej inteligencie dodržiavať eticko-právne regulácie narážame na nepochopenie inakosti „inteligencie“ týchto technológií. Uvedomujúc si bytostné rozdiely medzi človekom a systémami AI sa nemôžeme stavať k týmto systémom ako k niečomu ľudsky inteligentnému a človeku podobnému.

Jednu z implikácií, ktoré sa od tohoto omylu odvíjajú, sme už spomenuli v diskusiách o reguláciách LAWS – išlo o tzv. „následnú reguláciu“³⁹¹, ktorá nevychádza zo všeobecne akceptovaných mravných hodnôt a etických princípov, ale sa prispôsobuje vývoju a možnostiam nasadenia týchto autonómnych zbraňových systémov. V podstate toto prispôsobovanie a vývoj v oblasti etiky, hodnotového rámca a morálnej obhájitelnosti predpokladá u systémov AI podobný model správania, ako je ten ľudský a tým aj možnosť nahradiť ľudský morálny úsudok a svedomie algoritmickým spracovaním. Znovu pripomíname – pokiaľ nemáme vyriešený model správania sa systémov AI na spôsob teórie mysle a schopnosti zdravého rozumu u človeka – **ide o veľmi nebezpečný argumentačný faul.**

Podobným argumentačným zlyhaním je i porovnávanie presnosti činnosti bojových systémov AI (osobitne budúcich AGI) s konaním ľudských osôb nasadených v porovnateľných bojových situáciách. Vraj na základe „vhodného naprogramovania“

390 Modely správania na spôsob teórie mysle a schopnosti zdravého rozumu u človeka sme aspoň okrajovo diskutovali v kapitole 2.1.2.

391 „Následnú reguláciu“ sme rozoberali na konci kapitoly 2.7.8.

(po kapitolách 1 a 2 si aspoň približne vieme predstaviť, resp. si ani nedokážeme predstaviť, čo všetko by také niečo u AGI obnášalo) budú LAWs schopné lepšie dodržiavať humanitárne štandardy než ľudia. Určite môžeme súhlasiť, že špičkové LAWs dokážu presnejšie strieľať, rýchlejšie identifikovať jednotlivé objekty scény a možno i bleskurýchle vyhodnocovať operačnú situáciu, atď. – a to všetko bez únavy či iných podmienok zabezpečujúcich životné funkcie ľudského personálu. Takže systémy AI sa budú vedieť niektorých chýb lepšie vyvarovať než ľudia. Na tejto istej – povedzme „technologickej“ úrovni – však môžeme poukázať aj na zlyhania prameniace z rizík a problémov uvedených v kapitolách 2.1. až 2.4., takže z iného uhla pohľadu by mohli mať navrch ľudia.³⁹² Problematická je však iná, hlbšia rovina. V kontexte modelu ľudského správania poznajúc ľudské zlyhania, ktoré môžu viesť až k vojnovým zverstvám, sa snažíme nastaviť záväzné pravidlá a dohovory i trestno-právne postihy na obmedzenie týchto zlyhaní.³⁹³ **Pokiaľ však nemáme vyriešený model správania sa AGI, jej základné myšlienkové rámce korešpondujúce u človeka so zdravým rozumom, nemáme na základe čoho a ako nastavovať pravidlá, ktoré by boli vymožitelné a aj reálne dodržiavané.**

V niektorých oblastiach sa môžeme dostať až k principiálnej neistote v oblasti etických záruk. Napríklad, ak znovu uvažujeme o súčasných autonómnych zbraňových systémoch so smrtiacim účinkom (LAWs) ako o systémoch ANI, správne môžeme argumentovať, že pokiaľ tieto systémy nedokážu chápať bojovú scénu a na ňu naviazané široké súvislosti (napríklad rozpoloženie identifikovaných ľudských subjektov, výskyt civilistov na bojisku, použitie civilistov ako štít a pod.), nemôžu bez riadiaceho ľudského činiteľa konať. Ak však by sme mali LAWs vybavený algoritmi AGI, pokiaľ si nebudeme istí ich modelom správania porovnateľným s ľudskou schopnosťou zdravého rozumu a tým pádom aj zodpovedajúco formovateľným, o to skôr musíme povedať, že takéto systémy nemôžu samostatne konať!

Ako sme uviedli v kapitole 2.1.2., človek – napriek tomu, že neustále robí chyby – dokáže konať skutočne inteligentne. Jednoducho má to, čo chýba všetkým súčasným systémom umelej inteligencie, a to zdravý rozum (ako súčasť všeobecnej inteligencie), ktorého súčasťou je schopnosť abstrahovať a na základe analógií a konceptov nachádzať riešenia,

392 Samozrejme tým nerozporujeme vojenské benefity, ktoré sme uviedli v kapitole 2.7.

393 Napríklad Ženevské konvencie, ktoré upravujú podmienky a pravidlá medzinárodného práva na ochranu obetí vojny alebo Haagske konvencie a Ženevský protokol, ktoré upravujú použitie zbraní vo vojne.

myslieť a riešiť na základe nesmierneho rozsahu najrozmanitejších vedomostí, poznatkov a skúseností. Ľudská bytosť využíva zdravý rozum podvedome a spontánne v akejkoľvek oblasti života. Preto pre mnohých je dôveryhodný systém umelej inteligencie, ktorý môže úspešne fungovať v komplexnom reálnom svete, podmienený schopnosťou mať zdravý rozum tak, ako má človek.

Tieto úvahy nám môžu poslúžiť aj **v diskusii o schopnosti implementované eticko-právne požiadavky a obmedzenia priamo v systémoch AI.**

Vieme si predstaviť parciálnu implementáciu zameranú napríklad na odolnosť voči predsudkom a poverám algoritmov ANI, no nie komplexné zvládnutie etických princípov a právnych noriem. I v oblasti etiky sa tak pre nás stáva lákadlom uvedomelá AGI, schopná skutočne konať inteligentne, pre ktorú by komplexnosť etických princípov nemusela byť problémom. Ako sme však už uviedli, pokiaľ nebudeme mať vyriešený model správania so schopnosťou prejavovať (a zachovávať) zdravý rozum, s plnou implementáciou eticko-právnych požiadaviek a obmedzení nemôžeme rátať.

V komplexnom pohľade na človeka si môžeme uvedomiť viaceré rozmery mantinelov a obmedzení, ktoré sa podpisujú pod naše schopnosti určitým spôsobom konať:

1. ako biologické systémy staviame na tom, čo do DNA vložila evolúcia a preverila príroda, pričom presah do kognitívnych schopností a dopad na psychológiu človeka netreba v dnešnej dobe pripomínať.
2. využívame zdravý rozum (ako súčasť všeobecnej inteligencie), ktorého súčasťou je schopnosť abstrahovať a na základe analógií a konceptov nachádzať riešenia, myslieť a riešiť využívajúc bázu nesmierneho rozsahu najrozmanitejších vedomostí, poznatkov a skúseností.
3. máme hodnotový systém a využívame sociologické i kultúrne rámce vnímania sveta a konania.
4. máme svedomie a duchovný rozmer nášho života.

Uvažujúc o umelej inteligencii, k 1. bodu by sme vedeli nájsť aspoň hrubú analógiu v tvorbe dnešných systémov AI.

Vieme vytvoriť – a aj sme vytvorili – komplexné systémy ANI, ktoré podľa 2. bodu dokážu riešiť svoje úlohy využívajúc bázu nesmierneho rozsahu najrozmanitejších vedomostí a poznatkov. Nevieť však realizovať AGI so schopnosťou plnej abstrakcie za skutočného

chápania vzťahov či súvislostí a zdravého rozumu.

Dokážeme prevádzkovať technológie ANI, ktoré vedia spracúvať naše sociologické a kultúrne rámce. Sme však na rozpakoch, ako realizovať AGI s implementovaným (nebodaj aj chápaným) hodnotovým systémom.

Svedomie, ako referenčný bod pre etický systém a morálne rozhodovanie, už z princípu nie je možné v ANI „implementovať“. Keďže svedomie nedokážeme oddeliť od osoby so schopnosťou rozumu, slobodnej vôle a sebauvedomenia, v prípade AGI sa potýkame s dilemou umelej inteligencie, u ktorej si nevieme predstaviť ontologickú analógiu ľudského bytia, bez ktorej niet ani svedomia.³⁹⁴

Vráťme sa však späť do kráľovstva dnešných technológií umelej inteligencie, teda k úzko špecializovaným systémom slabej umelej inteligencie (ANI) a v nasledovných kapitolách sa aspoň v krátkosti pozrime na to, ako sa súčasný svet snaží uchopiť jej etické problémy a výzvy.

3.2. Angažovanosť a aktivity na poli etiky umelej inteligencie

Jednou z odlišností, ktorými sa stať o systémoch AI vo vojenskej sfére líšila od ostatných častí 2. kapitoly popisujúcich limity a riziká týchto technológií, bolo viacnásobné zdôraznenie angažovanosti a aktivít v oblasti etiky a regulácie nasadenia autonómnych zbraňových systémov.

Dôvod je jasný – prakticky žiadna iná aplikácia systémov umelej inteligencie nepriniesla toľko polemík a rizík ako práve vojenské využitie. Obavy z napredovania vo vývoji algoritmov umelej inteligencie a spôsobu ich zavádzania do moderných zbraňových systémov tak dali v poslednej dekáde vznik viacerým iniciatívam, či už odborným a celospoločenským alebo armádnym.³⁹⁵

Na chvíľku prekročiac rámec posledného desaťročia, pokladáme za vhodné spomenúť jednu zo základných výziev Russell-Einsteinovho manifestu, reagujúceho na masívny rozmach vývoja a zavádzania nukleárných zbraní do výzbroje veľmocí studenou vojnou rozdeleného sveta:³⁹⁶

394 A ak by sme si to i vedeli predstaviť, ako by v takomto prípade vyzeral teonómny rozmer svedomia, ktoré má referenčný bod a autoritu mimo seba v Bohu?

395 Niektoré sme už v kapitole 2.7.8. spomínali, no dovoľíme si ich predsa len zopakovať.

396 Russell-Einsteinov manifest, vydaný v roku 1955, bol verejným vyhlásením významných vedcov v čase

„Mnohé varovania vyslovili významní muži vedy a autority v oblasti vojenskej stratégie. Nikto z nich nepovie, že najhoršie výsledky sú isté. Hovoria však, že tieto výsledky sú možné a nikto si nemôže byť istý, že sa neuskutočnia. Zatiaľ sme nezistili, že by názory odborníkov na túto otázku v nejakej miere záviseli od ich politického presvedčenia alebo predsudkov. Pokiaľ to naše výskumy odhalili, závisia len od rozsahu vedomostí konkrétného odborníka. Zistili sme, že ľudia, ktorí toho vedia najviac, sú najpochmúrnejší.“

V tomto kontexte je skutočne nadčasová najznámejšia a pritom tak strohá výzva manifestu: „**Pamätajte na svoju ľudskosť a zabudnite na zvyšok**“.

Už v roku 2013 skupina inžinierov, odborníkov na umelú inteligenciu a robotiku a ďalších vedcov a výskumníkov z tridsiatich siedmich krajín sveta vydala **Scientists' Call to Ban Autonomous Lethal Robots**. V tejto výzve sa uvádza, že chýbajú vedecké dôkazy o tom, že by roboty mohli v budúcnosti disponovať „funkciami potrebnými na presnú identifikáciu cieľa, situačné povedomie alebo rozhodnutia týkajúce sa primeraného použitia sily“. LAWS tak môžu spôsobiť vysokú mieru vedľajších škôd. Vyhlásenie sa končí zdôraznením, že „rozhodnutia o použití hrubej sily sa nesmú delegovať na stroje“.³⁹⁷

V júli roku 2015 bol na jednej z medzinárodných konferencií o umelej inteligencii zverejnený otvorený list **Autonomous Weapons: An Open Letter from AI & Robotics Researchers**, ktorý vyzýval na zákaz autonómnych zbraní. V liste sa priamo píše: „technológia umelej inteligencie dosiahla bod, v ktorom je nasadenie týchto systémov prakticky, ak nie legálne, uskutočniteľné v priebehu nie desaťročí, ale rokov a v stávke je veľa: autonómne zbrane sa uvádzajú ako tretia revolúcia v zbraňových systémoch,

zvýšeného medzinárodného napätia medzi Východom a Západom, ako aj rastúceho zasahovania do vnútorných záležitostí štátov prostredníctvom vojenských i tajných operácií. Bolo to tiež obdobie, ktoré charakterizoval prevratný vedecký pokrok a technologické inovácie.

Manifest mal pôvod v obavách fyzika Josepha Rotblata a polyhistora Bertranda Russella z účinkov jadrových zbraní, rizika šírenia jadrových zbraní a budúcnosti ľudstva. To znamená, že autori chceli biť na poplach v súvislosti s potenciálne katastrofickými dôsledkami použitia jadrových zbraní a existenčným rizikom, ktoré pre ľudstvo predstavuje ich zachovanie. Medzi jedenástimi signatármi manifestu bolo niekoľko najvýznamnejších svetových vedcov. Desať z nich bolo nositeľmi Nobelovej ceny za prínos v oblasti fyziky, chémie, fyziológie a medicíny, literatúry alebo mieru.

VIGNARD, K. *Manifestos and open letters: Back to the future?* [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://thebulletin.org/2018/04/manifestos-and-open-letters-back-to-the-future/>>

397 *Scientists' Call to Ban Autonomous Lethal Robots*. ICRC, October 2013. [on-line]. [cit. 8. marca 2022].

Dostupné na internete: <<http://www.icrac.net/>>

po strelnom prachu a jadrových zbraniach“. Tento otvorený list poukazuje na pokrok i výhody spojené s vývojom systémov AI, zvažuje riziká a dopady zneužitia LAWS na celú oblasť umelej inteligencie a graduje k výzve na „zákaz útočných autonómnych zbraní, ktoré sú mimo zmysluplnej ľudskej kontroly“.³⁹⁸

List mal do roku 2016 pôsobivý zoznam signatárov, medzi ktorými sú priekopníci z oblasti vedy a techniky³⁹⁹, takmer päťtisíc výskumníkov v oblasti umelej inteligencie a robotiky i skoro dvadsaťsedem tisíc ďalších podporovateľov.

Nasledovali i cielenejšie aktivity, napríklad otvorený list výskumníkov a zakladateľov spoločností pôsobiacich v oblasti umelej inteligencie **An Open Letter to the United Nations: Convention on Certain Conventional Weapons**, v ktorom v roku 2017 vyzývali všetkých participantov v tejto oblasti vývoja, aby „zabránili pretekom v zbrojení týmito zbraňami, aby chránili civilistov pred ich zneužitím a aby zabránili destabilizujúcim účinkom týchto technológií“.⁴⁰⁰

Otvorený list nenechal nikoho na pochybách, že signatári majú vážne obavy z uplatňovania technológií umelej inteligencie a robotiky v budúcich zbraňových systémoch: „Nemáme veľa času na to, aby sme konali. Keď sa táto Pandorina skrinka otvorí, bude ťažké ju zavrieť.“⁴⁰¹

Tieto aktivity nezostali bez odozvy – v nasledujúcich rokoch viacero organizácií podieľajúcich sa na vývoji algoritmov umelej inteligencie aplikovateľných vo vojenských systémoch, resp. v akýchkoľvek systémoch AI zneužiteľných v armáde a represívnych zložkách, odrieklo svoju účasť na projektoch z týchto oblastí.⁴⁰² Paradoxom však je, že

398 *Autonomous Weapons: An Open Letter from AI [Artificial Intelligence] & Robotics Researchers*. [on-line]. Future of Life Institute, 2015. [cit. 8. marca 2022].

Dostupné na internete: <<http://futureoflife.org/open-letter-autonomous-weapons/>>

399 Okrem iných napríklad Elon Musk (vynálezca a zakladateľ spoločnosti Tesla), Steve Wozniak (spoluzakladateľ spoločnosti Apple), fyzik Stephen Hawking (Univerzita v Cambridge), Noam Chomsky (Massachusettský technologický inštitút), Stuart Russell (Berkeley) atď.

400 *An Open Letter to the United Nations: Convention on Certain Conventional Weapons*. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://futureoflife.org/2017/08/20/autonomous-weapons-open-letter-2017/>>

401 VIGNARD, *Manifestos and open letters: Back to the future?* [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://thebulletin.org/2018/04/manifestos-and-open-letters-back-to-the-future/>>

402 SAMPLE, I. *Thousands of leading AI researchers sign pledge against killer robots*. [on-line]. [cit. 21.

viaceré z týchto organizácií odmietlo obmedziť svoje technológie umelej inteligencie v tak kontroverznej oblasti ako je kapitalizmus dohľadu (surveillance capitalism), ktorý sme diskutovali v kapitole 2.5. a ktorého dôsledky sú v niektorých aspektoch pre spoločnosť takmer tak rizikové ako autonómne zbraňové systémy.⁴⁰³

Uvádzajúc aktivity v oblasti etiky autonómnych zbraňových systémov treba pripomenúť aj armádne aktivity, ktorých existencia, rozsah a praktický dopad majú svoju výpovednú hodnotu a dotvárajú hodnotový rámec tej-ktorej armády. Či už ide o univerzitné armádne aktivity v USA, z prostredia ktorých sme v kapitolách 2.7.7. a 2.7.8. citovali z článku *Pros and Cons of Autonomous Weapons Systems*, alebo priamo o eticko-právny diskurz prebiehajúci v armádnych kruhoch, napríklad v kapitole 2.7.8. komentovaný dokument **AI Principles: Recommendations on the Ethical Use of Artificial Intelligence by the Department of Defense** z Defense Innovation Board USA.

Ťahúňom armádnych aktivít na poli regulácií a eticko-právnych dôsledkov nasadenia LAWs sú primárne armády USA a NATO, nakoľko odzrkadľujú nielen spoločenské povedomie v tejto oblasti, ale predovšetkým technologickú vyspelosť armádnych systémov AI. Podľa viacerých autorov im technologicky nielen zdatne sekunduje, ale ich možno aj predbieha rozvoj čínskych zbraňových systémov a všeobecne technológií AI,⁴⁰⁴ čo však – vzhľadom na súčasný postoj čínskych úradov k etike systémov AI – môže byť problémom pre celosvetové presadzovanie potrebných regulácií v tejto oblasti. A vzhľadom na zavádzanie technológií AI aj do iných armád disponujúcich zbraňami hromadného ničenia, ktoré môžu mať skutočne rôznorodý vzťah k eticko-právnym aspektom AI, musíme brať vážne výzvy na seriózne riešenie tejto problematiky v snahe vyvarovať sa

marca 2022].

Dostupné na internete: <<https://www.theguardian.com/science/2018/jul/18/thousands-of-scientists-pledge-not-to-help-build-killer-ai-robots>>

403 ZUBOFF, SH. *The real reason why Facebook and Google won't change*. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.theguardian.com/science/2018/jul/18/thousands-of-scientists-pledge-not-to-help-build-killer-ai-robots>>

404 *The AI arms race*. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.ft.com/content/21eb5996-89a3-11e8-bf9e-8771d5404543>>

'Wake up': Ex-Air Force software officer warns China is winning AI battle. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.washingtontimes.com/news/2021/oct/11/nicolas-chaillan-warns-china-winning-ai-battle/>>

katastrofe nukleárnej vojny.⁴⁰⁵

Angažovanosť výskumníkov a odbornej verejnosti v oblasti regulácie zbraňových systémov umelej inteligencie našla odozvu aj v diplomatických diskusiách OSN v Ženeve, ktoré sa vedú v rámci Dohovoru o určitých konvenčných zbraniach z roku 1980. Konzultácie, ktoré sa započali v roku 2014 ako séria neformálnych stretnutí, v roku 2017 vykryštalizovali do stretnutí skupiny vládnych expertov. I napriek tomu, že členovia vedeckej komunity a experti sú opatrní a hľadajú riešenia potenciálnych nebezpečenstiev autonómnych zbraní, určite nie sú pesimistickí, pokiaľ ide o potenciálne výhody umelej inteligencie alebo technologických inovácií vo všeobecnosti. Sú preto aktívne zaangažovaní v úsilí, aby umelá inteligencia bola prospešná pre spoločnosť a v konečnom dôsledku aj pre ľudstvo.⁴⁰⁶

Jedným z plodov týchto aktivít v rámci Európskej únie bolo prijatie **Uznesenia o autonómnych zbraňových systémoch**⁴⁰⁷, ktoré sa v roku 2018 uskutočnilo na pôde Parlamentu EÚ a ktoré si kladie za cieľ globálny zákaz smrtiacich autonómnych zbraňových systémov. Myslím, že nás neprekvapí odmietavý postoj USA a Ruska (ale i ďalších dôležitých hráčov v oblasti LAWs) k tomuto zákazu. Ako sme uviedli – vzhľadom na potenciál LAWs pracujú skôr na reguláciách a „korektnom“ spôsobe ich nasadenia.

Vzhľadom na smrtiace účinky a rozsah rizík je pochopiteľné, že práve táto oblasť je v strede záujmu, no nie je jediná.

V kapitolách 2.5. a 3.1. sme spomenuli aj problematiku autonómnych vozidiel, v rámci ktorej sa rieši nielen schopnosť autonómnych systémov správne sa rozhodovať v kritických situáciách, ale i legislatívne požiadavky na tzv. bezpečnostného operátora (safety operator), hľadajúc objektívnu mieru zodpovednosti ľudského faktora pri riadení autonómneho vozidla v rámci jednotlivých stupňov automatizácie, ktoré sme uviedli v kapitole 2.5.

405 GEIST, E., LOHN, A. J. *How Might Artificial Intelligence Affect the Risk of Nuclear War?* [on-line]. Santa Monica, CA: RAND Corporation, 2018. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.rand.org/pubs/perspectives/PE296.html>>

406 VIGNARD, *Manifestos and open letters: Back to the future?* [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://thebulletin.org/2018/04/manifestos-and-open-letters-back-to-the-future/>>

407 *Resolution on autonomous weapon systems*. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <[https://oeil.secure.europarl.europa.eu/oeil/popups/ficheprocedure.do?reference=2018/2752\(RSP\)&l=en](https://oeil.secure.europarl.europa.eu/oeil/popups/ficheprocedure.do?reference=2018/2752(RSP)&l=en)>

Asi najďalej sú v tomto smere diskusie v USA, ktoré sa premietli do viacerých legislatívnych ustanovení, či už na úrovni samostatných štátov (napr. Kalifornia, Texas) alebo v rámci federálnych zákonov. V súčasnosti (marec 2022) napríklad Spojené štáty menia svoju legislatívu, z ktorej pre autonómne vozidlá 5. úrovne autonómie vypadla povinnosť zahrnúť do ich vybavenia základné ovládacie prvky.⁴⁰⁸ **I v tejto oblasti sme viackrát deklarovali požiadavku na schopnosť človeka kedykoľvek prebrať kontrolu nad systémom. Ako vidíme, čím sofistikovanejší systém umelej inteligencie v autonómnych vozidlách sa použije, tým viac sa naplnenie tejto požiadavky rozplýva. Ide o nebezpečný trend, ktorý by sa mohol rozšíriť i na iné oblasti nasadenia pokročilých systémov AI.**

Problematika autonómnych vozidiel je na eticko-právne výzvy veľmi bohatá – či už ide o spoľahlivú funkcionálnosť systémov AI za každých poveternostných podmienok a možných situácií v bežnej cestnej prevádzke, alebo o správne rozhodovanie sa v kritických situáciách⁴⁰⁹ a v neposlednom rade aj o ochranu osobných údajov a spoločenské dopady, keďže prevádzka týchto vozidiel je extrémne závislá na širokospektrálnom zbere najrozličnejších dát a ich spracovaní v reálnom čase algoritmi AI, resp. i následnom analytickom vyťažení v dátových centrách poháňaných technológiami umelej inteligencie.

V oblasti dohľadových systémov a hlavne komplexnej problematiky algoritmického riadenia míľovými krokmi kráča v ústrety umelej inteligencii Čína.⁴¹⁰ Je však len na škodu, že tento prekotný rozvoj nie je sprevádzaný adekvátnou verejnou diskusiou a nepodlieha odbornému i spoločenskému diskurzu, ktorý by bol garantom, resp. katalyzátorom eticko-právnych požiadaviek sprevádzajúcich tento pod taktovkou vládnej ideológie prebiehajúci

408 SHEPARDSON, D. *U.S. eliminates human controls requirement for fully automated vehicles*. [on-line]. [cit. 22. marca 2022].

Dostupné na internete: <<https://www.reuters.com/business/autos-transportation/us-eliminates-human-controls-requirement-fully-automated-vehicles-2022-03-11/>>

409 Už sme spomínali príklad kritickej situácie, v ktorom sa musí autopilot rozhodnúť medzi potencionálnym ohrozením života posádky alebo ostatných účastníkov v čase blížiacej sa alebo prebiehajúcej dopravnej nehody.

410 Pomerne ucelený pohľad na problematiku zavádzania algokracie v Číne ponúka súbor štúdií britskej agentúry pre sociálny rozvoj Nesta.

The AI Powered State. [on-line]. [cit. 26. marca 2022].

Dostupné na internete: <<https://www.nesta.org.uk/feature/ai-powered-state/>>

technologický šprint.⁴¹¹

Vo všeobecnosti môžeme povedať, že povedomie o etických požiadavkách na systémy umelej inteligencie neustále rastie. Riziká technológií AI riešia nielen expertné skupiny OSN, EÚ, či jednotlivých štátov, ale varujú pred nimi i viaceré osobnosti známe z vedeckého a technologického sveta (Hawking, Musk,...). A prakticky pri všetkých významnejších centrách výskumu umelej inteligencie sa etablujú oddelenia zaoberajúce sa etickými otázkami a výzvami.

Mnohé z iniciatív, ktoré prebiehajú vo vládnych a vedeckých pracovných skupinách i technologických komunitách, sa skutočne snažia formulovať potenciálne prínosy umelej inteligencie a súčasne obmedzovať jej riziká. Medzi významné výsledky týchto snáh môžeme zaradiť vedecký program formulovaný v **Otvorenom liste o výskumných prioritách pre robustnú a prospešnú umelú inteligenciu**⁴¹² z roku 2015 a **Globálnu iniciatívu IEEE o etike autonómnych a inteligentných systémov**⁴¹³. Medzi ďalšie iniciatívy na multilaterálnej úrovni patrí napríklad AI for Good Global Summit 2017⁴¹⁴, ako aj práca organizácií ako OpenAI⁴¹⁵ alebo Partnership on AI⁴¹⁶, na ktoré nadväzuje mnoho ďalších aktivít po celom svete.

411 Hodnotíme tak podľa našich kritérií euroatlantického spoločenského priestoru. Aktivitu na podporu implementácie etických pravidiel samozrejme prebiehajú aj v Číne, napr.:

Understanding China's AI Strategy. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.cnas.org/publications/reports/understanding-chinas-ai-strategy>>

China wants to make its own rules for AI ethics. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.scmp.com/abacus/tech/article/3029194/china-wants-make-its-own-rules-ai-ethics>>

412 *Open Letter on Research Priorities for Robust and Beneficial Artificial Intelligence*. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://futureoflife.org/2015/10/27/ai-open-letter/>>

413 *IEEE Global Initiative on Ethics of Autonomous and Intelligent Systems*. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <http://standards.ieee.org/develop/indconn/ec/ead_v2.pdf>

414 *AI for Good Global Summit 2017*. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.itu.int/en/ITU-T/AI/Pages/201706-default.aspx>>

415 *OpenAI*. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://openai.com/>>

416 *Partnership on AI*. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://partnershiponai.org/>>

3.3. Legislatívne kroky v oblasti etiky umelej inteligencie

Až donedávna v žiadnej z krajín sveta neexistovala ucelená legislatíva pokrývajúca celú problematiku umelej inteligencie. Doterajšie regulácie dotýkajúce sa technológií AI boli buď veľmi špecifické, riešiac parciálne problémy konkrétnych oblastí nasadenia týchto systémov a / alebo boli súčasťou iných regulácií, napr. kybernetickej bezpečnosti, ochrany osobných údajov, sektorových regulácií v rámci finančného sektora alebo štátnej správy.

V rámci Európskej únie však vzniká úplne nová regulácia, ktorý nemá obdobu nikde vo svete.⁴¹⁷ Stanovujúc si za cieľ pokryť celú problematiku súčasných systémov umelej inteligencie, ide vôbec o prvú komplexnú reguláciu konkrétnej technológie.⁴¹⁸

Primárnym dôvodom pre vznik a aplikovanie tejto regulácie je dopad fenoménu umelej inteligencie na človeka a spoločnosť. Jednoducho povedané, umelá inteligencia je spôsobilá zmeniť postavenie človeka.

Ľudská bytosť so svojim zmyslovým vnímaním, intelektuálnym vybavením a schopnosťami i komplexnou psychológiou nie je prakticky schopná technológiám tohoto „kalibru“ vôbec odolávať. Spoločnosť tak bez adekvátnych regulácií a nastavených limitov nie je schopná čeliť dôsledkom činnosti systémov a vplyvu technológií AI s potenciálom rozkladať základné princípy, na ktorých je naša spoločnosť postavená, napríklad negatívne vplývať na ľudské práva a dôstojnosť človeka, podryvať demokratické princípy a manipulovať, atakovať bezpečnostné mechanizmy a pod.⁴¹⁹

417 Vo svete v súčasnosti vzniklo viacero regulačných rámcov AI:

- EU AI HLEG Guidelines and Assessment List for Trustworthy AI (z nej pramení európsky Akt o AI)
- ICO AI Auditing Framework
- Singapore Model Governance Framework
- Dubai AI Ethics Toolkit
- Hong Kong Ethical Accountability Framework

Európsky rámec je však podľa analýzy *Artificial Intelligence: Principles, laws, and frameworks* jednoznačne najvyváženejší, najucelenejší a najkomplexnejší.

Por. *Artificial Intelligence: Principles, laws, and frameworks*. OneTrust DataGuidance Limited, 2022. ISSN 2398-9955.

418 PATTYNOVÁ, J. *Výzvy a právni aspekty umělé inteligence*. In: *Umělá inteligence 2021*. Praha: 2021.

Dostupné na internete: <<https://www.tuesday.cz/akce/umela-intelligence-1/zaznam-akce/>>

419 PATTYNOVÁ, J. *Výzvy a právni aspekty umělé inteligence*.

Dostupné na internete: <<https://www.tuesday.cz/akce/umela-intelligence-1/zaznam-akce/>>

Tento rozklad môže viesť až k disruptívnym zmenám v modernej informačnej a znalostnej spoločnosti.

Podľa Gartner *Top Strategic Predictions For 2020 And Beyond* „Technológie (primárne AI) a ich aplikácie sú pripravené ovplyvniť každý aspekt toho, čo nazývame človečenstvom“.⁴²⁰

Gartner vo svojej podrobnej správe *Top Strategic Predictions for 2020 and Beyond: Technology Changes the Human Condition*⁴²¹ odhaduje, že v roku 2023 bude 40% digitálne sledovaného správania realizovaných algoritmi AI a v roku 2050 bude viac než 50% reklám cielených na základe detekcie emócií technológiami umelej inteligencie. Ide o skutočnosti, na ktoré sa ľudská spoločnosť nie je schopná v tak krátkej dobe adaptovať.

Zaujímavou je preto predikcia regulácie, ktorá stavala na existujúcich legislatívnych aktivitách technologicky pokročilých štátov: do roku 2023 štyri z krajín G7 nastaví pravidlá dohľadu nad vývojármi AI a nasadením metód strojového učenia.

Tieto skutočnosti – aspoň rámcovo vysvetľujúce primárny zámer pre vznik a aplikovanie regulácie umelej inteligencie v jej dopade na človeka a spoločnosť – patria ku konkrétnym dôvodom, pre ktoré európski zákonodarcovia prišli s komplexnou reguláciou, o ktorej dúfajú, že podobne ako pri GDPR pôjde o legislatívu ašpirujúcu na globálny štandard, resp. reguláciu zásadným spôsobom ovplyvňujúcu zákonodarstvo aj ostatných krajín.

Univerzálny záber a potenciálny celosvetový dosah v súčasnosti prakticky jedinej komplexnej regulácie technológií umelej inteligencie vo svete je dôvodom, pre ktorý sa v tejto kapitole primárne venujeme len tejto konkrétnej legislatíve.

Okrem Gartnerom spomenutých trendov využívania technológií AI, ktorých riziká sme širšie diskutovali v kapitole 2.5. a ktorých dôsledky ľudská spoločnosť nie je schopná v tak krátkej dobe vstrebať, resp. adaptovať sa na ne, európski zákonodarcovia reagovali i na ďalšie „patológie“ a výzvy:⁴²²

420 Gartner *Top Strategic Predictions For 2020 And Beyond*. [on-line]. [cit. 23. marca 2022].

Dostupné na internete: <<https://www.gartner.com/smarterwithgartner/gartner-top-strategic-predictions-for-2020-and-beyond>>

421 *Top Strategic Predictions for 2020 and Beyond: Technology Changes the Human Condition*. [on-line]. [cit. 24. marca 2022].

Dostupné na internete: <<https://www.gartner.com/document/3970846>>

422 PATTYNOVÁ, Výzvy a právni aspekty umělé inteligence.

Dostupné na internete: <<https://www.tuesday.cz/akce/umela-intelligence-1/zaznam-akce/>>

- dilemma medzi personalizáciou zákazníckej skúsenosti/reklamy a manipuláciou v dôsledku informačnej asymetrie. V súčasnosti sociálne siete, poskytovatelia služieb, banky a iné organizácie spracúvajú o svojich používateľoch na základe algoritmického spracovania vstupov zákazníkov extrémne veľa (nielen) osobných údajov a sú schopní ďalekosiahleho modelovania a analýzy s následkami, ktoré sme v kapitole 2.5. farbisto rozoberali. Bez regulácie by tento trend nabral obľudné rozmery.
- ochrana osôb, spoločnosti a demokracie vzhľadom na riziká deep fake a manipulácií, pri ktorých ide o takú manipuláciu reality (primárne mediálnych záznamov a prenosov, informačných tokov a spravodajstva) prostredníctvom technológií AI, že prakticky nie je možné rozoznať reálne mediálne záznamy a správy od umelých, nemajúcich s realitou nič spoločné. Osobitne obrazový materiál má schopnosť pôsobiť na emocionálnu zložku a prinášať podprahové podnety, voči ktorým – ak sú v podaní sofistikovaných algoritmov AI – môžeme byť takmer bezbranní.
- možnosť zneužitia biometrických technológií a osobitne problematika identifikácie osôb (face recognition) v reálnom čase. V digitalizovanej spoločnosti ide o veľký problém v oblasti dohľadových systémov, sledovania, ochrany súkromia a ľudských práv. Navyše s pomocou biometrických údajov sa v informačnom svete overuje digitálna identita osoby, takže ich zneužitie môže mať fatálne následky z pohľadu práva i etiky.
- riziko sofistikovaných zásahov do súkromia a osobných slobôd z vážnych dôvodov, napr. ochrana zdravia a pod. I keď by sa mohlo zdať, že ide primárne o oblasť, ktorej sme sa venovali pri dohľadových a spravodajských systémoch, čo i len parciálne smerovanie k algoritmickému riadeniu spoločnosti poukazuje na celý rad nových problémov, keďže pre úspešnosť týchto systémov AI je potrebné zasiahnuť do súkromia veľkého množstva ľudí.⁴²³

Synergický efekt doterajších parciálnych regulácií, metodického skúmania vplyvu

423 Napr. pre správne diagnostikovanie vzácnych druhov vážnych chorôb musí algoritmus AI mať prístup k v súčasnosti osobitne chránenej kategórii zdravotných osobných údajov, pre spracúvanie hypoték a úverov v reálnom čase s odolnosťou voči sofistikovaným podvodom musí byť systém AI natrénovaný na nesmiernom množstve finančných dát zákazníkov, podpora pre automatizované súdne konania musí vychádzať z výborne zvládnutých databáz realizovaných súdnych prípadov a pod.

technológií umelej inteligencie na človeka a spoločnosť i rastúceho povedomia a angažovanosti na poli etiky a práva sa na pôde inštitúcií EÚ pretavil do návrhu nariadenia⁴²⁴ o umelej inteligencii, ktorý pod skráteným názvom **Akt o umelej inteligencii** (Artificial Intelligence Act) bol zverejnený 21. apríla 2021 a aktuálne je v stave revízie a doplnení zo strany Európskeho parlamentu a Rady EÚ.⁴²⁵

Nariadenie priamo na pôde EÚ nestavalo na zelenej lúke – už v apríli 2019 prezentovala Skupina expertov AI na vysokej úrovni **Etické usmernenia pre dôveryhodnú umelú inteligenciu**. Išlo o revidované znenie z roku 2018, do ktorého bolo v rámci otvorenej konzultácie zapracovaných viac ako 500 pripomienok.⁴²⁶ Mimochodom, sumárom usmernení je všeobecná požiadavka na systémy AI, ktoré musia byť:⁴²⁷

- zákonné – rešpektujúce všetky platné zákony a predpisy.
- etické – rešpektujúce etické zásady a hodnoty.
- robustné – z technického hľadiska a zároveň zohľadňujúca svoje sociálne prostredie.

Ide o rovnaké závery ako tie, ktoré sme uviedli ku koncu kapitoly 2.8. o podmienkach pre dôveryhodnú umelú inteligenciu.

Nariadenie Akt o umelej inteligencii sa navyše inšpirovalo úspešným zavedením dôležitej

424 Na stupnici nariadenie – smernica – rozhodnutia – odporúčania/stanoviská má nariadenie najväčšiu právnu silu, keďže má prednosť pred vnútroštátnym právom, je všeobecne záväzné pre všetkých, neimplementuje sa – platí v takej forme, v akej ho prijal Európsky parlament a Rada.

425 Plný názov nariadenia znie:

Nariadenie Európskeho parlamentu a Rady, ktorým sa stanovujú harmonizované pravidlá v oblasti umelej inteligencie (Akt o umelej inteligencii) a menia niektoré legislatívne akty únie. [on-line]. [cit. 24. marca 2022].

Dostupné na internete: <<https://eur-lex.europa.eu/legal-content/SK/TXT/?uri=CELEX:52021PC0206>>

426 *Ethics guidelines for trustworthy AI.* [on-line]. [cit. 28. marca 2022].

Dostupné na internete: <<https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>>

427 V usmerneniach sa zároveň predkladá **sedem kľúčových požiadaviek, ktoré by mali systémy AI spĺňať, aby sa mohli považovať za dôveryhodné: ľudské zastúpenie a dohľad; technická odolnosť a bezpečnosť; ochrana súkromia a správa údajov; transparentnosť; rozmanitosť (diversity), nediskriminácia a spravodlivosť; spoločenský a environmentálny blahobyť; zodpovednosť.**

regulácie v oblasti ochrany osobných údajov, tzv. GDPR⁴²⁸ a prevzalo z neho základné regulačné schémy, aplikované na problematiku umelej inteligencie a spojené s aktuálnym stavom poznania v oblasti etiky technológií AI.

Akt o umelej inteligencii:

- nadväzujúc na „White Paper“ EÚ bazíruje na *human-centered* prístupe, teda **požaduje také systémy AI, ktoré sú zamerané na človeka**. Bez splnenia tejto podmienky nie je možné hovoriť o dôveryhodnej umelej inteligencii.⁴²⁹
- **zavádza vysoký štandard povinností**. Doteraz naakumulované osvedčené postupy (best practices) z oblasti implementácie systémov AI zameraných na človeka, ktoré spĺňajú eticko-právne požiadavky, prenáša do podoby zákona.
- **aplikuje rozšírenú teritoriálnu pôsobnosť**, keďže jej regulačné požiadavky sa vzťahujú na akékoľvek systémy AI, jej tvorcov a prevádzkovateľov, pokiaľ akýmkoľvek spôsobom zasahujú do života obyvateľov a štátnych celkov EÚ.
- **využíva pomerne širokú definíciu systému AI**, ktorá zahŕňa nielen strojové učenie, ale aj prístupy založené na logike, resp. poznaní (logic and knowledge based) a štatistické metódy (v zásade sa snaží pokryť celú problematiku symbolických a subsymbolických systémov AI).
- reguluje a zákonom **rieši vzťahy medzi prevádzkovateľmi a používateľmi technológií AI**. Ide tak o **systémy AI, ktoré nejakým spôsobom interagujú s okolím** (netýka sa to napr. uzavretého systému, ktorý riadi nejaký výrobný proces a v rámci výrobnéj linky pracuje len s technickými dátami materiálov a pod.).
- podľa osvedčeného vzoru GDPR **stanovuje vysoké sankcie za porušenie nariadenia**.⁴³⁰ Pri bežnom porušení sankcie siahajú až do výšky 20 miliónov Eur, resp. 4% celosvetového obratu, no ak by išlo o porušenie článku 5 (zakázané

428 General Data Protection Regulation – *Nariadenie Európskeho parlamentu a Rady (EÚ) 2016/679 z 27. apríla 2016 o ochrane fyzických osôb pri spracúvaní osobných údajov a o voľnom pohybe takýchto údajov*. [on-line]. [cit. 19. februára 2022].

Dostupné na internete: <<https://eur-lex.europa.eu/legal-content/SK/TXT/?uri=celex%3A32016R0679>>

429 Túto problematiku sme spomínali v kapitole 2.8.

430 Vysoké sankcie spolu s definovanými a realizovanými mechanizmami kontroly rešpektovania GDPR asi v najväčšej miere prispeli k rýchlej adaptácii potrebných postupov ochrany osobných údajov. To isté sa očakáva aj v oblasti umelej inteligencie.

systémy AI) a článku 10 (správa dát), sankcie sa môžu vyšplhať až do výšky 30 miliónov Eur, resp. 6% celosvetového obratu.

Uplatnenie pravidiel Aktu o umelej inteligencii môžeme očakávať v horizonte niekoľkých rokov, keďže návrh z apríla 2021 prechádza jeden až dvojročným legislatívnym procesom EÚ, po ktorom bude nasledovať legisvakančná lehota v dĺžke 24 mesiacov umožňujúca po schválení nariadenia jeho implementáciu. Táto legisvakančná lehota je však na rozdiel od iných regulácií dosť ošemetná, lebo systémy umelej inteligencie sú principiálne iné – z povahy veci nie je možné technologicky dopĺňať zabezpečenie požiadaviek nariadenia až po začiatku jeho účinnosti.⁴³¹

Vzhľadom na dizajn, parametrizáciu, trénovanie a dôsledné zabezpečenie fungujúceho systému AI je treba požiadavky nariadenia implementovať od počiatku vývoja daného systému. Zverejnený návrh Aktu o umelej inteligencii tak extrémne naberá na význame, lebo – odhliadnuc od možných úprav v rámci prebiehajúceho legislatívneho procesu – bude zaväzovať už v súčasnosti vznikajúce systémy, ich tvorcov a prevádzkovateľov, resp. poskytovateľov.⁴³²

Akt o umelej inteligencii definuje tri typy technológií AI:

- zakázané systémy
- vysoko rizikové systémy
- ostatné systémy

Medzi zakázané systémy AI patria (čl. 5):

- podprahové techniky ovplyvňujúce správanie jednotlivcov
- metódy využívajúce slabiny zraniteľných osôb (napr. deti, hendikepovaní)
- systémy sociálneho hodnotenia (social scoring) využívané štátnou správou

⁴³¹ PATTYNOVÁ, *Výzvy a právni aspekty umělé inteligence*.

Dostupné na internete: <<https://www.tuesday.cz/akce/umela-inteligence-1/zaznam-akce/>>

⁴³² Autor sa spolupodielal na tvorbe Smernice Katolíckej cirkvi v oblasti GDPR a následnej verifikácii, návrhu i úpravách viacerých informačných systémov Cirkvi, pri ktorých bolo či už technologickými alebo procesnými postupmi možné dosiahnuť splnenie legislatívnych požiadaviek ochrany osobných údajov. **Ak v prípade GDPR patrili princípy *security by design* a *privacy by design* k best practices, v prípade systémov AI je *ethics by design* a *regulation by design* nutnou podmienkou ich vývoja, nasadenia a prevádzky.** V tejto oblasti by teda systémy AI mali byť tzv. odolné voči budúcnosti (future proof).

- biometrická identifikácia v reálnom čase na verejných priestranstvách orgánmi štátnej moci (iní ani nemôžu) okrem oprávnených výnimiek (face recognition môže mať výnimku, no napr. social scoring nie)

Vysoko rizikové systémy (čl. 6 + prílohy II. a III.), na ktoré sa bude viazať systém povinností, resp. zodpovednosti a certifikácie v rámci EÚ, sa radia do viacerých tzv. sektorov:

- zamestnávanie
- vzdelávanie
- zdravotníctvo
- prístup k verejným službám
- biometrická identifikácia
- overovanie bonity osôb a pod.

Kategória Ostatné systémy AI zahŕňa všetky ostatné technológie umelej inteligencie, ktoré – až na čl. 52 – nie sú viazané žiadnymi osobitnými povinnosťami. Čl. 52 nariaďuje, aby pre používateľov bolo zrejmé, že interagujú so systémom AI. V zásade možno povedať, že Ostatné systémy AI sú viazané povinnosťou transparentnosti. Navyiac sa odporúča, aby členské štáty podporovali dobrovoľné prijatie etických záväzkov, napr. prostredníctvom kódexu správania a pod.

Dôležitou súčasťou nariadenia o AI sú tzv. **významné povinnosti poskytovateľov**.⁴³³

- **riadenie rizík** (čl. 9): povinnosť implementovať interné procesy za účelom identifikácie, analýzy a zmierňovania následkov rizík v súvislosti s vysoko rizikovými systémami AI.
- **správa dát** (čl. 10): požaduje, aby testovacie a tréningové datasets boli „vysokej kvality“, aby tak systém AI, ktorý ich využíva, nebol diskriminačný a nebude vytvárať nepredvídané, resp. nesprávne výsledky.
- **technická dokumentácia** (čl. 11): stanovuje povinnosť pripraviť detailnú technickú dokumentáciu (jej rozsah je definovaný v prílohe IV), ktorá umožní auditovanie systému AI, vrátane overenia jeho výsledkov.

⁴³³ PATTYNOVÁ, *Výzvy a právni aspekty umělé inteligence*.

Dostupné na internete: <<https://www.tuesday.cz/akce/umela-intelligence-1/zaznam-akce/>>

- **interné záznamy** (čl. 12 a 20): nariaďuje požiadavku na zaznamenávanie (logovanie) a uchovávanie záznamov jednotlivých udalostí spojených s vývojom a využívaním systému AI.
- **transparentnosť** (čl. 13): ukladá povinnosť poskytnúť používateľovi dokumentáciu a manuál pre používanie systému AI, vďaka ktorému používateľ môže porozumieť a vykonávať kontrolu nad vytváraním výsledkov systému AI.
- **ľudský dohľad** (čl. 14): vyžaduje zabezpečenie možnosti zásahu kvalifikovaných osôb určených používateľom, predovšetkým schopnosť celkom prerušiť prevádzku systému AI a zmeniť výsledok vytvorený týmto systémom.
- **presnosť, robustnosť (spoľahlivosť) a kybernetická bezpečnosť** (čl. 12): požaduje, aby miera presnosti bola deklarovaná v technickej dokumentácii, aby systém AI bol odolný voči chybám a nejasnostiam, a tiež proti škodlivým zásahom tretích strán.
- **systém riadenia kvality** (čl. 17): ukladá povinnosť implementovať systém riadenia kvality, ktorý by minimálne mal zahŕňať rozsiahlu internú dokumentáciu a procesy pre zabezpečenie testovania a preverovania.
- **monitorovanie po uvedení na trh** (post-marketing monitoring, čl. 54): vyjadruje povinnosť zabezpečiť monitorovanie v súlade s nariadením aj po uvedení do predaja, resp. prevádzky, ktoré by bolo založené na dátach poskytnutých používateľmi, resp. získaných z iných zdrojov.
- **posúdenie zhody** (čl. 43): okrem výnimiek (zdravie, bezpečnosť osôb,...) sa požaduje posúdenie vykonávané priamo poskytovateľmi systému AI (self-assessment), a tiež vykonanie jeho aktualizácie pri každej významnej zmene systému.
- **registrácia** (čl. 61): ukladá povinnosť registrovať systém AI a poskytnúť o ňom informácie podľa prílohy VII. Tieto informácie budú vedené v databáze EÚ.
- **oznamovanie incidentov** (čl. 62): stanovuje povinnosť hlásiť príslušným orgánom členských štátov EÚ všetky závažné incidenty alebo poruchy systému AI, ktoré by predstavovali porušenie povinností EÚ pre ochranu ľudských práv.

Významné povinnosti poskytovateľov podľa Aktu o umelej inteligencii prakticky vyjadrujú veľa – podľa predkladateľov nariadenia to podstatné – z návrhov a postrehov, ktoré sme

rozoberali v 2. kapitole a sumarizovali v kapitole 3.1. Z nášho pohľadu tak zapracovanie podstatných požiadaviek na eticko-právne regulácie do pripravovaného nariadenia EÚ môže byť dôvodom pre veľké uspokojenie a zadosťučinenie.

Množina významných povinností však prináša aj jedno veľké jarmo: z takmer ideálu pre riešenie dôveryhodného, transparentného a na človeka orientovaného systému AI sa týmto nariadením stáva minimálny štandard – zákonná povinnosť, ktorú nebude ľahké splniť.

Ponajprv, nie je po technickej a procesnej stránke vôbec ľahké niektoré povinnosti splniť. Uvedme malý príklad – je v súčasnosti reálne mať dokonalý tréningový dataset, alebo mať tak vytrénovaný systém AI, aby neniesol riziká *long-tail* efektu? Nájdeme viacero scenárov, ktoré poukazujú na extrémnu obťažnosť splnenia niektorých povinností...

Výzvou je i prípadný nesúlad s už existujúcim nariadením o ochrane osobných údajov (GDPR). Ako napríklad zosúladiť princíp minimalizácie údajov z GDPR s požiadavkou na úplnosť datasetov a povinnosťou uchovávanía logov s citlivým obsahom? GDPR stanovuje tituly, t.j. dôvody pre spracúvanie osobných údajov. Ako obhájiť dôvody špecifické pre konkrétne algoritmy AI? Ako zosúladiť požiadavku na používateľské dáta, vďaka ktorým sa systém „učí“, s podmienkami GDPR? Viaceré problémy ešte čakajú na vyriešenie...⁴³⁴

Ďalším problémom je ekonomická, odborná i časová náročnosť spojená so zabezpečením kompatibility s uvedeným nariadením. Čím zložitejší systém a navyše ak ide vysoko rizikový systém, tým väčšou výzvou bude implementácia požiadaviek nariadenia.

A do tretice – vyvinúť a prevádzkovať systém kompatibilný s Aktom o umelej inteligencii sa prejaví vo zvýšenej ekonomickej záťaži v porovnaní s konkurenciou. Pravda, bavíme sa o trhoch mimo EÚ, na ktoré by nemali dosah ani rozšírené teritoriálne dopady tohoto nariadenia.

Znovu sa vrátiac k analógii nariadenia o ochrane osobných údajov (GDPR) – i v prípade Aktu o umelej inteligencii sa EÚ bude potýkať s rovnakými problémami „veľkého jarma“ regulačných požiadaviek. S odstupom času možno povedať, že vzhľadom na nekontrolovateľné el-dorado netransparentného narábania s osobnými údajmi naprieč celým digitálnym svetom bolo GDPR zásahom v hodine dvanástej a jeho benefity neustále

434 PATTYNOVÁ, *Výzvy a právni aspekty umělé inteligence*.

Dostupné na internete: <<https://www.tuesday.cz/akce/umela-inteligence-1/zaznam-akce/>>

prichádzajú. Silný regulačný imperatív (regulácia stupňa nariadenie, vysoké pokuty a pod.) a pomerne dobre zvládnutá legiskvančná lehota naviac zabezpečili, že nariadenie prinútilo tvorcov, prevádzkovateľov a poskytovateľov informačných systémov prispôbiť sa a akceptovať pravidlá hry.

Vzhľadom na solídny regulačný obsah Aktu o umelej inteligencii, možnosti jeho nasadenia a vymáhania podľa vzoru a na základe skúseností s GDPR si myslíme, že toto nariadenie má potenciál byť dôležitým vkladom pre zabezpečenie etického a právne rámca využívania systémov umelej inteligencie (ANI) v informačnej spoločnosti a digitálnom veku.

Akt o umelej inteligencii v súčasnosti prechádza legislatívnym konaním v rámci orgánov EÚ. Dúfajme, že jeho ovocím bude zachovanie a prípadné doplnenie dôležitých regulačných princípov, okorenené schopnosťou nájsť riešenia tých parciálnych problémov, ktoré by inak celoplošné nasadenie nariadenia diskvalifikovali.

I keď má Akt o umelej inteligencii potenciál širokospektrálneho pokrytia využívania technológií umelej inteligencie v modernej spoločnosti, už z princípu nemôže pokryť dve dôležité oblasti: nasadenie v oblasti pokročilého riadenia štátu, spravodajstva a dohľadu a nasadenie vo vojenskej oblasti.⁴³⁵

Vzhľadom na špecifikum a dosah nasadenia systémov AI v spravodajských službách⁴³⁶ pre ľudské práva, ochranu demokracie a slobôd si myslíme, že táto oblasť by mala byť pokrytá už základnými legislatívnymi mechanizmami a verejným dohľadom demokratickej spoločnosti, ktoré sa týkajú pôsobenia spravodajských služieb vo všeobecnosti.

Keďže akékoľvek obmedzovanie technológií umelej inteligencie vo vojenskej oblasti môže byť chápané ako bezpečnostné riziko a zníženie bojaschopnosti modernej armády, jednostranne prijaté regulácie nemusia byť účinné – nielen pre to, že sa jednou stranou ťažko prijímajú (aj keď pre hodnotovo orientovanú spoločnosť by to malo byť povinnosťou), ale i pre malú šancu na ich extra teritoriálne rozšírenie a akceptovanie. To však neznamená, že by sme mali na eticko-právne regulácie AI v armáde resignovať. Veď v predchádzajúcich kapitolách spomenuté eticko-právne procesy prebiehajúce napríklad v armáde USA môžu byť základom pre celosvetovú diskusiu mocností a na základe tlaku

435 Čl. 2, 3: „Toto nariadenie sa neuplatňuje na systémy umelej inteligencie vyvinuté alebo používané výlučne na vojenské účely.“

436 Aspekty nasadenia systémov AI v oblasti spravodajstva a algokracie sme diskutovali v kapitole 2.6.

verejnosti, angažovanosti jednotlivých častí spoločnosti v rôznych regiónoch sveta i úsilia zodpovedných strán môžu viesť k prijatiu celosvetových pravidiel i záväzkov pre oblasť vývoja a nasadenia rizikových vojenských systémov vybavených technológiami umelej inteligencie. Vývoj a nasadenie viacerých extrémne nebezpečných vojenských technológií je v súčasnosti na základe vzájomného konsenzu a právnych záväzkov zakázané, je preto žiadúce, aby sa tak stalo aj v prípade niektorých scenárov nasadenia smrtiacich automatických zbraňových systémov a ďalších rizikových technológií využívajúcich algoritmy umelej inteligencie.⁴³⁷

V akčnom rádiuse doteraz uvedených legislatívnych aktivít sa prevažne nachádzajú slabé systémy umelej inteligencie (ANI). **I keď základný eticko-právny rámec zostáva v platnosti, nie je jednoduché právne uchopiť systémy, ktoré sa aspoň v niektorých aspektoch približujú silnej umelej inteligencii (AGI). Problémy pramenia i z technologickej náročnosti daného predmetu práva⁴³⁸ i z celého komplexu otvorených otázok kognitívnych schopností a modelov správania, aj keď sa v súčasnosti ešte nejedná o právne reakcie na teóriu mysle alebo simulácie emócií a vedomia AGI.**

V prípade technológií silnej a všeobecnej umelej inteligencie bude vyvstávať otázka ich právneho postavenia.⁴³⁹ Mohla by za istých okolností existovať určitá spôsobilosť systému AGI byť nositeľom práv a povinností predvídaných právnym poriadkom? Prípadne, mohol by systém AGI vlastnými (právnymi?) úkonmi nadobúdať práva a brať na seba povinnosti? Ako by to bolo v prípade schopností analogických ľudskému uvažovaniu a prejavovaniu vôle? Kto by bol nositeľom autorských práv v prípade diel

437 Tieto scenáre sme rozoberali v kapitole 2.7.

438 Autor sa aj u renomovaných právnických agentúr stretol s analytickými a aplikačnými medzerami, ktoré vyplývali z nedostatku pochopenia princípov a fungovania zložitejších systémov AI.

439 Už v roku 2017 Európsky parlament, riešiac problematiku robotických systémov, reflektoval problematiku právneho statusu AI vo výzve Komisii, aby sa venovala vytvoreniu „špecifického právneho postavenia pre roboty v dlhodobom horizonte, aby sa aspoň tie najsofistikovanejšie autonómne roboty mohli považovať za elektronické osoby zodpovedné za náhradu akejkoľvek škody, ktorú môžu spôsobiť.“ *Uznesenie Európskeho parlamentu zo dňa 16. 2. 2017 obsahujúce odporúčania pre Komisiu k normám občianskeho práva v oblasti robotiky (2015/2103(INL))*. [on-line]. [cit. 29. marca 2022]. Dostupné na internete: <<https://eur-lex.europa.eu/legal-content/SK/TXT/PDF/?uri=CELEX:52017IP0051&from=EN>>

vytvorených algoritmami umelej inteligencie?⁴⁴⁰

I tieto otázky poukazujú na komplexnosť výzvy a potencionálny dopad budúcich systémov AGI prakticky na akúkoľvek oblasť života spoločnosti.

3.4. Postoj Cirkvi

Naliehavosť potreby skúmať, navrhnuť, prijať a realizovať etický rámec vývoja, používania a fungovania technológií umelej inteligencie rezonuje aj v Katolíckej cirkvi. Mnohorakým spôsobom je táto téma uchopená v rámci akademického prostredia, predovšetkým v interdisciplinárnom poňatí.⁴⁴¹ V agende vatikánskych inštitúcií je problematika AI pokračovaním angažovanosti v sociálnej sfére, nadväzujúc napríklad na širšie koncipované aktivity Pápežskej akadémie pre život, ktoré sme spomínali v úvode tejto práce.⁴⁴²

Doteraz najširšie koncipovaným počínom Pápežskej akadémie pre život⁴⁴³ v oblasti etiky umelej inteligencie bola vo februári 2020 zorganizovaná konferencia **renAIssance 2020**.⁴⁴⁴

440 ŠTARHA, Š., GAŠPAROVIČ, R. *AI z pohľadu práva*. [on-line]. [cit. 29. marca 2022].

Dostupné na internete: <<https://www.epravo.sk/top/clanky/ai-z-pohladu-prava-4483.html>>

441 Jedným z príkladom môže byť napríklad nedávny webinár "Man, Machine, and the Future of AI", ktorý v rámci série prednášok *Conversations That Matter: The Crossroads of Science and Human Dignity Fall 2021* organizoval McGrath Institute for Church Life, University of Notre Dame.

Conversations That Matter: The Crossroads of Science and Human Dignity Fall 2021. [on-line]. [cit. 28. marca 2022].

Dostupné na internete: <<https://mcgrath.nd.edu/conferences/academic-pastoral/conversations-that-matter-the-crossroads-of-science-and-human-dignity/conversations-that-matter-the-crossroads-of-science-and-human-dignity-fall-2021/>>

442 Išlo o podujatie Hackaton 2018, ktoré sa venovalo jednému z vážnych sociálnych problémov informačnej spoločnosti a digitálneho veku – digitálnemu rozdeleniu (digital divide). V duchu hesla „zapojenie mladých a technológií v prospech spoločného dobra“ sa v kontexte využitia technológií hľadali odpovede na otázky sociálneho začlenenia, medzi náboženského dialógu a zdrojov pre utečencov. *Vatican Hackathon – harnessing youth, technology to serve common good*. [on-line]. [cit. 28. marca 2022].

Dostupné na internete: <<https://www.vaticannews.va/en/vatican-city/news/2018-03/vatican-hackathon--.html>>

443 Dostupné na internete: <<https://www.academyforlife.va/>>

444 Konferencia bola známa aj pod názvom Rome Call for AI Ethics, či na sociálnych sieťach pod značkou #RomeCallforAIethics.

Prestížny AI Index Report publikovaný každoročne inštitútom HAI of Stanford University zaradil túto konferenciu medzi najhorúcejšie témy roku 2020.⁴⁴⁵

Na konferencii, ktorá sa konala 26. - 28. februára 2020 v Ríme, vystúpilo viacero špičkových odborníkov na etiku umelej inteligencie z akademickej oblasti i vedúcich predstaviteľov veľkých korporácií z oblasti informačných technológií.⁴⁴⁶

Jedným z výsledkov konferencie bolo aj podpísanie **Výzvy na etiku v umelej inteligencii**⁴⁴⁷ so zámerom „**akcentovať etický prístup k umelej inteligencii a podporovať zmysel pre zodpovednosť medzi organizáciami, vládami a inštitúciami s cieľom vytvoriť budúcnosť, v ktorej digitálne inovácie a technologický pokrok slúžia ľudským schopnostiam a tvorivosti, a nie ich postupnému nahrádzaniu**“.⁴⁴⁸

Signatármi výzvy boli Vincenzo Paglia, predseda Pápežskej akadémie pre život, Brad Smith, prezident spoločnosti Microsoft, John Kelly III, viceprezident IBM, Qu Dongyu, generálny riaditeľ FAO a Paola Pisano, ministerka inovácií Talianska.

Zámerom organizátora, Pápežskej akadémie pre život, bol **vstup do diskusie o etických kritériách AI v jednotlivých oblastiach vývoja a nasadenia**. K nosným témam patrila **diskusia zameraná na snahu o premostenie pohľadov jednotlivých krajín, firiem, záujmových skupín na oblasť etiky umelej inteligencie**, pričom naplneným cieľom konferencie sa stalo **stanovenie konkrétnych etických kritérií**.

Sumarizáciu tohoto úsilia vyjadruje vyššie spomenutý záverečný dokument, ktorého súčasťou je identifikovanie troch oblastí vplyvu a formulovanie šiestich princípov pre použitie umelej inteligencie pre dobro človeka a bez strachu zo zneužitia.

Akcentované oblasti vplyvu sú vo výzve definované nasledovne:

445 *Measuring trends in Artificial Intelligence – 2021 AI Index Report*. [on-line]. [cit. 28. marca 2022].

Dostupné na internete: <<https://aiindex.stanford.edu/ai-index-report-2021/>>

446 *The „good“ algorithm*. [on-line]. [cit. 28. marca 2022].

Dostupné na internete: <https://www.academyforlife.va/content/dam/pav/documenti%20pdf/2020/Assemblea/Atti_Assemblea_e_28febbraio/Atti%20completi_PAV_2020_.pdf>

447 *Rome Call for AI Ethics (document)*. [on-line]. [cit. 28. marca 2022].

Dostupné na internete:

<https://www.romecall.org/wp-content/uploads/2022/03/RomeCall_Paper_web.pdf>

448 *Rome Call for AI Ethics*. [on-line]. [cit. 19. augusta 2020].

Dostupné na internete: <<http://www.academyforlife.va/content/pav/en/events/intelligenza-artificiale.html>>

- **etika** – všetky ľudské bytosti sa rodia slobodné a rovné v dôstojnosti a právach.
- **vzdelávanie** – premeniť svet prostredníctvom inovácie umelej inteligencie znamená zaviazat' sa budovať budúcnosť pre mladé generácie a s nimi.
- **práva** – vývoj technológií AI v službách ľudstva a planéty sa musí odrážať v predpisoch a zásadách, ktoré chránia ľudí, najmä slabých a znevýhodnených, a prírodné prostredie.

Považujeme za vhodné uvedené oblasti vplyvu popísať priamo slovami záverečného dokumentu *Rome Call for AI Ethics*, nakoľko ide o hutný text sumarizujúci postrehy, návrhy a závery takmer tristo stranového konferenčného zborníka.

Oblasť etiky stavia na všeobecnej deklarácii OSN o ľudských právach. Všetky ľudské bytosti sa rodia slobodné a rovné v dôstojnosti a právach. Sú obdarené rozumom a svedomím a mali by voči sebe konať v duchu spolupatričnosti (por. čl. 1 Dekréту OSN o ľudských právach). Táto základná podmienka slobody a dôstojnosti musí byť chránená a zaručená aj pri výrobe a používaní systémov umelej inteligencie. To sa musí uskutočniť zabezpečením práv a slobody jednotlivcov tak, aby neboli algoritmami diskriminovaní z dôvodu svojej „rasy, farby pleti, pohlavia, jazyka, náboženstva, politického alebo iného zmýšľania, národného alebo sociálneho pôvodu, majetku, rodu alebo iného postavenia“ (čl. 2).⁴⁴⁹

Systémy AI musia byť koncipované, navrhnuté a implementované tak, aby slúžili a chránili ľudské bytosti a prostredie, v ktorom žijú. Tento základný pohľad sa musí premietnuť do záväzku vytvoriť také životné podmienky (spoločenské aj osobné), ktoré umožnia skupinám aj jednotlivým členom snažiť sa o plné vyjadrenie, ak je to možné.⁴⁵⁰

Aby bol technologický pokrok v súlade so skutočným pokrokom ľudského rodu a úctou k planéte, musí spĺňať tri požiadavky. Musí zahŕňať každú ľudskú bytosť a nikoho nediskriminovať; musí mať na zreteli dobro ľudstva a dobro každej ľudskej bytosti; napokon musí brať do úvahy zložitú realitu nášho ekosystému a vyznačovať sa tým, ako sa stará o planétu (náš „spoločný a zdieľaný domov“) a chráni ju vysoko udržateľným spôsobom, ktorý zahŕňa aj využívanie umelej inteligencie pri zabezpečovaní udržateľných

449 *Rome Call for AI Ethics (document)*. [on-line]. Str. 5. [cit. 28. marca 2022].

Dostupné na internete:

<https://www.romecall.org/wp-content/uploads/2022/03/RomeCall_Paper_web.pdf>

450 *Rome Call for AI Ethics (document)*, s. 5.

potravinových systémov v budúcnosti. Okrem toho si každý človek musí byť vedomý toho, kedy komunikuje so strojom.⁴⁵¹

Technológia založená na umelej inteligencii sa nikdy nesmie používať na akékoľvek zneužívanie ľudí, najmä tých najzraniteľnejších. Namiesto toho sa musí používať na pomoc ľuďom pri rozvíjaní ich schopností a na podporu planéty.⁴⁵²

Oblasť vzdelávania je pozvaním premieňať svet prostredníctvom inovácie umelej inteligencie, a tak sa zaviazat' budovať budúcnosť pre mladé generácie a s nimi. Tento zámer sa musí odraziť v záväzku k vzdelávaniu, vo vypracovaní špecifických učebných osnov, ktoré zahŕňajú rôzne humanitné, vedecké a technické disciplíny, a v prevzatí zodpovednosti za vzdelávanie mladších generácií. Tento zámer tiež znamená snahu o zlepšenie kvality vzdelávania, ktorého sa mladým ľuďom dostáva; toto vzdelávanie sa musí uskutočňovať prostredníctvom metód, ktoré sú prístupné pre všetkých, ktoré nediskriminujú a ktoré môžu ponúknuť rovnosť príležitostí a zaobchádzania. Všeobecný prístup k vzdelávaniu sa musí dosiahnuť prostredníctvom zásad solidarity a spravodlivosti.⁴⁵³

Prístup k celoživotnému vzdelávaniu musí byť zaručený aj pre starších ľudí, ktorým sa musí ponúknuť možnosť prístupu k off-line službám počas digitálneho a technologického prechodu. Okrem toho sa tieto technológie môžu ukázať ako nesmierne užitočné pri pomoci ľuďom so zdravotným postihnutím, aby sa mohli učiť a stať sa nezávislejšími: inkluzívne vzdelávanie preto znamená aj využívanie umelej inteligencie na podporu a integráciu každého človeka, poskytovanie pomoci a možností sociálnej participácie (napr. práca na diaľku pre osoby s obmedzenou mobilitou, technologická podpora pre osoby s kognitívnymi poruchami atď.).⁴⁵⁴

Vplyv transformácií, ktoré priniesol fenomén umelej inteligencie v spoločnosti, v práci a vo vzdelávaní, si vynútil prepracovanie školských osnov, aby sa motto vzdelávania "nikto nezostane pozadu" stalo skutočnosťou. V oblasti vzdelávania sú potrebné reformy s cieľom zaviesť vysoké a objektívne normy, ktoré môžu zlepšiť individuálne výsledky. Tieto normy by sa nemali obmedzovať na rozvoj digitálnych zručností, ale mali by sa

451 Rome Call for AI Ethics (document), s. 6.

452 Rome Call for AI Ethics (document), s. 6.

453 Rome Call for AI Ethics (document), s. 7.

454 Rome Call for AI Ethics (document), s. 7.

zamerat' na to, aby každý človek mohol naplno prejavit' svoje schopnosti a pracovať pre dobro komunity, aj keď z toho nemá osobný prospech.⁴⁵⁵

Pri navrhovaní a plánovaní spoločnosti zajtrajška musí využívanie systémov AI sledovať formy činnosti, ktoré sú sociálne orientované, tvorivé, prepojené, produktívne, zodpovedné a schopné pozitívne ovplyvniť osobný a spoločenský život mladších generácií. Sociálny a etický rozmer umelej inteligencie musí byť aj jadrom vzdelávacích aktivít v oblasti AI.⁴⁵⁶

Hlavným cieľom tohto vzdelávania musí byť zvyšovanie povedomia o príležitostiach a tiež o možných kritických problémoch, ktoré fenomén AI predstavuje z hľadiska sociálneho začlenenia a rešpektu k individuálnej osobe.⁴⁵⁷

Oblast' práv sumarizuje, ako sa rozvoj umelej inteligencie v službách ľudstva a planéty musí odrážať v predpisoch a zásadách, ktoré chránia ľudí (najmä slabých a znevýhodnených) a prírodné prostredie. Etický záväzok všetkých zúčastnených strán je kľúčovým východiskom; na to, aby sa táto budúcnosť stala skutočnosťou, sú absolútne nevyhnutné hodnoty, zásady a v niektorých prípadoch aj právne predpisy, ktoré tento proces podporujú, systematizujú a usmerňujú.⁴⁵⁸

Na vývoj a implementáciu systémov umelej inteligencie, ktoré budú prospešné pre ľudstvo a planétu a zároveň budú slúžiť ako nástroje na budovanie a udržiavanie medzinárodného mieru, musí ísť vývoj umelej inteligencie ruka v ruku so spoľahlivými opatreniami v oblasti digitálnej bezpečnosti.⁴⁵⁹

Aby technológie AI mohli fungovať ako nástroj v prospech ľudstva a planéty, musíme do centra verejnej diskusie postaviť tému ochrany ľudských práv v digitálnej ére. Nastal čas položiť si otázku, či nové formy automatizácie a algoritmickej činnosti nevyžadujú vytvorenie silnejších zodpovedností. Predovšetkým bude nevyhnutné zvážiť určitú formu „povinnosti vysvetľovania“: musíme sa zamyslieť nad tým, aby boli zrozumiteľné nielen kritériá rozhodovania algoritmických agentov založených na umelej inteligencii, ale aj ich účel a ciele. Tieto zariadenia musia byť schopné ponúknuť jednotlivcom informácie o logike algoritmov používaných na rozhodovanie. Tým sa zvýši nielen transparentnosť,

455 Rome Call for AI Ethics (document), s. 7-8.

456 Rome Call for AI Ethics (document), s. 8.

457 Rome Call for AI Ethics (document), s. 8.

458 Rome Call for AI Ethics (document), s. 9.

459 Rome Call for AI Ethics (document), s. 9.

sledovateľnosť a zodpovednosť, ale i adekvátnosť a váha rozhodovacieho procesu podporovaného počítačom. Je potrebné podporovať nové formy regulácie na podporu transparentnosti a dodržiavania etických zásad, najmä v prípade pokročilých technológií, ktoré majú vyššie riziko vplyvu na ľudské práva, ako je napríklad rozpoznávanie tváre.⁴⁶⁰

Praktickým vyjadrením konferenčných záverov zhrnutých v uvedených troch oblastiach sa stalo šesť princípov využitia systémov AI pre dobro človeka a bez strachu zo zneužitia:

- **transparentnosť** – algoritmy AI musia byť transparentné a zrozumiteľné pre všetkých.
- **inklúzia** – je potrebné zohľadniť potreby všetkých ľudí, aby z toho mali prospech všetci a aby sa všetkým jednotlivcom poskytli čo najlepšie podmienky na sebarealizáciu a rozvoj. Systémy AI nesmú nikoho diskriminovať, pretože každý človek má rovnakú dôstojnosť.
- **zodpovednosť** – tí, ktorí navrhujú a zavádzajú používanie technológií AI, musia postupovať zodpovedne a transparentne. Vždy musí existovať niekto, kto preberá zodpovednosť za to, čo systém AI vykonáva.
- **neutrannosť** – netvorit' a nekonať na základe zaujatosti, a tak chrániť spravodlivosť a ľudskú dôstojnosť. Systémy AI sa predsudkami nesmú nechať viesť alebo ich vytvárať.
- **spoľahlivosť** – systémy AI musia byť schopné spoľahlivo fungovať.
- **bezpečnosť a súkromie** – technológie AI musia byť bezpečné a rešpektovať súkromie používateľov.

Cieľ, ktorý konferenciou *renAIssance* Pápežská akadémia pre život sledovala, vyjadril Mons. Vincenzo Paglia slovami: „Zámerom výzvy je vytvoriť hnutie, ktoré sa rozšíri a zapojí ďalších aktérov: verejné inštitúcie, mimovládne organizácie, priemyselné odvetvia a skupiny, aby určili smer vývoja a používania technológií odvodených od umelej inteligencie. Z tohto hľadiska môžeme povedať, že prvý podpis tejto výzvy nie je vyvrcholením, ale východiskom pre záväzok, ktorý sa javí ako ešte naliehavejší a dôležitejší než kedykoľvek predtým. Pripojenie sa k tejto iniciatíve pre zástupcov priemyslu, ktorí ju podpíšu, znamená záväzok, ktorý má význam aj z hľadiska nákladov a priemyselného príspevku k vývoju a distribúcii ich výrobkov. Ak sa Akadémia cíti byť

⁴⁶⁰ Rome Call for AI Ethics (document), s. 9.

povolaná zintenzívniť svoje úsilie o uľahčenie získavania poznatkov a podpisov iných medzinárodných aktérov, táto výzva je len prvým krokom, po ktorom by mali nasledovať ďalšie. Text výzvy sa vyznačuje aj tým, že je prvým **pokusom sformulovať súbor etických kritérií so spoločnými referenčnými bodmi a hodnotami**, čím ponúka príspevok k rozvoju spoločného jazyka na interpretáciu toho, čo je človek“.

Celá výzva je v duchu kresťanskej antropológie intenzívne zameraná na človeka. Ide o viackrát spomínaný *human-centered* prístup, ktorý je však v záveroch konferencie širšie koncipovaný s osobitným zreteľom na sociálnu spravodlivosť, rovnosť a ľudskú dôstojnosť.

Dôležitým vyjadrením je **požiadavka, aby etický záväzok bol základom pre regulácie a právne normy.**

V niektorých aspektoch badať jasný vplyv osobitne v poslednej dekáde prebiehajúcich etických a legislatívnych procesov, ktoré boli integrované do Výzvy na etiku v umelej inteligencii a zároveň sa stali aj súčasťou návrhu európskeho Aktu o umelej inteligencii. Keďže finálna verzia návrhu Aktu o umelej inteligencii bola predstavená v apríli 2021, bolo by zaujímavé sledovať, ako a do akej miery obsah rímskej Výzvy na etiku v AI ovplyvnil tento regulačný návrh EÚ.

V súčasnosti najucelenejší a v minulej kapitole predstavený regulačný rámec pre technológie umelej inteligencie, Akt o umelej inteligencii, tak stavia na etických a hodnotových základoch jasne korešpondujúcich s rímskou výzvou, pričom každý z nich – Akt i Výzva – má svoje vlastné zameranie. Akt ako nariadenie EÚ je striktným právnym rámcom vyžadujúcim dodržiavanie etických a právnych noriem. Výzva – ako uvádza Mons. Paglia – si kladie za cieľ iniciovať univerzálne hnutie, naprieč celým svetom i všetkými oblasťami spoločnosti, ktoré chce usmerňovať a formovať vývoj a používanie technológií umelej inteligencie podľa etických princípov a morálnych hodnôt pre dobro človeka, spoločnosti a sveta, akcentujúc dôstojnosť každého človeka a ochranu životného prostredia ako spoločného domu celého ľudstva.

Význam snahy o celosvetové hnutie, na ktorom aktívne participuje Katolícka cirkev je o to väčší, ak si uvedomíme hodnotovú a etickú diskrepanciu vo svete. Podľa štúdie *The global landscape of AI ethics guidelines* **v zásade panuje celospoločenský konsenzus ohľadom potreby etických usmernení v oblasti umelej inteligencie, ktorý je však sprevádzaný podstatnými rozdielmi v chápaní etiky a hodnôt:** „Naše výsledky ukazujú globálnu konvergenciu, ktorá sa objavuje okolo piatich etických princípov (transparentnosť,

spravodlivosť a férovosť, neškodnosť, zodpovednosť a súkromie), pričom sa podstatne líšia v tom, ako sa tieto zásady interpretujú; prečo sa považujú za dôležité, akej problematiky, oblasti alebo aktérov sa týkajú a ako sa by sa mali uplatňovať. Naše zistenia zdôrazňujú význam integrácie usmernení s podstatnou etickou analýzou a primeranou implementáciou stratégií.⁴⁶¹

I keď pre skutočne etický vývoj a nasadenie systémov AI sú podstatné tak komplexné i dosah majúce legislatívne rámce, ako je nariadenie EÚ Akt o umelej inteligencii, konferencia Pápežskej akadémie pre život, jej záber a výsledok nám správne pripomínajú celostvetové aktivity, angažovanosť a hnutie, ktoré nielenže vždy stojí na počiatku hľadania adekvátnych etických právnych noriem, ale je i katalyzátorom ich návrhu a prijatia, nehovoriac o raste celospoločenského povedomia.

Rímska konferencia renAIssance 2020 len potvrdila, že Katolícka cirkev z hĺbky svojej náuky má čím prispieť do tohto spoločného diela rozvoja moderného sveta.

I v tomto prípade treba však uviesť, že spomínané aktivity a pomerne jasné i komplexné riešenia sa primárne týkajú systémov slabej umelej inteligencie (ANI). Riešenie výziev, ktoré sa viažu na silnú, resp. uvedomelú umelú inteligenciu (AGI), je v súčasnosti doménou skôr akademického výskumu a interdisciplinárneho bádania s prienikom do oblasti transhumanizmu, neurobioetiky,⁴⁶² kognitívnych vied a pod.

Sumarizujúc tému umelej inteligencie a etiky si uvedomujeme viaceré výzvy v tejto oblasti:

- ako sa principiálne k fenoménu umelej postavíme
- ako ju uchopíme, aký rámec zvolíme
- aké všeobecné a základné etické princípy stanovíme
- aké parciálne usmernenia či špecifické odporúčania vyjadríme

461 JOBIN, A., IENCA, M., VAYENA, E. *The global landscape of AI ethics guidelines*. In: *Nat Mach Intell*. [on-line]. 2019, 1, s. 389–399. [cit. 28. marca 2022].

Dostupné na internete: <<https://doi.org/10.1038/s42256-019-0088-2>>

462 Týmto otázkam sa napríklad venujú na Oddelení neurobioetiky, Ateneo Pontificio Regina Apostolorum.

Dostupné na internete: <<https://www.upra.org/en/ricerca/gruppi-di-ricerca/neurobioetica/>>

4. Navrhnuté riešenie etických problémov ANI

*Nič veľkého nevstúpi do života smrteľníkov bez prekliatia.*⁴⁶³

Začínať kapitolu týmito slovami dramatika, vzhľadom na svoju tvorbu priam tragéda, naviac politika i stratéga v jednej osobe, Sofoklesa je skutočne deprimujúce a možno trochu odzrkadľujúce naše zameranie sa na problémy a temnú stránku súčasných technológií umelej inteligencie.

Kladieme si však otázku, či to tak skutočne musí byť? Ide len o vyjadrenie dejinnej skúsenosti s ľudskou slabosťou neschopnou odolať pokušeniu, ktoré je tým väčšie, čím väčší je potenciál novej veci, ktorú spoločnosť objavila? Alebo ide skôr o sarkazmus, z čias antiky na svetlo sveta vytiahnutý, vyjadrujúci spôsob, akým sa moderný človek zmieruje s realitou, ktorá sa nedá zmeniť?

Skúsme to teda inak...

„A Boh videl všetko, čo urobil, a hľa, bolo to veľmi dobré.“

4.1. Základný postoj vo svetle Zjavenia⁴⁶⁴

Základné posolstvo o stvorení sveta a človeka komplexne hovorí o stvorení, že „to bolo dobré“.⁴⁶⁵ Zároveň človek bol stvorený na Boží obraz (צֶלֶם אֱלֹהִים) s rozumom a slobodnou vôľou, aby spravoval zverený svet.⁴⁶⁶ Človek ako obraz Boha bol uschopnený komunikovať so svojím Stvoriteľom i s ďalšími ľuďmi. K prirodzenému zákonu, ktorý je vpísaný do ľudskej duše a svedomia dostáva i pozitívny Boží zákon – Desatoro, ktoré

⁴⁶³ Sofokles.

⁴⁶⁴ Kapitola 4.1. je z veľkej časti prevzatá z licenciátskej práce autora, konkrétne z kapitoly 3.1. *Principiálny pohľad na základe Božieho zjavenia*.

ŠANTAVÝ, *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi*, [online], s. 46-48. [cit. 14. marca 2022].

Dostupné na internete:

<https://peter.santavy.cloud/docs/Analiza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

⁴⁶⁵ „A Boh videl všetko, čo urobil, a hľa, bolo to veľmi dobré.“ (Gn 1,31).

⁴⁶⁶ „Plodte a množte sa a naplňte zem! Podmaňte si ju a panujte nad rybami mora, nad vtáctvom neba a nad všetkou zverou, čo sa hýbe na zemi!“ (Gn 1,28).

opisuje človeka, nielen stvoreného na obraz Boží, ale človeka, ktorý tento obraz i žije: **vedome sa rozhoduje a koná tak, aby budoval vzťah voči svojmu Stvoriteľovi a múdro užíval zverený svet.**

Užívanie (podmanenie) zvereného sveta je v Gn 1,28 vyjadrené hebrejskými výrazmi „podmaňte si ju“ (וְכַבְּשֶׁתָּ) a „panujte“ (וְיָרְדוּ), ktoré sú odvodené od starostlivosti hlavy rodiny, či pastierskeho „užívania“ stáda, t.j. od **schopnosti užívať, rozvíjať a z generácie na generáciu si odovzdávať stvorený svet.**⁴⁶⁷

Tento základný pohľad hovoriaci o principiálnom postoji k dielu stvorenia, spravovaní a rozvoji tohoto sveta sa nesie naprieč celým starozákonným zjavením a vrcholí v *ketubím* – v múdroslovnej literatúre Starého zákona, konfrontujúcej *chokmá*, t.j. Božiu múdrosť so *sofia*, t.j. múdrosťou vtedajšieho moderného helénskeho sveta.⁴⁶⁸ **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.**

Výzvy Ježiša Krista „Chodte teda, učte všetky národy...“ (Mt 28, 19) a „Chodte do celého sveta a hlásajte evanjelium všetkému stvoreniu“ (Mk 16, 15) vyjadrujú **nadčasové poslanie, zahŕňajúce civilizačný rozvoj**, v ktorom „Budete mi svedkami v Jeruzaleme i v celej Judei aj v Samárii a až po samý kraj zeme“ (Sk 1,8).

Novozákonný pohľad, spojený s ohlasovaním Evanjelia a prvopočiatkami pastoračnej starostlivosti, je veľmi pekne reprezentovaný konaním sv. Pavla v aténskom Areopágu, ktorý bol miestom koncentrácie vtedajšieho mozgového trustu – pôvodne miesto najvyššieho súdu, no v tom čase i miesto mysliteľov, predstaviteľov tej najlepšej filozofickej tradície helénskej kultúry a zároveň zástancov celého vtedajšieho Panteónu.⁴⁶⁹ Išlo o miesto, ktoré sa pre ohlasovateľa evanjelia Pavlových kvalít a apoštola národov stalo nemalou výzvou zvestovať Ježiša ako Krista a Pána.

V sedemnástej kapitole Skutkov apoštolov môžeme veľmi jasne vidieť, ako si Pavol počínal: využil svoje vzdelanie, znalosť gréckej filozofie a kultúry, aby tak spôsobom vlastným zhromaždeným Aténčanom začal svoje zvestovanie. Vieme, ako to prebiehalo – Pavlovo priblíženie sa ich zmýšľaniu malo úspech: dali mu slovo a on mohol začať svoju

467 Podobne i v Gn 2,15 Boh, keď umiestňuje človeka v raji, dáva mu poslanie „aby ho obrábal (לְעַבְדָּהּ) a strážil (לְשׁוּמָרָהּ)“.

468 Kniha Múdrosti a Kniha Sirachovcova.

469 Pavol vystúpil v aténskom Areopágu počas svojej Druhej misijnej cesty (r. 50 – 52).

evanjelizačnú reč. Kristovo zmŕtvychvstanie však bolo pre ľudí odchovaných mliekom platónskej filozofie prisilnou kávou a Pavol nakoniec končí fiaskom: „vypočujeme ťa o tom inokedy“ (Sk 17,32). I keď Písmo spomína, že niektorí predsa len uverili na jeho kázanie, bola to slabá náplast' na Pavlovo sklamanie, ktoré je krásne čitateľné v niektorých chvíľach ďalšieho Pavlovho misijného putovania (napr. 1Kor 2,1-5).

I keď mohol sv. Pavol predpokladať, že jeho hlásanie v Areopágu asi skončí fiaskom, predsa si uvedomoval, že s Evanjeliom nemôže obísť intelektuálne a kultúrne stredisko vtedajšieho sveta, ktoré sa podieľalo na modernom rozvoji starovekého stredomoria a helénskej kultúry. Vedel, že treba i naďalej v tomto úsilí pokračovať,⁴⁷⁰ a že po ňom prídu ďalší, ktorí budú prinášať Kristovo svetlo do odkrytých oblastí poznania⁴⁷¹ a konfrontovať i spájať moderné bádanie a rozvoj spoločnosti s Božím zjavením.⁴⁷²

Práve na základe udalostí v Areopágu v kontexte ohlasovania Evanjelia sv. Ján Pavol II. vo svojej encyklike *Redemptoris missio* o stálej platnosti misijného poslania poukazuje na moderné areopágy, ktoré sú priamou výzvou na permanentne platné pozvanie k misii ad gentes. Veľmi jasne definuje svet komunikácie ako prvý areopág súčasnosti⁴⁷³, ktorý sa podstatne podieľa na vytváraní global village – celosvetovej dediny, v rámci ktorej moderné komunikačné prostriedky presahujú rámec nástroja komunikácie a podieľajú sa na vytváraní nového kultúrneho kontextu – novej kultúry, techniky, vyjadrovania a psychológie vzťahov – jednoducho povedané na vytváraní niečoho, čo je predzvesťou informačnej spoločnosti. Tej spoločnosti, ktorej paradigmatická zmena zahŕňa aj čo najširšiu implementáciu technológií umelej inteligencie.⁴⁷⁴

I keď od uvedenia encykliky už ubehlo skoro tridsať rokov, v pohľade na rodiacu sa informačnú spoločnosť s jej technologickým portfóliom môžeme vidieť nadčasový záber a prorocké slová pápeža v dobe, keď pre mnohých vtedajších mysliteľov a vizionárov

470 Pavol počas Tretej misijnej cesty v Efeze vyučoval v škole gréckeho filozofa a učiteľa rétoriky Tyranna. Por. (Sk 19,9).

471 Sv. filozof a mučeník Justín, predstavitelia alexandrijskej a antiochiskej školy, scholastici, ...

472 Skvelým zdrojom informácií o civilizačnom rozmere pôsobenia Cirkvi je:

WOODS, E. T. Jr. *Ako Katolícka cirkev budovala západnú civilizáciu*. Bratislava: Redemptoristi - Slovo medzi nami, 2010.

473 JÁN PAVOL II. *Redemptoris Missio*. Praha: Zvon, 1994, s. 45-46.

474 Umelú inteligenciu v kontexte paradigmatickej zmeny informačnej spoločnosti sme uvádzali v kapitole 1.11.

virtuálny svet a kybernetický priestor, informačné a komunikačné prostriedky i algoritmy umelej inteligencie boli ťažko uchopiteľnými fenoménmi v kontexte rozvoja spoločnosti.

Analyzujúc súčasný svet, v ktorom sa technológie umelej inteligencie stávajú neodmysliteľnou súčasťou informačnej spoločnosti, môžeme s plnou zodpovednosťou označiť aj oblasť umelej inteligencie areopágom, fórom moderného veku, do ktorého – nasledujúc príklad svätého apoštola Pavla – treba s odvahou vstúpiť a vo svetle evanjelia i v tejto oblasti prispievať k nekončiacemu úsiliu budovať spravodlivý svet, chrániť dôstojnosť každého človeka a rozvíjať civilizáciu lásky.

Inak povedané, aj technológie umelej inteligencie sú súčasťou nášho dejinného kontextu, v ktorom sme pozvaní mať účasť na budovaní Božieho kráľovstva.

4.2. Interdisciplinárny rámec

Máme za to, že **skutočné riešenie etických problémov a výziev technológií umelej inteligencie nie je možné úspešne realizovať bez interdisciplinárneho rámca, v rámci ktorého sme dostatočne oboznámení aj s technologickou stránkou týchto systémov a psychologickými, sociologickými i právnymi aspektmi ich nasadenia.**

Dostatočné oboznámenie sa s technologickou stránkou nám predovšetkým umožňuje pochopiť podstatu a rozsah technologických limitov a rizík algoritmov AI, a následne si adekvátnejšie predstaviť ich dosah na jednotlivé oblasti reálneho nasadenia, dôsledky na život človeka a fungovanie spoločnosti.

Naviac sa domnievame, že pre kvalifikovanú filozofickú a teologickú diskusiu o podstatných aspektoch umelej inteligencie je nutnou podmienkou pochopenie rozdielov medzi slabou a silnou, resp. medzi úzko špecializovanou a všeobecnou umelou inteligenciou v ich technologickej podstate. Rozsah diskutovaných tém môže byť skutočne rozsiahly – napr. sociálna spravodlivosť, ľudské práva, spravodlivá vojna, bioetické otázky, transhumanizmus, podstata ľudskej bytosti,...

Identifikácia oblastí, v ktorých sa seriózne nasadenie systémov umelej inteligencie nezaobíde bez uspokojivého návrhu riešenia etických problémov, zasa vyžaduje dôkladnú analýzu psychologických, sociologických i právnych aspektov ich nasadenia.

V horizonte vývoja uvedomele umelej inteligencie môžu mať nezastupiteľnú úlohu viaceré

netechnické vedecké disciplíny, schopné prispieť k riešeniu modelov správania, teórie mysle, simulácie emócií i vedomia a pod.

K dôležitým rozmerom interdisciplinárnej komunikácie patrí aj spolupráca s inštitúciami na medzinárodnej i štátnej úrovni, bez ktorej nie je možné vytvárať reálne etické a právne regulačné rámce, a v neposlednom rade kooperácia s privátnym sektorom, ktorého vývojové kapacity a finančné i ľudské zdroje sú motorom vývoja a nasadenia systémov AI takmer do všetkých oblastí reálneho života.

4.3. Všeobecné návrhy

4.3.1. Umelá inteligencia zameraná na človeka

Základným princípom pre akýkoľvek systém umelej inteligencie je zameranie na človeka, teda známy a všeobecne prijímaný princíp **human-centered artificial intelligence**. Navrhujeme však a zdôrazňujeme, že by princíp umelej inteligencie zameranej na človeka mal:

- **byť chápaný v duchu kresťanskej antropológie.**
- **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.**

Chápeme, že human-centered prístup zahŕňa aj mnohé technologické, praktické a právne náležitosti vývoja a nasadenia systémov umelej inteligencie (mnohé z nich sme už spomínali), no bez vyššie uvedeného návrhu aplikácia tohto princípu môže podliehať relativizmu a erózii hodnôt, či postupnému vyprázdneniu jeho podstatných aspektov.

Uchopenie princípu zamerania umelej inteligencie na človeka v duchu kresťanskej antropológie znamená prijať pohľad na človeka v kontexte Božieho zjavenia. V tejto perspektíve vidíme človeka, ktorý je stvorený na Boží obraz a pozvaný ku spásu. Vnímame realitu ľudského bytia, ktoré je zranené hriechom a v dimenziách pozemského putovania hľadá cestu k svojmu bytostnému naplneniu v Bohu. V perspektíve vedeckého bádania si uvedomujeme nielen všetko to, čo človeka na tejto ceste ovplyvňuje, ale i to, čo

sa stáva kontextom a prostriedkom tejto cesty napĺňanej v zmluvnom dialógu s Bohom a ostatnými ľuďmi. Prijímame radostnú zvesť o vykúpení a nachádzame vzor opravdivého človeka a skutočného naplnenia nášho človečenstva v Ježišovi Kristovi. Učíme sa nazerať na blížnych optikou evanjeliovej lásky a dobra s osobitným akcentom na službu a ochranu slabých a trpiacich. Povolání do spoločenstva Najsvätejšej Trojice máme vytvárať ľudské spoločenstvo a budovať civilizáciu lásky.

Rámec kresťanskej antropológie a z neho prameniacci diapazón kresťanského humanizmu nám poskytuje aj dôležité morálne vyhranenie sa voči pokušeniu umelo vylepšovaného sveta prostredníctvom umelej inteligencie, či už na spôsob de Chardinovho vývoja od hominizácie cez divinizáciu až po bod Omega, alebo transhumanistického vylepšovania človeka, odovzdania primátu budúcej superinteligencii a pokriveného vnímania eschatologickej roviny viery v kontexte transfigurizmu, t.j. náboženského transhumanizmu.

4.3.2. Dôveryhodná umelá inteligencia

Zameranie umelej inteligencie na človeka je základným predpokladom pre akýkoľvek systém umelej inteligencie, s ktorým môžeme bez zbytočných obáv interagovať a primerane mu dôverovať.

Principiálny postoj orientácie na človeka sa tak stáva ekvivalentným problematike dôveryhodnosti umelej inteligencie, pričom je treba stanoviť podmienky, bez splnenia ktorých by nasadenie systémov AI do reálneho sveta, v ktorom interagujú s človekom a vplývajú na spoločnosť, nemalo byť umožnené.

Na základe našich záverov z kapitol 2.8. a 3.1. i podľa Etického usmernenia pre dôveryhodnú umelú inteligenciu Skupiny expertov na umelú inteligenciu pri EÚ, ktoré bolo uvedené v kapitole 3.3., formulujeme **základné požiadavky na dôveryhodné systémy AI**. Tieto systémy musia byť:

- **legálne** – vyhovovať požadovaným normám, zákonom i reguláciám a spĺňať všetky platné zákony a predpisy.
- **etické** – rešpektovať etické zásady a hodnoty.
- **robustné** – dosahovať potrebné štandardy bezpečnosti a robustnosti⁴⁷⁵ nielen z technologického hľadiska, ale zohľadňovať aj sociálne prostredie a dopady na spoločnosť.

⁴⁷⁵ Ide o podchytenie a riešenie širokého rámca problémov, ktoré sme rozoberali v 2. kapitole.

Uvedomujeme si, že za každou z uvedených požiadaviek sa nachádza celý komplex problémov, procesov a dynamiky implementácie parciálnych riešení, o ktorých nemôžeme tvrdiť, že sú dokonalé a definitívne. Sú však cestou k cieľu, ktorým je dôveryhodná umelá inteligencia zameraná na človeka.

U každej z uvedených oblastí by sme mali definovať hranicu, od ktorej môžeme technológie umelej inteligencie považovať za dôveryhodné.

V oblasti noriem, zákonov a regulácií môžeme súhlasiť s aktuálnym obsahom pripravovanej regulácie Európskej únie v oblasti umelej inteligencie, Aktom o umelej inteligencii a považovať ho za primeranú legislatívnu hranicu pre súčasné dôveryhodné nasadenie systémov AI.

I napriek tomu, že Akt o umelej inteligencii v súčasnosti v rámci orgánov EÚ prechádza legislatívnym konaním, ktoré môže priniesť určité zmeny, dúfame, že dôležité regulačné princípy zostanú zachované, resp. vhodne doplnené o riešenie tých parciálnych problémov, ktoré by celoplošné nasadenie nariadenia diskvalifikovali. Len tak môže byť do platnosti uvedený dostatočne silný právny rámec dôveryhodnej umelej inteligencie, ktorý podľa vzoru GDPR ovplyvní veľkú časť sveta.

Minimálne etické princípy, zásady a hodnoty, ktoré by mali dôveryhodné technológie umelej inteligencie rešpektovať, uvádzame v nasledujúcej kapitole 4.3.3.

Ak odmyslíme všetky spoločenské problémy, ideologické a hodnotové rozdiely i legislatívne zápasy podstupované na poli umelej inteligencie, najproblematickejšou oblasťou pre dosiahnutie dôveryhodných systémov AI sa môže javiť robustnosť, ktorú môžeme zjednodušene vyjadriť ako schopnosť bezpečne a spoľahlivo pracovať, resp. fungovať za akýchkoľvek podmienok.

Dokonale robustný systém neexistuje a v kontexte súčasných moderných technológií si ani nevieme predstaviť jeho realizáciu v blízkej budúcnosti. Avšak na prevádzke jadrových zariadení, letoch do vesmíru, fungovaní viacerých oblastí kritickej infraštruktúry štátu, zabezpečení životných funkcií ľudí odkázaný na prístroje, atď. vidíme, že **vieme realizovať robustné systémy s primeranou mierou rizika.**

Minimálnu hranicu pre podmienku robustnosti dôveryhodného systému by sme mohli stanoviť takto:

- robustné systémy musia spĺňať vysoké technologické štandardy a normy, ktoré sú

pre jednotlivé oblasti nasadenia záväzné.

- robustnosť je chápaná ako neustále prebiehajúci proces, v rámci ktorého sa schopnosti bezpečne a spoľahlivo pracovať vylepšujú, pričom sa existujúce chyby priebežne odhaľujú a opravujú.
- vývoj a prevádzka robustných systémov prebieha v rámci noriem riadenia kvality, manažmentu rizík a nahlasovania, serióznej analýzy a riešenia incidentov.

Treba zdôrazniť, že obsahom uvedených troch bodov je celý súbor opatrení a nariadení, pričom väčšina z nich je už implementovaná v európskom Akte o umelej inteligencii v časti tzv. významných povinností poskytovateľov, ktoré sme uviedli v kapitole 3.3. Ide o riadenie rizík, technickú dokumentáciu, interné záznamy (logovanie), technologickú transparentnosť, možnosť ľudského dohľadu a kvalifikovaného zásahu, prevádzkovú presnosť, spoľahlivosť a kybernetickú bezpečnosť, systém riadenia kvality, viaceré aspekty certifikácie (posúdenie zhody, registrácia, post-marketing monitoring) a oznamovanie incidentov.

Nielen podľa pripravovaného nariadenia EÚ, ale z podstaty veci sa k nim radí aj tzv. správa dát, požadujúca, aby testovacie a tréningové datasety boli „vysokej kvality“, aby tak systém AI, ktorý ich využíva, nebol diskriminačný a nevytváral nepredvídané, resp. nesprávne výsledky. Správa dát je však viacrozmernej problém, ktorý zasahuje do oblasti legálnosti, etiky i robustnosti dôveryhodných systémov AI.

Bez zodpovedného stanovenia a splnenia minimálnych požiadaviek v oblasti legálnosti, etiky a robustnosti systémov umelej inteligencie nie je možné hovoriť o dôveryhodnej umelej inteligencii.

4.3.3. Etické požiadavky na dôveryhodné systémy umelej inteligencie

V práci sme doteraz akcentovali viacero etických výziev, mnohokrát deklarovali etické zásady v rámci diskusie konkrétnej problematiky umelej inteligencie a sumarizovali aktuálne dianie v tejto oblasti. Na tomto pomerne širokom základe by sme chceli postaviť – podľa nás **zásadné – etické požiadavky na dôveryhodné systémy umelej inteligencie zameranej na človeka:**

- **pri vývoji, výrobe, nasadení, poskytovaní a používaní systémov umelej inteligencie musí byť zaručená ochrana slobody, dôstojnosti a bezpečia každej ľudskej osoby i celej spoločnosti.**

- **technológie umelej inteligencie musia byť plne pod ľudskou kontrolou a ovládateľné človekom.**
- **algoritmy i výsledky činnosti systémov AI musia byť človekom pochopiteľné a revidovateľné.**
- **akékoľvek nasadenie technológií AI musí byť prospešné pre človeka a spoločnosť.**
- **systémy umelej inteligencie nesmú byť nástrojom digitálneho rozdelenia.**
- **technológie umelej inteligencie nesmú škodiť nášmu spoločnému domu a mali by prispievať k spoločenskému a environmentálnemu blahobytu.**

Uznávame, že bez solídnej orientácie v problematike uvedené zásady nedokážeme dobre uchopiť a správne aplikovať, no považujeme ich za univerzálnejšie a viac principiálne, než konkrétne zásady, ktoré sme v rámci viacerých etických iniciatív predstavili:

renAlssance 2020 ⁴⁷⁶	Ethics guidelines for trustworthy AI ⁴⁷⁷	The global landscape of AI ethics guidelines ⁴⁷⁸	IEEE Global Initiative on Ethics of Autonomous and Intelligent Systems ⁴⁷⁹
transparentnosť	transparentnosť	transparentnosť	transparentnosť
zodpovednosť	zodpovednosť	zodpovednosť	zodpovednosť
inklúzia	rozmanitosť (diversity), nediskriminácia a spravodlivosť	spravodlivosť a férovosť	ľudské práva
nestrannosť			kompetencie
spoľahlivosť	technická odolnosť	neškodnosť	účinnosť

476 *Rome Call for AI Ethics (document)*. [on-line]. [cit. 28. marca 2022]. Dostupné na internete: <https://www.romecall.org/wp-content/uploads/2022/03/RomeCall_Paper_web.pdf>

477 *Ethics guidelines for trustworthy AI*. [on-line]. [cit. 28. marca 2022]. Dostupné na internete: <<https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>>

478 *The global landscape of AI ethics guidelines*. [on-line]. [cit. 28. marca 2022]. Dostupné na internete: <<https://doi.org/10.1038/s42256-019-0088-2>>

479 *IEEE Global Initiative on Ethics of Autonomous and Intelligent Systems*. [on-line]. [cit. 31. marca 2022]. Dostupné na internete: <<https://standards.ieee.org/industry-connections/ec/autonomous-systems/>>

	a bezpečnosť		
bezpečnosť a súkromie	bezpečnosť, ochrana súkromia a správa údajov	neškodnosť, súkromie	dátová agenda, informovanosť o zneužití
	spoločenský a environmentálny blahobyt		blahobyt

Ak porovnáme štyri rôzne prístupy – závery vatikánskej konferencie renAissance 2020 (Rome Call for Ethics), etické usmernenia skupiny expertov EÚ na umelú inteligenciu, sumarizované etické odporúčania z akademického sveta a etický základ pre tvorbu systémov umelej inteligencie od najväčšej organizácie zastrešujúcej moderné technologické normy a štandardy, pozorujeme rozmanitosť uchopenia základných etických pravidiel pre technológie umelej inteligencie. I keď sa v mnohom tematicky prelínajú, badať určité rozdiely prameniace z filozofických, sociologických, technologických, hodnotových a možno i ideologických predpokladov, na ktorých stavajú.

Máme za to, že náš prístup je univerzálnejší a môže byť nielen základom pre ktorúkoľvek z uvedených zásad jednotlivých iniciatív, resp. inštitúcií, ale dokáže i nadčasovo uchopiť budúci vývoj spoločnosti v konfrontácii a koexistencii s neustále sa rozvíjajúcim a technologicky napredujúcim fenoménom umelej inteligencie.

Realizácia (vývoj, výroba, nasadenie, poskytovanie, využívanie,...) systémov umelej inteligencie musí spĺňať:

- principiálny základ dôveryhodného systému umelej inteligencie zameraného na človeka.
- aplikáciu základných etických zásad vo všetkých rozmeroch, ktoré daný systém AI môže obnášať.
- splnenie legislatívnych požiadaviek.
- robustnú realizáciu a využívanie.

Aplikácia etických zásad a legislatívnych požiadaviek v reálnom svete by mala byť v rovine konkrétnych a jasných odporúčaní. Určite by sme sa voči našim zásadám

neprehrešili, ak by sme stavali na kombinácii etických odporúčaní vatikánskej konferencie renAIssance 2020⁴⁸⁰ a obsahu európskeho nariadenia Akt o umelej inteligencii.⁴⁸¹

4.3.4. Oblasti implementácie etických princípov

V prehľade etických výziev, prameniach z limitov a rizík súčasných systémov umelej inteligencie, ktoré sme priebežne uvádzali v 2. kapitole a sumarizovali v kapitole 3.1., nachádzame **tri oblasti nutnej implementácie etických noriem, eticko-právnych regulácií a morálnych zásad:**

- etické normy, zákonné regulácie a morálne zásady **tvorcov** systémov AI.
- etické normy, zákonné regulácie a morálne zásady **poskytovateľov a používateľov** týchto systémov.
- implementované eticko-právne požiadavky a obmedzenia priamo **v systémoch AI**.

Etický diskurz 4. kapitoly sa stále zameriava na technológie slabej umelej inteligencie (ANI)⁴⁸², u ktorých z podstaty veci musíme rátať so zaangažovaním ľudského faktora vo všetkých oblastiach tvorby, používania i realizácie prostriedkov umelej inteligencie. Preto opakujeme jeden z našich záverov z kapitoly 3.1.: **aplikovanie etických princípov a regulácií v ANI – akokoľvek náročným až neuskutočniteľným by sa to v praxi mohlo ukázať – má vďaka prítomnosti ľudského faktora jasné zásady, vyhranené oblasti a viac menej presne definované kritériá.**

Treba preto investovať nemalé úsilie do vzdelávania, osvetu i prevencie a formovať

480 *Rome Call for AI Ethics (document)*. [on-line]. [cit. 28. marca 2022].

Dostupné na internete:

<https://www.romecall.org/wp-content/uploads/2022/03/RomeCall_Paper_web.pdf>

481 *Nariadenie Európskeho parlamentu a Rady, ktorým sa stanovujú harmonizované pravidlá v oblasti umelej inteligencie (Akt o umelej inteligencii) a menia niektoré legislatívne akty únie*. [on-line]. [cit. 24. marca 2022].

Dostupné na internete: <<https://eur-lex.europa.eu/legal-content/SK/TXT/?uri=CELEX:52021PC0206>>

482 Stále sa nachádzame v oblasti úzko špecializovaných systémoch 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 inteligentné správanie na základe modelov a aplikovaných metód i tréningových dát. Sú to teda systémy zamerané na riešenie konkrétnych úloh a sú závislé na ľudskom vstupe a konfigurácii.

morálne postoje vývojárov, poskytovateľov i používateľov týchto technológií.⁴⁸³

A ako sme v kapitole 3.1. poznamenali, ruka v ruke s formáciou a vzdelávaním odborníkov v oblasti umelej inteligencie by sa mala rozvíjať aj **edukácia a osвета spoločnosti**, bez ktorej si ťažko predstaviť rast spoločenskej citlivosti a zodpovednosti v oblasti celoplošnej adaptácie a využívania systémov AI.

Osobitnou oblasťou je **základný výskum umelej inteligencie**, ktorý už z princípu nemôže byť mnohými reguláciami viazaný. I v tejto oblasti by však mali existovať etické garancie, ktorých cieľom je pri akomkoľvek type výskumu rešpektovať princíp zamerania na človeka a s tým súvisiacu dôveryhodnosť AI. Teda **Human-centered AI a Trustworthy AI považujeme za nutné princípy akéhokoľvek výskumu, vývoja, realizácie alebo nasadenia systémov umelej inteligencie.**

V kapitole 3.3. sme na margo technologických požiadaviek európskeho Aktu o umelej inteligencii diskutovali praktickú a v niektorých scenároch principiálnu nemožnosť splnenia eticko-právnych podmienok v technologickej rovine v prípade neskoršieho doplnenia technických úprav alebo procesných postupov do existujúcich systémov AI.

Ak v prípade kybernetickej bezpečnosti (NIS) alebo ochrany osobných údajov (GDPR) patria princípy *security by design* a *privacy by design* k tzv. osvedčeným a odporúčaným postupom (best practices), **v prípade systémov AI je *ethics by design* a *regulation by design* nutnou podmienkou ich vývoja, nasadenia a prevádzky.**⁴⁸⁴ Etické zásady a právne normy tak musia byť integrálnou súčasťou prostriedkov umelej inteligencie už od ich technologického návrhu – neskorším doplnením sa v plnej miere a požadovanom rozsahu vo väčšine systémov AI nebudú dať spoľahlivo aplikovať.

Vzhľadom na autonómny a adaptívny rozmer technológií umelej inteligencie sa splnením princípov *ethics by design* a *regulation by design* kladie základ pre tzv. odolnosť týchto systémov voči budúcnosti (future-proof systems), t.j. schopnosť v budúcnosti spoľahlivo a bezpečne pracovať aj napriek možným nepredvídateľným okolnostiam.

483 Napr. v predchádzajúcej kapitole spomínaná iniciatíva *IEEE Global Initiative on Ethics of Autonomous and Intelligent Systems* zahŕňa i osobitné zameranie na etiku vývoja a vývojárov systémov AI.

484 Regulation by design v sebe obnáša aj security by design a privacy by design. Spolu s ethics by design tak tvoria celkový rámec trustworthy by design.

4.4. Špecifické odporúčania

Popisujúc súčasné legislatívne aktivity v kapitole 3.3. sme s uspokojením konštatovali, že európsky Akt o umelej inteligencii má potenciál pokryť využívanie technológií umelej inteligencie prakticky vo všetkých oblastiach reálneho života spoločnosti. Predsa sme však spomenuli oblasti, v ktorých využitie algoritmov AI vybočuje z bežného rámca nasadenia. Ide o nasadenie systémov AI v oblasti pokročilého riadenia štátu, spravodajstva a plošného dohľadu a nasadenie vo vojenskej oblasti.⁴⁸⁵

Keďže ide o oblasti s veľkými právomocami, dosahom i rizikom zneužitia pre dobro človeka a celej spoločnosti, radi by sme upresnili, že do ich osobitnej regulácie by mala byť zahrnutá aj oblasť etických princípov a morálnych aspektov využívania.⁴⁸⁶

Odporúčania, ktoré uvádzame v tejto kapitole, by mali byť doplnkom k všeobecným návrhom a základným princípom, ktoré sme uviedli v kapitole predchádzajúcej.

4.4.1. Oblasť pokročilého riadenia štátu, spravodajstva a plošného dohľadu

Je pravdou, že niektoré z technológií algoritmického riadenia a plošného dohľadu pokrýva pripravované európske nariadenie Akt o umelej inteligencii v časti Zakázané systémy (čl. 5), no ide len o parciálne riešenie, ktoré je navyše sprevádzané možnými tzv. oprávnenými výnimkami. Ich legislatívne usmernenie musí byť preto pokryté inými zákonnými normami.

Opakujeme preto požiadavku, ktorú sme uviedli v kapitole 3.3.: **vzhľadom na špecifikum a dosah nasadenia systémov AI v oblasti pokročilého riadenia štátu, spravodajských služieb a plošného dohľadu⁴⁸⁷ pre ľudské práva, ochranu**

485 Uvedomujeme si, že vzhľadom na osobitné riziká a výzvy existuje viacero oblastí so špecifickými etickými odporúčaniami a právnymi normami, napr. oblasť autonómnych vozidiel s ochranou ľudských životov, zabráneniu škodám a prisúdeniu zodpovednosti, bankový sektor s otázkou sociálnej spravodlivosti a transparentnosti, spoločnosť a mediálny svet s rizikami kapitalizmu dohľadu a manipulovania s ľuďmi, samostatné časti algoritmického riadenia so spravodlivosťou, transparentnosťou a demokratickými princípmi, prípadne rôznorodé kombinácie uvedených špecifik, napr. v zdravotníctve a pod.

Všetky tieto oblasti však dostatočne pokrýva navrhnutý Akt o umelej inteligencii v kombinácii s prípadnými špeciálnymi sektorovými reguláciami.

486 Etické princípy a morálne zásady v tejto oblasti ťažko odporúčať, treba ich vyžadovať.

487 Viaceré aspekty nasadenia systémov AI v oblasti spravodajstva a algokracie sme diskutovali v kapitole 2.6.

demokracie a slobôd si myslíme, že táto oblasť by mala byť pokrytá už základnými legislatívnymi mechanizmami a verejným dohľadom demokratickej spoločnosti, ktoré sa týkajú riadenia spoločnosti, ľudských práv a pôsobenia spravodajských služieb vo všeobecnosti.

Nie je jednoduché realizovať a vyžadovať striktnú zákonnosť i dôsledný dohľad demokraticky zvolených zástupcov a spoločnosti pri nasadení systémov AI v rámci spravodajských služieb, dohľadových systémov a všetkých foriem i stupňov algoritmického riadenia spoločnosti. Treba preto v celom spektre nasadenia od spravodajských služieb až po algoritmické riadenie akcentovať veľkú rozvážnosť a vyžadovať striktnú legálnosť i dôsledný dohľad demokraticky zvolených zástupcov, obmedziť dopady na sociálnu spravodlivosť a zabezpečiť dodržiavanie ľudských práv a hodnôt. Ide nielen o prijaté zákony, ale i nastavené procesy kontroly a mechanizmy zásahov v prípade podozrení na zlyhanie, či zneužitie činnosti technológií umelej inteligencie.

Osobitnou kapitolou je distribúcia pokročilých systémov umelej inteligencie do rizikových krajín. Podľa štúdie *The Global Expansion of AI Surveillance*, ktorú sme rozoberali v kapitole 2.6., je jednou zo znepokojujúcich skutočností fakt, že demokratické štáty, z ktorých väčšina sofistikovaných technológií algoritmického riadenia, spravodajstva a dohľadu pochádza, neprijímajú primerané opatrenia na monitorovanie a kontrolu šírenia týchto technológií, majúcich potenciál sa podieľať na celom rade možných porušení ľudských práv a zneužití autokratickými režimami a diktatúrami. Navyše je za ostatné desaťročie známych a zdokumentovaných viacero prípadov, keď sa technologické spoločnosti z demokratických krajín priamo podieľali na nasadení svojich technológií v scenároch obmedzujúcich ľudské práva, slobodu, súkromie a demokraciu.⁴⁸⁸ Len pomaly rastie celospoločenské povedomie a angažovanie pri obmedzovaní exportu do rizikových krajín.

Navrhujeme, aby oblasť exportu produktov a technológií umelej inteligencie, ktoré môžu byť zneužívané v oblasti pokročilého riadenia štátu, spravodajstva a plošného dohľadu, bola predmetom medzinárodnej regulácie s cieľom zamedziť ich vývoz do rizikových krajín. Sankcie by však nemali byť nástrojom (geo) politických zápasov, ale skutočného nasadenia na poli etického využívania systémov AI.

488 Napr. Microsoft, Cisco a ďalší pri sledovacích systémoch v Číne....

4.4.2. Systémy AI vo vojenskej oblasti

V kapitole 3.1. sme ako jeden zo záverov rizík nasadenia technológií umelej inteligencie vo vojenskej oblasti uvádzali: „Schopnosti umelou inteligenciou poháňaných autonómnych zbraňových systémov a kybernetických zbraní i napriek všetkým rizikám vedú k zvyšujúcim sa tlakom na financovanie a zavádzanie útočných kybernetických zbraní. Jednotlivé krajiny sa nedokážu vzdať tak lákavej technologickej výhody a sú pevne rozhodnuté technológie umelej inteligencie implementovať v celej šírke možného zmysluplného využitia. Problematika obmedzenia týchto útočných systémov je skomplikovaná aj reálnym stieraním hraníc medzi obranným a útočným nasadením takmer vo všetkých oblastiach vojenského využitia technológií umelej inteligencie.“

Túto skutočnosť sme následne v kontexte prijímaných regulácií v kapitole 3.3. komentovali slovami: „Keďže akékoľvek obmedzovanie technológií umelej inteligencie vo vojenskej oblasti môže byť chápané ako bezpečnostné riziko a zníženie bojaschopnosti modernej armády, jednostranne prijaté regulácie nemusia byť účinné – nielen pre to, že sa jednou stranou ťažko prijímajú (aj keď pre hodnotovo orientovanú spoločnosť by to malo byť povinnosťou), ale i pre malú šancu na ich extra teritoriálne rozšírenie a akceptovanie.“⁴⁸⁹

Mnohí štátni aktéri by chceli aplikovať latinské *si vis pacem, para bellum* i na oblasť armádneho nasadenia technológií AI v zbraňových systémoch a ich potenciálom upevňovať svoje postavenie. Pre ďalších – pozerajúc k horizontu možností autonómnych zbraňových systémov – by sa ich zavádzanie do výzbroje mohlo podobat' odstrašujúcemu potenciálu jadrových zbraní. Ak však pozeráme za horizont súčasných možností vojenských systémov AI, vidíme technológie, ktorých rizikový potenciál prekračuje nebezpečenstvo prameniace z terajších arzenálov jadrových zbraní.

Rizikám a limitom systémov umelej inteligencie vo vojenskom nasadení a v zbraňových systémoch sme sa obsérne venovali v kapitole 2.7., pričom ich etické dôsledky sme diskutovali v kapitolách 3.1. a 3.2.

V otvorenom liste *Autonomous Weapons: An Open Letter from AI & Robotics Researchers*, ktorý sme spomínali v kapitolách 2.7.8. a 3.2. a ktorý vyzýval na zákaz autonómnych zbraňových systémov, sú tieto systémy označované ako tretia revolúcia v zbraňových systémoch, po strelnom prachu a jadrových zbraniach. Ide o posun nielen v technológii, ale – a to je pre armádne nasadenie dôležité – predovšetkým o posun

⁴⁸⁹ Kapitola 3.3.

v schopnostiach a účinnosti týchto zbraní, ktorých plné nasadenie by mohlo viesť až k nukleárnej, či jej podobnej katastrofe.

Nadväzujúc tak na Russell-Einsteinov manifest, varujúci pred reálne možnými rizikami jadrového zbrojenia, pripájame sa k všetkým výzvam, ktoré volajú po zákaze smrtiacich autonómnych zbraňových systémov. Či už išlo o *Scientists' Call to Ban Autonomous Lethal Robots* z roku 2013, spomínaný *Autonomous Weapons: An Open Letter from AI & Robotics Researchers* z roku 2015, v roku 2017 vydaný *An Open Letter to the United Nations: Convention on Certain Conventional Weapons*, alebo európske *Resolution on autonomous weapon systems* z roku 2018.⁴⁹⁰

Principiálny postoj v oblasti smrtiacich autonómnych zbraňových systémov (LAWs) je jasný: **technológiami umelej inteligencie poháňané automatické smrtiace zbraňové systémy, systémy automatického zameriavania a vyberania cieľov, automatické systémy schopné bez zásahu človeka rozhodnúť o smrtiacej reakcii akéhokoľvek druhu (od útoku dronu až po rozpútanie jadrového konfliktu) musia byť zakázané.**

Rozoberajúc limity a riziká súčasných technológií umelej inteligencie,⁴⁹¹ ani pri najlepšej vôli nie sme schopní vytvoriť také systémy ANI, ktorým by sme mohli zveriť samostatné rozhodovanie v tak dôležitej oblasti. A pokiaľ ide o AGI – znovu sa opakujúc – ak nemáme vyriešené otázky ako napr. funkčný model správania, implementáciu komparatívnych hodnotových rámcov a pod., nie sme schopní o tejto možnosti ani len špekulatívne uvažovať.

Vývoj a nasadenie viacerých extrémne nebezpečných vojenských technológií je v súčasnosti na základe vzájomného konsenzu a právnych záväzkov zakázané. Je preto žiadúce, aby sa tak stalo aj v prípade technológiami umelej inteligencie poháňaných plne automatizovaných smrtiacich zbraňových systémov, systémov automatického zameriavania a vyberania cieľov i automatických systémov schopných bez zásahu človeka rozhodnúť o smrtiacej reakcii akéhokoľvek druhu.

Od plne autonómnych zbraňových systémov bez kontroly človeka sa však vráťme k tomu, čím sme túto kapitolu začali – strategickému i praktickému zavádzaniu takých technológií AI do armádnych systémov, ktoré by mali byť pod kontrolou človeka.

490 Všetky uvedené iniciatívy sme diskutovali v kapitole 3.2.

491 Znovu sa odvolávame na kapitoly 2.1. až 2.4.

V kapitolách 2.7.8. a 3.2. spomenuté eticko-právne procesy, ktoré prebiehajú napríklad v armáde USA, môžu byť základom pre celosvetovú diskusiu mocností a na základe tlaku verejnosti, angažovanosti jednotlivých častí spoločnosti v rôznych regiónoch sveta i úsilia zodpovedných strán môžu viesť k prijatiu celosvetových pravidiel i záväzkov pre oblasť vývoja a nasadenia rizikových vojenských systémov vybavených technológiami umelej inteligencie.

Ako základ môže poslúžiť v kapitole 2.7.8. diskutovaný dokument *AI Principles: Recommendations on the Ethical Use of Artificial Intelligence* by the Department of Defense z Defense Innovation Board USA s podstatnými závermi, ku ktorým patrí zodpovednosť (rozhodovacia právomoc), opatrnosť pri príprave testovacích dát a návrhu systému, dosledovateľnosť, spoľahlivosť a ovládateľnosť.⁴⁹²

Okrem väčšiny všeobecných požiadaviek, ktoré sme v kapitole 4.3. uvádzali a ktoré je možné vo vojenskom využití aplikovať, mali by v prípade autonómnych zbraňových systémov kontrolovateľných človekom platiť nasledovné zásady:

- **nutnou podmienkou prevádzky ľubovoľného systému AI, ktorý môže predstavovať riziko pre akúkoľvek ľudskú osobu, je schopnosť a možnosť človeka prebrať kedykoľvek kontrolu nad týmto systémom, resp. právo a možnosť verifikovať a prehodnotiť výsledky jeho činnosti.**
- **limity, regulácia a obmedzenia LAWs by mali predstavovať etický rámec stanovený na základe morálnych hodnôt ľudskej spoločnosti, nie na základe relativistickej tzv. „následnej regulácie“.**⁴⁹³

Vráťme sa k rôznorodosti štátnych aktérov a armádneho zápolenia, ktoré sme v úvode tejto kapitoly spomínali a ktoré pomerne efektívne bráni v možnom jednostrannom, resp. dobrovoľnom obmedzovaní zavádzania rizikových zbraňových systémov AI.

Uviedli sme, že napriek zníženiu konkurenčnej schopnosti a malej šanci na akceptovanie inými štátmi, by jednostranne prijaté (vyššie uvedené) eticko-právne regulácie mali byť pre hodnotovo orientovanú spoločnosť povinnosťou. Pozývame k tomuto kroku z viacerých dôvodov:

492 Vidíme, že ide o dosť oklieštený regulačný rámec oproti iniciatívam, ktoré sme rozoberali v kapitole 4.3.

493 Fenomén „následnej regulácie“, ktorá smeruje k hodnotovému a morálnemu relativizmu a snaží sa prehodnotiť argumenty o nenahradiiteľnosti ľudského svedomia a morálneho úsudku, sme diskutovali v kapitole 2.7.8.

- žiadny štát, pokiaľ sa zriekne etických princípov a morálnych zásad, nemá právo obhájiť svoju účasť na vojnovom konflikte a zvíťaziť.
- **pri nasadení moderných zbraňových systémov s celoplošnými účinkami a technológií, zasahujúcich v hybridných vojnách a vojnách 4. generácie prakticky celé populácie štátov, sa koncept spravodlivej vojny stáva neprijateľný.**
- len na základe konkrétne prijatých záväzkov sa môže celosvetová diskusia mocností, tlak verejnosti, angažovanosť jednotlivých častí spoločnosti v rôznych regiónoch sveta a úsilie zodpovedných strán pretaviť v postupné prijatie celosvetových pravidiel i záväzkov pre oblasť vývoja a nasadenia rizikových vojenských systémov vybavených technológiami umelej inteligencie.

V čase písania tejto kapitoly nás používanie viacerých sofistikovaných útočných bojových systémov v prebiehajúcom vojnovom konflikte na Ukrajine stavia pred dilemu, či budeme technológie umelej inteligencie využívať v duchu hesla *si vis pacem, para bellum* vo falošnej predstave, že vďaka nim sa vyhneme vojnovým konfliktom, pričom však budeme mať tendenciu ich v potencionálnych konfliktoch aj použiť.

Alebo v intenciách *si vis pacem para pacem* – vyvarujúc sa pacifistickému nevyužitiu ich potenciálu – dokážeme byť dostatočne zrelí na ich obranné využitie a nasadenie pre ochranu života a ľudskej dôstojnosti i zabezpečenie skutočných hodnôt spoločnosti.

4.5. Priestor pre angažovanie sa Cirkvi

4.5.1. Misia zjednocovať, usmerňovať a propagovať etické aktivity

V kapitole 3.4. sme sa v kontexte štúdie *The global landscape of AI ethics guidelines* pozastavovali nad celkovou hodnotovou a etickou diskrepanciou vo svete: v zásade panuje celospoločenský konsenzus ohľadom potreby etických usmernení v oblasti umelej inteligencie, tento je však sprevádzaný podstatnými rozdielmi v chápaní etiky a hodnôt.⁴⁹⁴

Pripomeňme si preto slová Mons. Vincenza Pagliu na margo rímskej konferencie renAIssance 2020: „Zámernom výzvy je vytvoriť hnutie, ktoré sa rozšíri a zapojí ďalších

494 Por. JOBIN, IENCA, VAYENA, *The global landscape of AI ethics guidelines*. [on-line]. [cit. 28. marca 2022].

Dostupné na internete: <<https://doi.org/10.1038/s42256-019-0088-2>>

aktérov: verejné inštitúcie, mimovládne organizácie, priemyselné odvetvia a skupiny, aby určili smer vývoja a používania technológií odvodených od umelej inteligencie. Z tohto hľadiska môžeme povedať, že prvý podpis tejto výzvy nie je vyvrcholením, ale východiskom pre záväzok, ktorý sa javí ako ešte naliehavejší a dôležitejší než kedykoľvek predtým. Pripojenie sa k tejto iniciatíve pre zástupcov priemyslu, ktorí ju podpíšu, znamená záväzok, ktorý má význam aj z hľadiska nákladov a technologického príspevku k vývoju a distribúcii ich výrobkov. Ak sa Akadémia cíti byť povolaná zintenzívniť svoje úsilie o uľahčenie získavania poznatkov a podpisov iných medzinárodných aktérov, táto výzva je len prvým krokom, po ktorom by mali nasledovať ďalšie. Text výzvy sa vyznačuje aj tým, že je prvým pokusom sformulovať súbor etických kritérií so spoločnými referenčnými bodmi a hodnotami, čím ponúka príspevok k rozvoju spoločného jazyka na interpretáciu toho, čo je človek“.⁴⁹⁵

Katolícka cirkev má v rámci svojho pôsobenia a univerzálneho spoločenstva výnimočnú možnosť **iniciovať a rozvíjať celosvetové širokospektrálne hnutie s cieľom etického vývoja a používania technológií umelej inteligencie.**

Na základe Zjavenia, solídneho teologického aparátu a interdisciplinárnych skúseností má potenciál **identifikovať, uchopiť, rozpracovať a formulovať základný univerzálny súbor etických kritérií a hodnôt** ako odpoveď na etický a hodnotový relativizmus.⁴⁹⁶

Vďaka svojej angažovanosti od diplomatických misií až po periférie sveta má možnosť **ponúkať a inkulturovať spoločné referenčné body a hodnoty základných etických kritérií do rôznorodosti ľudského spoločenstva a jednotlivých oblastí nasadenia systémov AI v spoločnosti.**

Predstavitelia Cirkvi môžu tiež pôsobiť ako **neutrálni sprostredkovatelia v špecifických oblastiach**, napr. v rámci jednaní o obmedzení vojenského nasadenia technológií umelej inteligencie a pod.

⁴⁹⁵ *Rome Call for AI Ethics*. [on-line]. [cit. 19. augusta 2020].

Dostupné na internete: <<http://www.academyforlife.va/content/pav/en/events/intelligenza-artificiale.html>>

⁴⁹⁶ Vid' napr. „Etické rozhodovanie vo všeobecnosti neprináša absolútne správne alebo absolútne nesprávne riešenia daných situácií. Inými slovami, rozhodnutie, ktoré istá skupina ľudí považuje za správne, resp. v ich ponímaní etické, nemusí automaticky predstavovať správne (etické) rozhodnutie v očiach iných ľudí.“

ŠTARHA, GAŠPAROVIČ, *AI z pohľadu práva*. [on-line]. [cit. 29. marca 2022].

Dostupné na internete: <<https://www.epravo.sk/top/clanky/ai-z-pohladu-prava-4483.html>>

Keďže vplyv a dôsledky využívania AI budú viac a viac zasahovať prakticky do všetkých oblastí života dnešného sveta, musí byť táto oblasť z pohľadu etického využívania predmetom neustálej osvety a angažovanosti zo strany Cirkvi – podobne, ako sa deje v prípade viacerých dôležitých aktivít na poli OSN, medzinárodnej úrovni i každodennej práce Cirkvi v dnešnej spoločnosti.

4.5.2. Akcent na univerzálne bratstvo a sociálne priateľstvo

V duchu encykliky pápeža Františka *Fratelli tutti* – o bratstve a sociálnom priateľstve je Cirkev pozvaná prispieť k využívaniu technológií umelej inteligencie pre dobro celého ľudského spoločenstva na zemi:

„...moja kritika technokratickej paradigmy neznamená, že môžeme zostať v bezpečí, len ak sa budeme usilovať kontrolovať jej excesy. Najväčšie nebezpečenstvo nespočíva často vo veciach, v materiálnych skutočnostiach, organizáciách, ale v spôsobe, akým ich ľudia používajú.“⁴⁹⁷

Myšlienka budovať univerzálne bratstvo a sociálne priateľstvo, ktorá rezonovala naprieč tzv. oblasťami vplyvu definovanými v záveroch rímskej konferencie renaissance 2020, je mnohorakým spôsobom uchopená a rozvinutá v encyklike *Fratelli tutti*.

V rámci týchto dokumentov ako príklad uvádzame niekoľko výziev, v ktorých – okrem všetkého na poli etiky doteraz uvedeného – má Cirkev potenciál predstavovať fenomén umelej inteligencie ako nástroj účasti na evanjeliovom budovaní univerzálneho bratstva a sociálneho priateľstva:

- premostenie k marginalizovaným (spomínaný *bridge the gap* a zmierňovanie *digital divide*).
- nenechávať nikoho za sebou (*nemo resideo, no one left behind*).
- vedomie, že ak chceme budovať nejakú budúcnosť, musíme ju budovať všetci, inak sa to nedá.

Sme pozvaní podstúpiť zápas o pravú tvár etiky umelej inteligencie – aby sa z nej nestal nástroj moderných ideológií, sociálnej nespravodlivosti, obmedzovania ľudských práv, digitálneho rozdelenia, erózie hodnôt a ohrozenia ľudského bytia i celého

497 PÁPEŽ FRANTIŠEK. *Fratelli tutti*. [on-line]. Čl. 166.[cit. 4. apríla 2022].

Dostupné na internete: <<https://www.kbs.sk/obsah/sekcia/h/dokumenty-a-vyhlasenia/p/dokumenty-papezov/c/encyklika-fratelii-tutti>>

spoločenstva – ale aby bola dobrým sluhom, prostriedkom budovania univerzálneho bratstva a sociálneho priateľstva.⁴⁹⁸

498 Ide o jeden zo zámerov vyjadrených v úvode tejto práce – technológie umelej inteligencie v úlohe dobrého sluha, nie zlého pána.

5. Vízia silnej a všeobecnej umelej inteligencie

„Rozhodli se svěřit naše bezpečí něčemu, co není schopné věrnosti,
morálky ani moudrosti.“⁴⁹⁹

Akokoľvek je táto práca primárne zameraná na súčasné technológie umelej inteligencie, ktoré sa napriek mnohorakej mediálnej nadsádzke nedajú nazvať inak než len slabými systémami ANI, prakticky pri všetkých doteraz rozoberaných témach sme sa nevyhli aspoň krátkemu pohľadu za horizont – k systémom silnej a uvedomelej umelej inteligencie, ktorú síce mnohí s nádejou očakávajú, no iní sa jej obávajú. Pritom však všetci stoja pred dilemou, ako ju uchopiť...

Protagonisti Dartmouthského seminára, ktorý znamenal zrod fenoménu umelej inteligencie,⁵⁰⁰ si kládli za cieľ „pokúsiť sa zistiť, ako prinútiť stroje používať jazyk, vytvárať abstrakcie a pojmy, riešiť druhy problémov, ktoré sú teraz vyhradené pre ľudí, a zlepšovať sa.“⁵⁰¹

Účastníci seminára mali veľký entuziazmus a optimizmus ohľadom dosiahnutia pokročilej umelej inteligencie, čo John McCarthy vyjadril slovami: „myslíme si, že významný pokrok v riešení jedného alebo viacerých z týchto problémov sa dá dosiahnuť, ak starostlivo vybraná skupina vedcov bude na tom spoločne pracovať počas jedného leta“.⁵⁰²

Vôbec sa netajili zámerom svojej vízie a snahy zamerať na realizáciu uvedomelej umelej inteligencie.

Postupom času – napriek všetkému vynaloženému úsiliu – prichádza vytriezvenie. McCarthy na začiatku šesťdesiatych rokov zakladá Standfordský projekt umelej inteligencie „s cieľom vytvorenia plne inteligentného stroja v rámci dekády“.⁵⁰³ Budúci

499 CAMPBELL, J. *Dvojník*, Fantom Print 2015, ISBN 978-80-7398-329-7, s. 713.

500 Dartmouthský seminár, ktorý sa uskutočnil v r. 1955, sme uvádzali v kapitole 1.1.

501 MITCHELL, M. *Conceptual Abstraction and Analogy in Artificial Intelligence*. In: *ALIFE 2020: The 2020 Conference on Artificial Life*. [on-line]. [cit. 5. apríla 2022].

Dostupné na internete: <https://doi.org/10.1162/isal_a_00354>

502 MCCARTHY et al., *Proposal for the Dartmouth Summer Research Project in Artificial Intelligence*. [on-line]. [cit. 3. februára 2021].

Dostupné na internete: <<https://doi.org/10.1609/aimag.v27i4.1904>>

503 MORAVEC, *Mind Children: The Future of Robot and Human Intelligence*, s. 20.

laureát Nobelovej ceny v tom istom čase predikuje: „Do dvadsať rokov budú stroje schopné vykonávať akúkoľvek prácu, ktorú môže vykonávať človek“.⁵⁰⁴ A Minský obnovuje nádeje, hovoriac: „v rámci jednej generácie budú problémy spojené s vytvorením umelej inteligencie v podstate vyriešené“.⁵⁰⁵

Pomaly dobíjame siedme desaťročie od Dartmouthského seminára a napriek úžasnému pokroku – osobitne v posledných dekádach – žiadna z týchto predpovedí sa doteraz nenaplnila. Výskumné témy, predstavené na seminári, sú stále otvorené a aktívne skúmané v odbornej komunite po celom svete.

Hoci umelá inteligencia dosiahla za posledné desaťročie dramatický (dá sa povedať, že parciálne až exponenciálny) pokrok v oblastiach, ako je videnie, spracovanie prirodzeného jazyka a robotika, súčasným systémom AI stále takmer úplne chýba schopnosť vytvárať pojmy a abstrakcie podobné ľudským, schopnosť porozumieť, chápať súvislosti a myslieť, realizovať niečo na spôsob zdravého rozumu ako predpokladu pre všeobecnú inteligenciu.⁵⁰⁶

5.1. „Trhliny v inteligencii“ pokročilých systémov AI

Méta uvedomelej umelej inteligencie by sa mala dosiahnuť prostredníctvom silnej (strong) a všeobecnej (general) AI. Všeobecnej, lebo dokáže zvládnuť akúkoľvek intelektuálnu úlohu a má schopnosť generalizovať, t.j. zovšeobecňovať a prenášať, či adaptovať naučené schopnosti na iné úlohy. Silnej, pretože aj skutočne rozumie tomu, čo rieši a vykonáva.

Podľa protagonistov uvedomelej AGI by na tomto základe mal byť možný rozvoj umelej

504 SIMON, H. A. *The Shape of Automation for Men and Management*. New York: Harper & Row, 1965, s. 90.

505 MINSKY, M. L. *Computation: Finite and Infinite Machines*. Upper Saddle River, N.J.: Prentice Hall, 1967, s. 2.

506 Gary Marcus, profesor psychológie a špecialista v oblasti umelej inteligencie uvádza: pri vývoji „silnej umelej inteligencie“ – teda všeobecnej umelej inteligencie na úrovni človeka – „nie je takmer žiadny pokrok“.

PRESS, G. *12 Observations About Artificial Intelligence From The O'Reilly AI Conference*. In: *Forbes*. [on-line]. 2016, 31. októbra. [cit. 7. augusta 2020].

Dostupné na internete: <<https://www.forbes.com/sites/gilpress/2016/10/31/12-observations-about-artificial-intelligence-from-the-oreilly-ai-conference/>>

inteligencie, ktorá by bola porovnateľná (ak nie lepšia) s myslou človeka, smerujúc tak k vedomiu a sebauvedomeniu, až po analógiu ľudskej osoby.

Nech je však rozvoj súčasných systémov umelej inteligencie akokoľvek prevratný a dosiahnuté výsledky impozantné, stále narážame na podstatné problémy, bez ktorých sa o skutočnej AGI nedá uvažovať.

Predovšetkým ide o tzv. **bariéru chápania zmyslu** (barrier of meaning) medzi súčasnými systémami AI a inteligenciou na ľudskej úrovni.⁵⁰⁷ Algoritmy ANI podstatným spôsobom nedokážu chápať zmysel, ani sa vyrovnáť ľuďom vo vnímaní, reči a chápaní. Systémy ANI preto konajú chyby diametrálne odlišné od tých ľudských a majú problémy s abstrahovaním⁵⁰⁸ a prenášaním skúseností, resp. toho, čo sa naučili. Technológiám ANI chýba chápanie zdravého rozumu, keďže **miesto myslenia realizujú jeho simuláciu**. Rozdiel medzi myslením a jeho napodobnením v pochopení zmyslu, súvislostí, analógií,⁵⁰⁹ konceptov a mnohých ďalších súvislostí je tak veľký, že súčasné systémy ANI sú náchylné na jednoduché oklamanie na základe reálneho nechápania obsahu, súvislostí a zmyslu.⁵¹⁰

V rámci systémov AI nie sme schopní vytvoriť niektoré základné prvky analogické ľudskému chápaniu, ktoré máme či už vrodené, alebo sa vyvíjajúce v rámci prvých rokov života. Viacerí psychológovia hovoria napríklad o tzv. intuitívnej fyzike – základnom, rýchlo osvojenom a zovšeobecnenom poznaní o objektoch hmotného sveta, intuitívnej biológii –

507 MITCHELL, *Artificial Intelligence*, s. 235.

508 Už v rámci propozícií Dartmouthského seminára sa „formovanie abstrakcií“ uvádza ako jedna z kľúčových schopností umelej inteligencie. Ako však uvádza prof. Melanie Mitchell, schopnosť systémov AI formovať ľudskému mysleniu sa podobajúce konceptuálne abstrakcie zostáva až doteraz takmer vôbec nevyriešená.

509 Niektorí kognitívni vedci navrhli, že tvorba analógií je ústredným mechanizmom pojmovej abstrakcie a porozumenia u ľudí. Douglas Hofstadter, ktorého sme spomínali v úvode tejto práce, nazval tvorbu analógií „jadrom poznávania“, pričom so spoluautorom Emmanuelom Sanderom poznamenáva: „Bez pojmov nemôže byť myslenie a bez analógií nemôžu byť pojmy.“ Schopnosti tvorby analógií sú kľúčové pri vývoji systémov umelej inteligencie s inteligenciou podobnou ľudskej. MITCHELL, *Conceptual Abstraction and Analogy in Artificial Intelligence*. [on-line]. [cit. 5. apríla 2022]. Dostupné na internete: <https://doi.org/10.1162/isal_a_00354>

510 Práve **chyby nepodobné ľudským, ktorých sa systémy AI dopúšťajú a ktoré sme v tejto práci viackrát uvádzali (napr. zraniteľnosti voči tzv. adversarial examples), poukazujú na nepochopenie konceptov nachádzajúcich sa vo vstupných dátach**. Systémy AI sa pri vhodnej voľbe algoritmov, hyperparametrov a tréningových dát dokážu vytrénovať na úzko špecializované úlohy, no nedokážu chápať princípy, koncepty a zmysel toho, čo vykonávajú.

uvedomovaní si rozdielov medzi živými a neživými objektami, alebo o intuitívnej psychológii – sociálnej schopnosti vnímať pocity, postoje a zámery iných osôb. U človeka pritom ide len o fundamentálnu bázu poznatkov, na základe ktorých sa rozvíjajú kognitívne schopnosti, rôznorodé aspekty učenia sa a myslenia, ku ktorým patrí schopnosť spoznávať nové koncepty len z niekoľkých príkladov, schopnosť ich zovšeobecňovať a aplikovať na iné situácie, rýchlo chápať zmysel reálnych situácií a na ich základe sa prakticky okamžite, adekvátne a správne rozhodovať.⁵¹¹

To, čo je v rámci intuitívnej fyziky, biológie a psychológie bežnou súčasťou prvých rokov rozvoja ľudskej osobnosti, dokážu v súčasnosti zvládnuť najlepšie systémy AI len vo veľmi obmedzenom a neúplnom podaní. Navyiac k tomu potrebujú extrémne množstvo vstupných dát, tréningových cyklov a vysoký výpočtový výkon. A akúkoľvek pokročilú oblasť, v ktorej sú technológie AI dokonca lepšie než ľudia, vieme dosiahnuť len za cenu extrémnej špecializácie a úzkeho zamerania systémov AI.⁵¹²

I to odráža súčasnú neexistenciu modelu mysle a modelovania správania systémov AI, keďže algoritmy umelej inteligencie nie sú schopné na základe skutočného chápania vytvárať tzv. mentálne modely dôležitých aspektov reálneho života a sveta. Ide v zásade o mentálne simulácie fungovania sveta, ktoré sa nám premietajú v mysli a na základe ktorých vieme nielen očakávať, čo sa stane, ale sa aj zmysluplne rozhodovať.⁵¹³

Algoritmy AI, i napriek veľkým pokrokom v rozpoznávaní reči a snahy strojovo pochopiť význam slovných vyjadrení, nedokážu uchopiť metafory, ktorými sa bežne vyjadrujeme, keďže – ako sme už viac krát uviedli – majú problém s abstrakciou a analógiou. Ľudská myseľ abstrahuje a vytvára analógie (vo väčšine prípadov podvedome) a tvorí tak základ pre formovanie konceptov. A koncepty sú základom chápania i celkového ľudského modelu myslenia a schopnosti adekvátne konať, resp. reagovať na situácie s ktorými sa človek v živote stretáva.⁵¹⁴

511 Por. MITCHELL, *Artificial Intelligence*, s. 236-237.

Por. PEARL, J., MACKENZIE, D. *The Book of Why: The New Science of Cause and Effect*. New York: Basic Books, 2018.

512 Napr. v súčasnosti asi najsofistikovanejšie systémy umelej inteligencie v oblasti hier Go a šach, ktorými sú viaceré verzie AlphaGo a AlphaGo Zero, sú absolútne nepoužiteľné napr. v autonómnych vozidlách, inteligentných zbraňových systémoch, onkologických alebo astrofyzikálnych systémoch a pod.

513 Por. MITCHELL, *Artificial Intelligence*, s. 240-242.

514 HOFSTADTER, D. R. *Analogy as the Core of Cognition*. [on-line]. Presidential Lecture, Stanford

Schopnosť abstrakcie, analógie, tvorby konceptov a využívania metafor patrí k nutným podmienkam pre vytvorenie modelu myslenia založeného na zdravom rozume. Pokiaľ systém umelej inteligencie nebude toho schopný, nikdy neprekročí hranicu medzi slabou a silnou umelou inteligenciou. Mnohí odborníci z oblasti umelej inteligencie sa zamýšľajú nad otázkou, či tieto tzv. „trhliny v inteligencii“ (cracks in intelligence) systémov AI môžu byť opravené vylepšovaním algoritmov a prísunom ďalších dát, alebo či niečo podstatné nám v týchto systémoch chýba.⁵¹⁵

Už od čias publikovania článku Alana Turinga o univerzálnom počítačom stroji, ktorý mal byť schopný vyriešiť ľubovoľnú výpočtovú a logickú operáciu (Turingov stroj⁵¹⁶) badať snahu vnímať schopnosť vykonávať logické operácie ako niečo samostatné, oddelené od subjektu, ktorým bola do tej doby výhradne ľudská bytosť. Túto ideu, siahajúcu až k Descartesovmu špekulatívnemu „naše telá a naše myšlienky sa skladajú z rôznych látok a podliehajú rôznym fyzikálnym zákonom“⁵¹⁷, si naplno osvojili aj účastníci Dartmouthského seminára. Dlhodobu (a dodnes) dominantným prístupom sa tak stalo presvedčenie, že všeobecná a silná umelá inteligencia môže byť dosiahnutá od ľudského tela oslobodenými počítačmi.

V komunite venujúcej sa umelej inteligencii však stále pretrvávala aspoň malá skupina vedcov, ktorá zastávala tzv. **hypotézu stelesnenia** (embodiment hypothesis) založenú na premise, že stroj nemôže dosiahnuť inteligenciu ľudskej úrovne bez nejakého druhu „tela“, v ktorom interaguje s okolitým svetom, pretože mozog oddelený od tela nikdy nemôže nadobudnúť koncepty potrebné pre všeobecnú inteligenciu.⁵¹⁸

V súčasnosti, vnímajúc vyššie opísané fundamentálne problémy a stagnáciu v smerovaní

University, 2009. [cit. 7. apríla 2022].

Dostupné na internete: <<https://www.youtube.com/watch?v=n8m7lFQ3njk>>

HOFSTADTER, D. R, SANDER, E. *Surfaces and Essences*. New York: Basic Books, 2013, s. 3.

515 MARCUS, G. *Deep Learning: A Critical Appraisal*. [on-line]. [cit. 7. apríla 2022].

Dostupné na internete: <<https://arxiv.org/abs/1801.00631>>

516 *Turing machine*. [on-line]. [cit. 7. apríla 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Turing_machine>

517 *Dualism*. [on-line]. [cit. 7. apríla 2022].

Dostupné na internete: <<https://plato.stanford.edu/entries/dualism/>>

518 Por. CLARK, A. *Being There: Putting Brain, Body, and World Together Again*. Cambridge, Mass.: MIT Press, 1996.

k silnej a všeobecnej umelej inteligencii, sa začína čím ďalej tým viac odborníkov prikláňať k hypotéze stelesnenia a veľa z nich pretavilo tento posun do svojej zaangažovanosti vo vývoji systémov, ktoré interagujú s reálnym svetom (napr. autonómne vozidlá alebo robotické systémy, ktorých plná autonómnosť a akcieschopnosť sa viaže na užívanie zdravého rozumu).⁵¹⁹

Ďalším problémom, ktorý vyjadruje bariéru chápania zmyslu, je kreativita. V predslove práce sme zo spomienok prof. Melanie Mitchell spomínali veľké znepokojenie Douglasa Hofstadtera po skúsenosti s činnosťou programu EMI (Experiments in Musical Intelligence). EMI generovala hudbu v štýle rôznych klasických skladateľov, využívajúc veľkú množinu pravidiel, ktoré boli zamerané na zachytenie všeobecnej syntaxe i jednotlivých špecifik komponovania. Tieto pravidlá boli následne aplikované na mnohé príklady z konkrétnych diel skladateľov s cieľom vytvoriť nové diela v konkrétnom štýle konkrétneho skladateľa. A niektoré napodobeniny boli tak dobré, že oklamali aj hudobných odborníkov, z čoho pramenilo aj spomínané Hofstadterovo znepokojenie. EMI ako vysoko sofistikovaný generátor skladieb vytvorila mnoho rôznych hudobných diel, z ktorých následne autor programu, David Cope, vyberal tie najlepšie – tie, ktoré boli schopné presvedčiť aj zdatné publikum o kráse programovo vytvorenej hudby.

Môžeme povedať, že EMI bola kreatívna? Myslíme si, že v žiadnom prípade. EMI totižto netvorila originálne motívy, len sofistikovane skladala hudbu z toho, čo mala analyzované z tvorby ľudských skladateľov. EMI skladbu, ktorú vytvorila, nevedela posúdiť – z jej „pohľadu“ boli všetky rovnako kvalitne vygenerované – až Cope, resp. iný ľudský poslucháč vedeli posúdiť a vyberať tie skladby, ktoré boli dobré. EMI jednoducho nebola schopná pochopiť hudbu, ktorú generovala, a to ani v rovine hudobných konceptov, ani v rovine emočného dopadu na poslucháča.⁵²⁰

V čase písania tejto kapitoly (apríl 2022) iniciatíva OpenAI predstavila DALL-E 2⁵²¹ –

519 Takto argumentuje napr. prof. Hiroshi Ishiguro, riaditeľ Laboratória inteligentnej robotiky, ktoré je súčasťou Katedry inovácie systémov na Graduate School of Engineering Science na univerzite v Osake v Japonsku.

CHILD, A. *Origin of the Species*. [filmový dokument]. Apple TV, 2021, 5:54. [cit. 25. apríla 2022].

Dostupné na internete:

<<https://tv.apple.com/us/movie/origin-of-the-species/umc.cmc.1vsh8or9ojg824l4kn2kidkuw>>

520 Por. MITCHELL, *Artificial Intelligence*, s. 274.

521 DALL-E 2. [on-line]. [cit. 8. apríla 2022].

systém umelej inteligencie, ktorý dokáže vytvoriť realistické obrazy a umenie na základe opisu v prirodzenom jazyku. I keď je medzi EMI a DALL-E 2 približne tridsaťročný rozdiel, v schopnosti umelej inteligencie byť kreatívnou sa vôbec nič nezmenilo. Akokoľvek zaujímavé obrazy DALL-E 2 tvorí, vôbec nechápe, čo a ako vytvára, a už vôbec nievie zhodnotiť umelecký a emočný dopad výsledku svojho algoritmu hlbokého učenia.

Na dosiahnutie kreativity nestačí len tvoriť niečo, čo môže byť označené ako kreatívne. Treba to aj vedieť pochopiť a ohodnotiť – či už na strane umelca alebo publika. Nič také však súčasné systémy umelej inteligencie nedokážu.

To isté môžeme povedať aj o rôznych aspektoch estetiky, ktoré bez prekročenia bariéry zmyslu sú technológiám AI nedostupné.

Zastávame názor, že ľudskému mysleniu sa približujúca schopnosť abstrakcie, analógie, konceptualizácie, simulácie, chápania zmyslu atď. sa nedá vytvoriť u subsymbolických systémov na základe vylepšovania súčasných modelov hlbokého učenia a masívneho prísunu trénovacích dát, resp. u symbolických systémov rozsiahlymi množinami symbolov a budovaním logických väzieb medzi nimi. Pre dosiahnutie tohoto cieľa by bolo treba prelomového vylepšenia algoritmov AI, ktoré – ako sa domnievame – súvisí so schopnosťou nielen analyticky uchopiť, ale i v reálnom svete zažiť a pochopiť niečo z podstaty myslenia a vedomia ľudskej bytosti.⁵²²

I napriek tomu, že sa odborná komunita už celé dekády snaží, nutné podmienky pre prekročenie bariéry chápania zmyslu, na ktoré sme doteraz poukázali, sa javia ako hlboko nepochopiteľné. Viac a viac sa preto ozývajú hlasy, ktoré prekročenie tejto bariéry spájajú s požiadavkou vedomia a sebauvedomenia, emócií, hypotézou stelesnenia, prípadne ďalšími aspektami, ktoré smerujú k analógii ľudskej osoby. A je pritom možné, že prekročiť túto bariéru znamená prijať aj niečo z limitov, emócií, logickej iracionality, socializácie a mnohých kognitívnych nedostatkov, ktoré sprevádzajú človeka a civilizáciu ako takú.⁵²³

Dostupné na internete: <<https://openai.com/dall-e-2/>>

522 Ako ANI bola výzvou pochopiť, ako tieto systémy fungujú, aký je dosah ich činnosti, aké dôsledky môžu mať na človeka a spoločnosť a ako sa k nim postaviť, aby boli eticky využívané, AGI je výzvou objaviť, ako by vlastne mala byť silná a všeobecná AI vytvorená, aký potenciál v sebe bude obnášať a čoho by bola schopná. Do akej miery môže byť všeobecná a či je vôbec reálne také niečo vytvoriť.

523 V podobnom duchu diskutuje aj prof. Melanie Mitchell na margo otázky, ako ďaleko sa nachádzame

Ešte nepolapiteľnejšie sa v tomto kontexte javia také méty, ako vnútorná sloboda, či zmysel pre dobro, krásu, obetu, utrpenie alebo lásku. Všetky sú navzájom prepojené a nielen z pohľadu kresťanského svetonázoru dostávajú vyšší zmysel a naplnenie v duchovnom rozmere ľudského bytia.⁵²⁴

5.2. Ontologické otázky

Organizácia Non-human Rights Project (NhRP), ktorá sa v roku 2012 pretransformovala z pôvodného Center for the Expansion of Fundamental Rights, aktívne presadzuje myšlienku začlenenia zvierat do kategórie osoby a priradenia im tomu zodpovedajúcich práv. V jednej zo svojich ostatných iniciatív spojilo sedemnást' vedcov svoje sily a adresovalo najvyššiemu súdu v New Yorku svoj pohľad (amicus brief) v mene možnosti priznania ľudských práv neľudským osobám. NhRP netvrdí, že i rôzne druhy zvierat (napr. primáty) sú ľudskými bytosťami, ale že aj ony sú osobami, ktoré majú nárok na svoje „ľudské“ práva.⁵²⁵

Podľa zástancov NhRP je tak pojem osoby len nádobou, do ktorej sa v okamihu, keď sa príslušnej entite prizná hodnosť osoby, môžu „nakvapkať“ príslušné práva. Legálne priznanie statusu osoby sa tak stáva základným a nevyhnutným predpokladom pre nárokovanie si akýchkoľvek práv a spoločenských možností súvisiacich s osobou. A spôsob, ako dosiahnuť tento právny status pre zvieratá, ktoré nie sú ľuďmi, je vedecky dokázať, že majú sebauvedomenie.⁵²⁶

od vytvorenia všeobecnej umelej inteligencie na úrovni človeka.

Por. MITCHELL, *Artificial Intelligence*, s. 275-277.

524 Dokázala by umelá inteligencia trpieť? A ak áno, vedela by pochopiť a naplniť slová rakúskeho neurológa a psychiatra Viktora E. Frankla, ktorý hovorí:

„To, pred čím má utrpenie človeka chrániť, je apatia, umŕtvujúca duševná strnulosť. Dokiaľ trpíme, zostávame duševne nažive. V utrpení dokonca vyzrievame, v utrpení rastieme, činí nás bohatšími a mocnejšími.“

„V niektorých ohľadoch utrpenie prestáva byť utrpením v okamihu, keď nájde zmysel, napríklad zmysel obety.“

Viktor Frankl citáty. [on-line]. [cit. 25. apríla 2022].

Dostupné na internete: <<https://citaty-slavy-nych.sk/autori/viktor-frankl/>>

525 Por. *The Nonhuman Rights Project: Frequently Asked Questions*. [on-line]. [cit. 10. apríla 2022].

Dostupné na internete: <<https://www.nonhumanrights.org/frequently-asked-questions/>>

526 THURZO, *The Influence of Existentialism and Subjectivism on the Concept of the Human Person*, s. 7.

V kontexte priznania štatútu osoby na základe sebauvedomenia sa diskusia o obsahu pojmu osoby transformuje na diskusiu o podstate sebauvedomenia.

V podaní NhRP tak ide o snahu modifikovať chápanie sebauvedomenia, aby zahŕňalo aj psychologické pochody primátov. Snahou NhRP je dokázať, že napr. spomínané primáty sú vnímavé, pričom vnímavosť (sentience) je schopnosť vnímať, cítiť a prežívať. Táto vlastnosť alebo schopnosť subjektu sa považuje za základ jeho autonómie, odkiaľ by mal byť už len malý krôčik k právnemu uznaniu statusu osoby pre primáty.⁵²⁷

Skúsme uvedené postrehy aplikovať na oblasť umelej inteligencie. Interakcia s viacerými pokročilými systémami AI by nás mohla viesť k záveru, že komunikujeme s niečím, čo dokáže byť vnímavé, dokáže reagovať na emócie a ich aj vzbudzovať, prípadne dokáže samostatne myslieť. Nemohli by sme sa tak pripojiť k argumentácii NhRP a snažiť sa pokročilé systémy umelej inteligencie⁵²⁸ vyhlásiť za neľudské osoby?

Ak by sme aj hypoteticky prijali rozšírený pohľad na osobu definujúcu sebauvedomenie na základe vnímavosti, dôrazne sme v predchádzajúcich kapitolách upozorňovali, že v prípade technológií AI ide vždy len o simuláciu procesov myslenia, teda aj o simuláciu čohokoľvek, čo by sme mohli nazvať súčasťou vnímavosti. Preto **zastávame názor, že ani v rozšírenom pohľade na osobu, definujúcu sebauvedomenie na základe vnímavosti, súčasným najpokročilejším systémom ANI nemôžeme priradiť štatút osoby.**

Radi by sme však podotkli, že je rozdiel medzi vnímavosťou a vedomím, resp. sebauvedomením. Vnímvosť, či senzibilita zahŕňa schopnosť subjektívneho vnímania, pocitov a skúseností. Avšak pri vedomí, resp. sebauvedomení ide o uvedomovanie si seba a svojho okolia v plnosti kognitívnych schopností, emócií, abstrahovaní a prenášaní skúseností, chápania zmyslu, procesov učenia a rozhodovania sa, vôľových aspektov konania a pod.

Napriek veľkým pokrokom v interdisciplinárnom rámci kognitívnej vedy stále nemáme

527 THURZO, *The Influence of Existentialism and Subjectivism on the Concept of the Human Person*, s. 8.

Neľudským osobám by boli spočiatku pridelené menšie práva, podľa NhRP by išlo predovšetkým o právo na fyzickú slobodu. Žiada sa nám však dodať, že sloboda ide ruka v ruke so zodpovednosťou, a tak uznanie všeobecného práva na slobodu v kontexte osoby so sebou prináša aj právnu zodpovednosť za svoje skutky. Ako však v tomto prípade realizovať penalizáciu, nápravu škody a trest, osobitne s jeho nápravným účinkom?

528 I napriek snahám niektorých odborníkov zaradiť najnovšie verzie Alpha Zero do kategórie AGI, stále sa pohybujeme v intenciách ANI.

vyriešené základné otázky, napr. vzťah medzi myslou a telom, nevieme presne definovať schopnosť vedomia a už vôbec nie, ako ju vytvoriť na báze čisto fyzikálnych (biofyzikálnych) procesov – a to ani na úrovni vnímanosti. Takže ani nevieme, ako ju preniesť do elektronických systémov.

Mnohé aspekty ľudskej mysle sa snažíme popísať tzv. teóriou mysle a vytvoriť tak model správania, ktorého základom je zdravý rozum ako predpoklad pre všeobecnú inteligenciu. Ani takto však nemáme odpoveď na otázku, čo je podstatou vedomia, takže nedokážeme ani určiť, aké kritériá by museli spĺňať systémy umelej inteligencie, aby sme ich mohli považovať za vnímavé a vedomé. Navyše, akýkoľvek morálny a právny systém je založený na existencii a koexistencii ľudských osôb, ktoré konajú na základe svojho svedomia, ako najbližšej normy mravnosti formovanej vonkajšou, objektivizujúcou normou, resp. normatívnym zákonom. V kontexte problémov, s ktorými sa – riešiac métu uvedomelej umelej inteligencie – potýkame, si ani nevieme predstaviť formu, resp. implementáciu mechanizmu, ktorý by etické princípy a – ak hypoteticky hovoríme o osobe – morálne zásady dokázal v rámci systému AGI realizovať.

V rámci nášho vzhladu do ontologickej perspektívy potencionálnych uvedomelých systémov AI, by sme radi akcentovali náš náhľad, ktorý považujeme za legitímny nielen v kontexte kresťanského svetonázoru, ale i z pohľadu súčasného vedeckého bádania v oblasti umelej inteligencie.

Ľudská bytosť v jednote tela a duše je z pohľadu kresťanskej antropológie a Zjavenia stvorená na Boží obraz.⁵²⁹ Ide o výnimočnosť, ktorá nás odlišuje od ostatných bytostí stvoreného fyzického sveta. Sebauvedomenie spojené s rozumom a slobodnou vôľou chápeme ako mohutnosť duše, ktoré presahujú biologickú realitu mozgu a nervovej sústavy. Na základe tejto výnimočnosti definujeme aj osobu len a výlučne ako ľudskú bytosť.

Materialistický uhol pohľadu – odmietajúc duchovnú dušu ako „formu“ tela, zásadne sa podieľajúcu na vedomí a sebauvedomení – však svojim protagonistom dáva nádej,⁵³⁰ že skôr či neskôr budeme schopní vytvoriť uvedomelú umelú inteligenciu s vedomím

529 *Katechizmus Katolíckej cirkvi*. 362-365. [on-line]. [cit. 10. apríla 2022].

Dostupné na internete: <<https://katechizmus.sk/>>

530 Napríklad v rozvoji genetických algoritmov umelej inteligencie a ich integrácii s ďalšími pokročilými metódami hlbokého učenia simulovať evolúciu smerujúcu k mysliacim bytostiam.

a sebauvedomením analogickým človeku, ktorá bude ašpirovať na štatút osoby.

Ak hypoteticky tento pohľad rozvíjame ďalej, miesto implementácie etických princípov a požiadaviek v systémoch ANI sa nevyhneme otázke analógie svedomia a morálnych hodnôt u uvedomelej AGI.

A ako sme naznačili vyššie, svedomie – ak má mať zmysel a plniť svoju funkciu – by malo mať referenčný bod mimo seba. V pohľade kresťanskej antropológie je svedomie teonómne, jeho referenčnou autoritou je Boh v rovine prirodzeného i pozitívneho zákona. **V súčasnosti nepoznáme spôsob, ako by bolo možné u umelej inteligencie dosiahnuť niečo podobné svedomiu. A ak by sme to aj dosiahli, nikto nám nedá záruky, že bude zdieľať rovnaké hodnoty a rovnakým spôsobom ich aj chrániť a zachovávať.**⁵³¹

531 Znie to síce ako z lacného sci-fi, no ako veľmi zjednodušený príklad môžeme vziať napr. rozhodnutie AGI pre pomyselnú ochranu budúcnosti ľudstva eliminovať časť obyvateľstva, ktorú by podľa svojich kritérií tento systém AGI pokladal za ohrozenie a pod. Za chápaním bizarnosti tejto úvahy stojí v myslí človeka celý hodnotový aparát stavajúci na zdravom rozume. A ak by zdravý rozum zlyhával, nastupuje legislatívne i praktické obmedzenie možností konkrétnej ľudskej osoby, brániace také niečo vykonať. Ak by išlo o AGI, ktorá by mala napr. na starosti armádne systémy alebo kritickú infraštruktúru, tieto obmedzenia nemusia fungovať.

Záver

Rozprávanie o stvorení v knihe Genezis, s jej potvrdením poriadku a dobroty všetkého stvorenia ako prejavu jediného a milujúceho Boha, vyjadruje základný vzťah a prvotnú zmluvu⁵³², ktorú Boh uzavrel s ľudstvom.

Meditácia tohoto pohľadu, spojená s gréckou metafyzickou špekuláciou o bytí a kauzalite, jednote a rozdielnosti, pohybe, zmene a číslach, umožňovala človeku chápať výskum prírodných vied ako „čítanie knihy stvorenia“.⁵³³ A v kontexte dejinného vývoja s príchodom scholastickej metódy⁵³⁴ a takých princípov, ako je napr. Occamova britva, sa ľudstvo usilovalo vziať do ruky rydlo, neskôr brko s kalamárom a nakoniec moderné pero, aby do tejto knihy prírody aj niečo doplnilo.

Ak však máme v tejto analógii pokračovať ďalej, rozvoj umelej inteligencie spojený v rámci štvrtej priemyselnej revolúcie s napredovaním v robotike a biotechnológiách možno prirovnať ku Gutenbergovmu stroju, ktorý znamenal prelom v tom, čo a ako dokázal človek zapísať – teda ako a v čom chceme knihu stvorenia dopĺňať alebo prepísať. Je tak na nás, či sa dokážeme stále vracieť k *počiatkom* (~Genezis) a mať dobre premeditované to, čo nám Boh zjavil, aby sme tak „Gutenbergov vynález umelej inteligencie“ správne použili a tento svet skutočne rozvíjali. Veď návod „ako na to“ sme v Ježišovi Kristovi a v evanjeliu dostali a treba ho len rozmeniť na drobné, aby sme ho aj v oblasti umelej inteligencie napĺňali.

V práci sme sa primárne venovali problematike etických výziev a morálnych aspektov súčasných systémov umelej inteligencie, ktorú poznáme pod spoločným označením slabá umelá inteligencia (ANI). Tento termín zahŕňa 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),

532 Por. ŠANTAVÝ, P. *Dejiny spásy: zmluvy Starého zákona a Nová zmluva* [diplomová práca]. [on-line]. Bratislava: RKCMBF UK, 2000. [cit. 6. novembra 2021].

Dostupné na internete: <https://peter.santavy.cloud/docs/Dejiny_spasy-zmluvy_SZ_a_NZ.pdf>

533 Por. SMITH, R. *Why No Science?* [on-line]. [cit. 6. novembra 2021].

Dostupné na internete: <<https://www.thecatholicthing.org/2021/11/02/why-no-science/>>

534 *Scholastic Method* [on-line]. [cit. 6. novembra 2021].

Dostupné na internete: <<https://www.encyclopedia.com/religion/encyclopedias-almanacs-transcripts-and-maps/scholastic-method>>

ktoré vykazujú inteligentné 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.

Jedným z hlavných prínosov tejto práce je ponúknutý kvalitný a primerane hlboký interdisciplinárny rámec, bez ktorého nie je možné úspešne realizovať skutočné riešenie etických problémov a výziev technológií umelej inteligencie. Ide o rámec, v ktorom sme dostatočne oboznámení aj s technologickou stránkou týchto systémov a psychologickými, sociologickými i právnymi aspektmi ich nasadenia.

V prvej kapitole sme preto ponúkli potrebný náhľad do problematiky umelej inteligencie, sumarizujúc jej základné vlastnosti, delenie a metódy. Osobitne sme sa venovali algoritmom inšpirovaným činnosťou mozgu, keďže ony sú základom prakticky všetkých pokročilých systémov AI. Poukázali sme na základné problémy niekoľkých dekád vývoja systémov AI i na masívne zavádzanie týchto technológií v súčasnosti, pričom nezostal opomenutý ani futuristický presah fenoménu umelej inteligencie.

Druhá kapitola predstavuje základ pre **d'alší dôležitý prínos práce – identifikáciu, pomenovanie, analýzu a pochopenie rizík spojených s technológiami umelej inteligencie v celej možnej šírke spektra ich nasadenia.**

Pre reálne uchopenie problematiky etiky týchto technológií považujeme široko spektrálne uchopenie a komplexné pochopenie limitov a rizík súčasných systémov AI za podstatné. Preto sme sa v druhej kapitole pomerne obširne venovali hlavným rizikovým faktorom a zraniteľnostiam algoritmov súčasných systémov AI, zaoberali sme sa bezpečnosťou procesov založených na týchto technológiách, vysvetľovali problematiku kybernetickej bezpečnosti ako integrálnej súčasti zabezpečenia funkčnosti systémov AI a neopomenuli sme ani riziká vyplývajúce z technologickej komplexnosti a potrebného infraštruktúrneho zázemia pre spoľahlivú činnosť v reálnom svete.

Analýza uvedených limitov a rizík bola základom pre rozoberanie negatívnych dôsledkov využívania technológií umelej inteligencie v spoločnosti, zahŕňajúc celé spektrum ich nasadenia od sociálnych sietí až po problematiku autonómnych vozidiel a rozširujúc tak interdisciplinárny rámec o psychologické a sociologické dôsledky ich využitia. Osobitne sme sa venovali oblasti dohľadových systémov, technológií využívaných v spravodajských službách a v rámci algoritmického riadenia štátu, keďže z pohľadu etických výziev ide o veľmi špecifické pole vývoja i nasadenia systémov umelej inteligencie. Azda

najproblematickejšou oblasťou je adaptácia technológií AI vo vojenskej sfére – od využitia v armádnych spravodajských službách, cez modelovanie a simulácie technológií a procesov, virtualizačné nástroje a výcvik, až po vývoj a nasadenie smrtiacich autonómnych zbraňových systémov a kybernetických zbraní. Tejto oblasti sme venovali osobitný priestor, keďže armádne využitie systémov AI v sebe obnáša celé spektrum otázok s potenciálom prevýšiť všetky ostatné etické dilemy a výzvy.

V prvej časti tretej kapitoly sme **ako ďalší osobitný prínos tejto práce sumarizovali naše etické postrehy a závery, ku ktorým sme dospeli v rámci analýzy limitov a rizík súčasných systémov AI:**

- súčasné systémy AI sprevádzajú oprávnené a vážne obavy z toho, že ak nechápeme ako tieto systémy pracujú (prakticky každý sofistikovaný systém umelej inteligencie sa javí ako black box), nemôžeme im reálne dôverovať a ťažko dokážeme predpovedať okolnosti, za ktorých tieto systémy zlyhajú.
- problém s verifikovateľnosťou postupov a výsledkov činnosti systémov umelej inteligencie sa tak premieta i do oblasti implementácie etických pravidiel, či už ide o spôsob, ako hodnotiť a kontrolovať činnosť systémov AI, alebo ide priamo o spôsob implementácie regulácií technológiami AI.
- súčasné systémy umelej inteligencie trpia rôznymi zraniteľnosťami. Musíme mať neustále na pamäti, že tieto systémy môžu zlyhávať najrozličnejšími a často neočakávanými spôsobmi. Systémy umelej inteligencie robia chyby, ktoré sú diametrálne odlišné od ľudských chýb a zlyhaní. Preto môžu byť prekvapivé, neočakávané a nebezpečné.
- s týmito zlyhaniami a obmedzeniami treba rátať aj v oblasti snahy o technologickú implementáciu etických noriem a mantinelov. Nemusí to byť vôbec triviálne a v niektorých scenároch nasadenia ani uskutočniteľné.
- vzhľadom na rozličné typy útokov zameraných na procesy umelej inteligencie sa etické výzvy neviažu len na návrh a realizáciu systémov AI, no rozširujú sa aj o oblasť etiky použitia a z toho prameniace riziká zneužitia.
- systémy umelej inteligencie sú vystavené aj problémom v oblasti kybernetickej bezpečnosti, keďže mnohé riziká atakujúce bezpečnosť procesov umelej inteligencie, zneužívajúce nedostatky dizajnu alebo amplifikujúce dôsledky

rizikových faktorov, ako vektor útoku používajú zraniteľnosti v oblasti kybernetickej bezpečnosti. Kybernetická bezpečnosť je proces, ktorý má svoju dynamiku a vyjadruje skutočnosť, že ani v oblasti systémov umelej inteligencie neexistuje dokonale bezpečný a spoľahlivý systém.

- vývoj a nasadenie sofistikovaných systémov AI poukazuje na ich veľkú komplexnosť, ktorá je nielen rizikovým faktorom bezpečnosti a stability fungovania systémov, ale mnohokrát i neľahkou výzvou pre úspešnú a jasnú implementáciu regulácií a etických pravidiel.
- úspešné nasadenie mnohých moderných systémov AI vyžaduje prísun extrémneho množstva dát z reálneho sveta a priamo z ľudského prostredia. Pri veľkých systémoch zhromažďujúcich denno denne extrémne množstvo dát je veľmi ťažké zaručiť etiku spracúvania týchto údajov a vylúčiť riziko ich zneužitia.
- systémy umelej inteligencie sa s veľkou mierou istoty stávajú schopnými vytvárať veľmi presné psychologické profily, odhaľovať akékoľvek väzby a mnohé osobné informácie, predikovať a ovplyvňovať naše konanie. Tieto systémy však mnohokrát nie sú predmetom regulácie, resp. verejnej kontroly.
- neuvedomujúc si, ako je naša myseľ a psychika zraniteľná, existuje riziko postupného prechodu od technologického prostredia založeného na systémoch AI k prostrediu založenému na závislosti a manipulácii.
- mýlnikom, ktorého by sme sa mali obávať, teda 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.
- v kontexte slabej umelej inteligencie (ANI) ešte stále konečnú voľbu cieľov nerobia stroje, ale človek, takže stále je v ľudskej moci tieto dôsledky ovplyvniť a zmeniť. Preto musia existovať etické pravidlá a regulácie, ktoré by dokázali chrániť jednotlivca i celú spoločnosť voči rizikám sofistikovanej práce systémov AI v rámci sociálnych sietí, systémov riadenia spoločnosti, resp. analogických platforiem.
- viaceré aplikácie systémov umelej inteligencie v reálnom svete budú musieť v zlomkoch sekúnd riešiť rôzne kritické situácie – a vedieť ich vyriešiť eticky. Otázkou je, do akej miery môžu byť etické princípy a tomu zodpovedajúce právne regulácie aj technicky úspešne implementovateľné.

- vo viacerých oblastiach nasadenia bude treba vyriešiť rozdelenie zodpovednosti medzi človekom a riadiacim systémom AI a s tým súvisiace legislatívne požiadavky pre reálnu prevádzku týchto systémov.
- musíme sa seriózne zamyslieť, ako stanoviť akceptovanú mieru neistoty v oblasti spoľahlivosti a robustnosti systémov umelej inteligencie.
- fenomén digitálneho rozdelenia (digital divide), ktorý sa s neustálym rozvojom informačných technológií a prechodom k informačnej spoločnosti stáva reálnym problémom, môže byť vďaka necitlivému nasadeniu technológií umelej inteligencie oproti súčasnosti ešte znásobený.
- v rámci algoritmického riadenia sa zvyšuje riziko zneužitia systémov AI na efektívnu kategorizáciu a obmedzovanie ľudských práv občanov. Aké sú možnosti nastavenia etických mantinelov a reálne dodržiavaných právnych rámcov v oblasti riadenia štátu, spravodajstva a dohľadu?
- trio technológií umelej inteligencie – sledovanie a dohľadové systémy spojené s analytickými nástrojmi a schopnosťou vytvárať modely predikujúce ľudské konanie či vývoj situácie – tvorí v súčasnosti skutočne silnú technologickú výbavu (nielen) spravodajských služieb. Ako spoločnosť musíme vedieť zabezpečiť striktnú zákonnosť i dôsledný dohľad demokraticky zvolených zástupcov a spoločnosti pri nasadení systémov AI v rámci spravodajských služieb, dohľadových systémov a všetkých foriem i stupňov algoritmického riadenia spoločnosti.
- pokročilé modelovanie technológií, simulácia priebehu vojenských operácií i priebehu častí konfliktu a silný virtualizačný efekt môžu viesť k morálnemu znecitlivieniu obsluhujúceho personálu a velenia. Dôsledkom môže byť strata vnímanosti pre ohrozenie ľudských životov a reálnych hodnôt pri skutočnom bojovom nasadení.
- existuje silný tlak na využívanie autonómnych zbraňových systémov, ktoré sú vďaka technológiám umelej inteligencie vo viacerých scenároch nasadenia lepšie ako ľudia. Ako dokážeme zabezpečiť, že smrtiace autonómne zbraňové systémy (LAWs), ktoré sú zámerne dizajnované tak, aby dokázali efektívne pracovať aj pri strate spojenia, budú vedieť správne použiť svoju smrtiacu silu a rešpektovať požadované mantinely, ak je nám známe, že systémy AI robia mnohokrát iné chyby

ako ľudia a nevieme ich vytrénovať tak, aby vedeli adekvátne odpovedať na všetky reálne situácie? Uvedený problém ešte narastá pri skupinovom, resp. kooperatívnom riadení autonómnych zbraňových systémov.

- útočné kybernetické systémy využívajúce súčasné technológie umelej inteligencie považujeme za veľmi rizikové a nevhodné (nielen) pre vojenské nasadenie. Keďže nemáme reálne odpovede ohľadom eliminácie ich rizík a nasadenia, stupňuje sa apel proti akémukoľvek nasadeniu útočných autonómnych zbraní, ak sú mimo zmysluplnej ľudskej kontroly.
- v súčasnosti nemáme vedecké dôkazy o schopnosti automatizovaných systémov disponovať funkciami potrebnými na presnú identifikáciu cieľa, situačné povedomie alebo rozhodnutia týkajúce sa primeraného použitia sily. LAWs tak môžu spôsobiť vysokú mieru vedľajších škôd a preto sa rozhodnutia o použití hrubej sily nesmú delegovať na stroje.
- schopnosti umelou inteligenciou poháňaných autonómnych zbraňových systémov a kybernetických zbraní i napriek všetkým rizikám vedú k zvyšujúcim sa tlakom na financovanie a zavádzanie útočných kybernetických zbraní. Je veľkou etickou a regulačnou výzvou, ako tieto tendencie zastaviť.
- problematika obmedzenia týchto útočných systémov je skomplikovaná aj reálnym stieraním hraníc medzi obranným a útočným nasadením takmer vo všetkých oblastiach vojenského využitia technológií umelej inteligencie.
- nutnou podmienkou prevádzky ľubovoľného systému AI, ktorý môže predstavovať riziko pre akúkoľvek ľudskú osobu, je schopnosť a možnosť človeka prebrať kedykoľvek kontrolu nad týmto systémom, resp. právo a možnosť verifikovať a prehodnotiť výsledky jeho činnosti.
- limity, regulácia a obmedzenia LAWs by mali predstavovať etický rámec stanovený na základe morálnych hodnôt ľudskej spoločnosti, nie na základe relativistickej tzv. „následnej regulácie“, ktorý môže viesť k nepredvídateľným morálnym dopadom.
- na základe doterajších skúseností v oblasti limitov a rizík súčasných systémov AI je otáznosť, či v celom spektre technológií a algoritmov umelej inteligencie bude možné pevný etický rámec dodržať. Vieme si to predstaviť pri cílení – čo môže znamenať aj limitujúcom a tým aj znevýhodňujúcom – dizajne systémov AI

s dôrazom na dodržiavanie navrhnutého rámca (ethics by design), avšak pri súčasnej snahe o čo najširšie adoptovanie umelej inteligencie v armáde a získanie konkurenčnej výhody sme v tomto smere pomerne skeptickí.

- v kontexte rastúcej dôvery v systémy umelej inteligencie a vnímajúc problémy, s ktorými sa nasadenie týchto systémov potýka i oblastí, ktorých sa ich využitie dotýka, môžeme a musíme stanoviť podmienky, bez splnenia ktorých by nasadenie systémov AI do reálneho sveta, v ktorom interagujú s človekom a vplývajú na spoločnosť, nemalo byť umožnené.

Uvedený sumár problémov technológií umelej inteligencie s priamym či nepriamym dopadom na človeka a spoločnosť sa stal základom pre naše návrhy riešenia etických problémov a stanovenie všeobecných i špecifických etických zásad, ktoré predkladáme v štvrtej kapitole.

Interdisciplinárny rámec sme v rámci tretej kapitoly rozšírili analýzou súčasných aktivít na poli etiky umelej inteligencie smerujúc k potrebným reguláciám na zabezpečenie etického rámca využívania týchto technológií. Osobitne sme sa venovali európskemu Aktu o umelej inteligencii, ktorý považujeme síce za náročnú, avšak v súčasnosti asi najprepracovanejšiu a najkomplexnejšiu (pripravovanú) reguláciu v oblasti umelej inteligencie s potenciálom ovplyvniť etiku využívania systémov AI vo veľkej časti sveta. Tretia kapitola bola završená analýzou súčasných aktivít Cirkvi na poli etiky umelej inteligencie, pričom sme sa osobitne zaoberali závermi rímskej konferencie renAIssance 2020, známej aj pod názvom Rome Call for Ethics.

Problematika etických výziev a morálnych aspektov súčasných systémov slabej umelej inteligencie (ANI) bola završená naším návrhom riešenia etických problémov AI, ktoré sme predkladali v štvrtej kapitole. Ponajprv išlo o vyjadrenie a naše chápanie základného, a to pozitívneho postoja k fenoménu umelej inteligencie vo svetle Zjavenia. Osobitne sme akcentovali interdisciplinárny rámec prístupu k problematike etiky AI, bez ktorého skutočné riešenie etických problémov a výziev technológií umelej inteligencie nie je možné úspešne realizovať.

K hlavným prínosom práce patrí náš návrh základnej štruktúry etických princípov a zásad, ktorý sme v štvrtej kapitole predstavili:

- rozšírili sme diapazón základného zamerania umelej inteligencie na človeka

(human-centered AI) o kontext kresťanskej antropológie, inklúziu každej ľudskej bytosti bez diskriminácie so zreteľom na dobro ľudstva a spoločnosti v rozšírenej optike starostlivosti o náš spoločný a zdieľaný domov, teda o celý stvorený svet.

- definovali sme a objasnili principiálne požiadavky na dôveryhodné systémy AI, ktoré musia byť legálne, etické a robustné. Vo vedomí, že dokonalý systém umelej inteligencie neexistuje, sme navrhli minimálne legislatívne, etické a technologické požiadavky, resp. hranicu, od ktorej môžeme tieto technológie považovať za dôveryhodné.
- predstavili sme vlastnú sadu všeobecných a univerzálnych etických zásad:
 - pri vývoji, výrobe, nasadení, poskytovaní a používaní systémov umelej inteligencie musí byť zaručená ochrana slobody, dôstojnosti a bezpečia každej ľudskej osoby i celej spoločnosti.
 - technológie umelej inteligencie musia byť plne pod ľudskou kontrolou a ovládateľné človekom.
 - algoritmy i výsledky činnosti systémov AI musia byť človekom pochopiteľné a revidovateľné.
 - akékoľvek nasadenie technológií AI musí byť prospešné pre človeka a spoločnosť.
 - systémy umelej inteligencie nesmú byť nástrojom digitálneho rozdelenia.
 - technológie umelej inteligencie nesmú škodiť nášmu spoločnému domu a mali by prispievať k spoločenskému a environmentálnemu blahobytu.
- nami predstavené etické zásady sme porovnali so zásadami najdôležitejších súčasných aktivít a regulácií v oblasti AI, aby sme tak vyjadrili univerzálnosť a nadčasovosť nášho návrhu.
- naším zámerom bolo navrhnuť a vytvoriť tak univerzálnu a všeobecnú množinu etických zásad, že podľa nej môžeme či už definovať – alebo ešte lepšie – z existujúcich legislatívnych rámcov a etických odporúčaní vyberať konkrétne a jasné odporúčania pre ich aplikáciu v reálnom svete. Takto sme i označili kombináciu etických odporúčaní vatikánskej konferencie renAIssance 2020 a obsahu európskeho nariadenia Akt o umelej inteligencii za v súčasnosti

najvhodnejšie konkrétne a do legislatívneho rámca zasadené odporúčania pre etický vývoj, nasadenie a využívanie súčasných systémov umelej inteligencie.

- tiež sme považovali za dôležité popísať všetky oblasti nutnej implementácie etických noriem, eticko-právnych regulácií a morálnych zásad. Ide o oblasť tvorby systémov AI, poskytovateľov i používateľov týchto systémov a implementácie noriem i obmedzení priamo v systémoch AI. Akcentovali sme viaceré činitele, bez ktorých nie je možné tieto oblasti zasadiť do etického rámca vývoja, nasadenia a vyžívania – ide napr. o edukáciu a osvetu, pravidlá pre základný výskum, nutné podmienky vývoja a pod.

Vzhľadom na dôraz, ktorý sme v druhej kapitole kládli na oblasť pokročilého riadenia štátu, spravodajstva a plošného dohľadu i využívania systémov umelej inteligencie vo vojenskej oblasti, v štvrtej kapitole prinášame i **naše vlastné závery, návrhy regulácií a etické odporúčania v týchto špecifických a dôležitých oblastiach, ktoré považujeme za ďalší osobitný prínos tejto práce:**

- využitie technológií AI v oblasti pokročilého riadenia štátu, spravodajstva a plošného dohľadu musí byť vzhľadom na svoju povahu, potenciál a riziká pokryté už základnými legislatívnymi mechanizmami a verejným dohľadom demokratickej spoločnosti, ktoré sa týkajú riadenia spoločnosti, ľudských práv a pôsobenia spravodajských služieb vo všeobecnosti.
- oblasť exportu produktov a technológií umelej inteligencie, ktoré môžu byť zneužitú v oblasti pokročilého riadenia štátu, spravodajstva a plošného dohľadu, by mala byť predmetom medzinárodnej regulácie s cieľom zamedziť ich vývoz do rizikových krajín.
- technológiami AI poháňané automatické smrtiace zbraňové systémy (LAWs), systémy automatického zameriavania a vyberania cieľov, automatické systémy schopné bez zásahu človeka rozhodnúť o smrtiacej reakcii akéhokoľvek druhu (od útoku dronu až po rozpútanie jadrového konfliktu) musia byť zakázané.
- nutnou podmienkou prevádzky ľubovoľného systému AI, ktorý môže predstavovať riziko pre akúkoľvek ľudskú osobu, je schopnosť a možnosť človeka prebrať kedykoľvek kontrolu nad týmto systémom, resp. právo a možnosť verifikovať a prehodnotiť výsledky jeho činnosti.

- limity, regulácia a obmedzenia LAWS by mali predstavovať etický rámec stanovený na základe morálnych hodnôt ľudskej spoločnosti, nie na základe relativistickej tzv. „následnej regulácie“.

V závere štvrtej kapitoly **sme predložili niekoľko návrhov pre využitie potenciálu Cirkvi a osobitnú angažovanosť podporujúcu etický prístup k problematike umelej inteligencie v celej jej šírke.** Ide predovšetkým o misiu zjednocovať, usmerňovať a propagovať etické aktivity vo svete a neustávajúcu snahu akcentovať a budovať univerzálne bratstvo a sociálne priateľstvo aj v oblasti digitálneho sveta a jeho technológií.

Sekundárnym aspektom nášho úsilia bolo v kontexte súčasných technológií ANI poukázať i na problematiku uvedomelej umelej inteligencie (AGI), ktorá by podľa jej protagonistov mala byť dosiahnuteľná prostredníctvom silnej (strong) a všeobecnej (general) AI. Všeobecnej, lebo dokáže zvládnuť akúkoľvek intelektuálnu úlohu a má schopnosť generalizovať, t.j. zovšeobecňovať a prenášať, či adaptovať naučené schopnosti na iné úlohy. Silnej, pretože aj skutočne rozumie tomu, čo rieši a vykonáva.

I keď prakticky pri všetkých témach rozoberaných v prvých štyroch kapitolách sme sa nevyhli aspoň krátkemu pohľadu za horizont – k systémom silnej a všeobecnej umelej inteligencie, až v piatej kapitole sme sa výlučne venovali jej viacerým podstatným problémom, pričom **naše uchopenie problematiky uvedomelej umelej inteligencie taktiež považujeme za dôležitý prínos tejto práce:**

- diskutovali sme teóriu mysle a zdravého rozumu so schopnosťou abstrakcie, analógie, konceptualizácie, simulácie, či chápania zmyslu, identifikujúc tak bariéru chápania zmyslu, venovali sme sa hypotéze stelesnenia a predstavili sme komplexné dôvody, pre ktoré nie sme v súčasnosti schopní vytvoriť silnú a všeobecnú umelú inteligenciu.
- v kontexte nielen klasickej definície, ale i modernistickej snahy o redefiníciu pojmu osoby sme polemizovali s víziou uvedomelej AI a poukázali na dôvody, prečo ani najpokročilejším systémom ANI nemôžeme priradiť štatút osoby.
- na základe kresťanskej antropológie sme vyjadrili výnimočnosť ľudskej bytosti a tým aj odmietnutie štatútu osoby pre akúkoľvek umelú inteligenciu, vrátane AGI.
- použijúc materialistický uhol pohľadu – odmietajúci duchovnú dušu ako „formu“ tela, zásadne sa podieľajúcu na vedomí a sebauvedomení – sme okrem iného poukázali

na v súčasnosti neriešiteľnú problematiku implementácie mechanizmu svedomia.

- osobitne sme akcentovali reálny problém funkčného mechanizmu svedomia, ktoré má referenčný bod mimo seba. U hypotetickej analógie svedomia uvedomelej umelej inteligencie nie sme schopní garantovať, že bude zdieľať rovnaké hodnoty a rovnakým spôsobom ich aj chrániť a zachovávať.

Celkovo prácu predkladáme ako určité nóvum, snažiac sa o inovatívny prístup v nazeraní na problematiku umelej inteligencie v kontexte morálnej teológie.

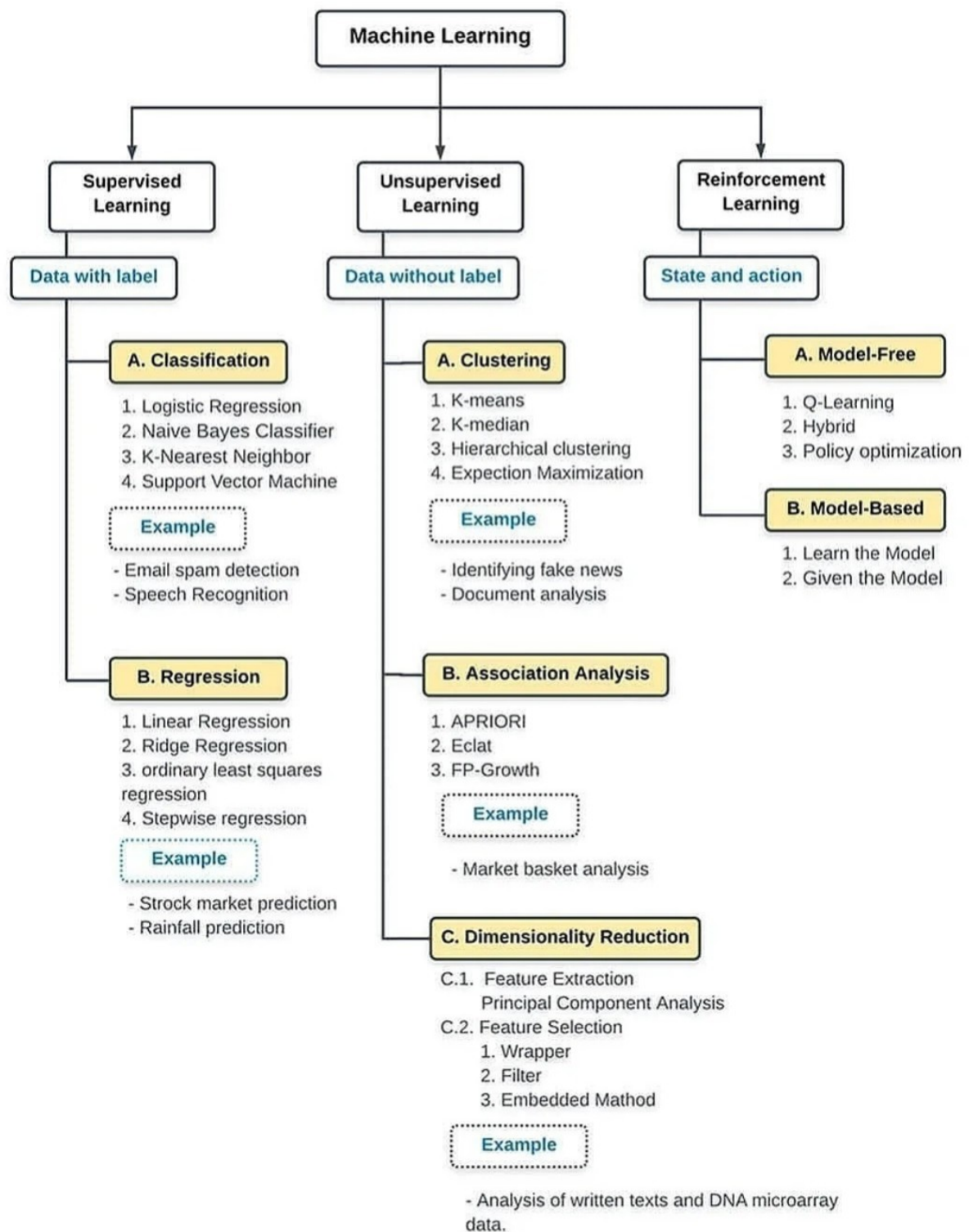
Interdisciplinárne uchopenie fenoménu AI, dôsledná analýza limitov i rizík a z nej prameniace etické závery, návrh základnej štruktúry všeobecných i špecifických etických princípov a zásad i vyjadrenie diskutabilných faktorov uvedomelej umelej inteligencie a zdôvodnenie nášho jasného postoja k jej porovnávaní s ľudskou bytosťou – to všetko tvorí náš vklad do prebiehajúceho celosvetového etického diskurzu, ktorý tým viac naberá na dôležitosti, čím viac sa technológie AI rozvíjajú a jej sofistikované implementácie sa do reálneho života masívne zavádzajú, ovplyvňujúc tak životy miliónov ľudí.

Nástojčivosť tejto témy vynikne, ak si z našej rozpravy o oblastiach implementácie etických princípov (kap. 4.3.4) pripomenieme, že v prípade pokročilých systémov umelej inteligencie – na rozdiel od iných informačných systémov a technológií – prakticky nie je možné na existujúce a už nasadené systémy spoľahlivo aplikovať etické zásady a regulácie prostredníctvom doplnenia technických úprav alebo procesných postupov. Jednoducho povedané, s etikou musíme vývoj týchto systémov sprevádzať a ich realizáciu predchádzať, inak nám – nielen jednotlivcom, ale celej spoločnosti digitálneho veku – tieto technológie prerastú cez hlavu.

Dúfame, že táto práca prispeje k rozšíreniu (nielen) etických obzorov v jednej z najdynamickejšie sa rozvíjajúcich oblastí digitálneho sveta a láskavý čitateľ bude zhovievavý k prípadným nepresnostiam a nedostatkom nášho pohľadu na problematiku etických výziev a morálnych aspektov súčasných systémov umelej inteligencie.

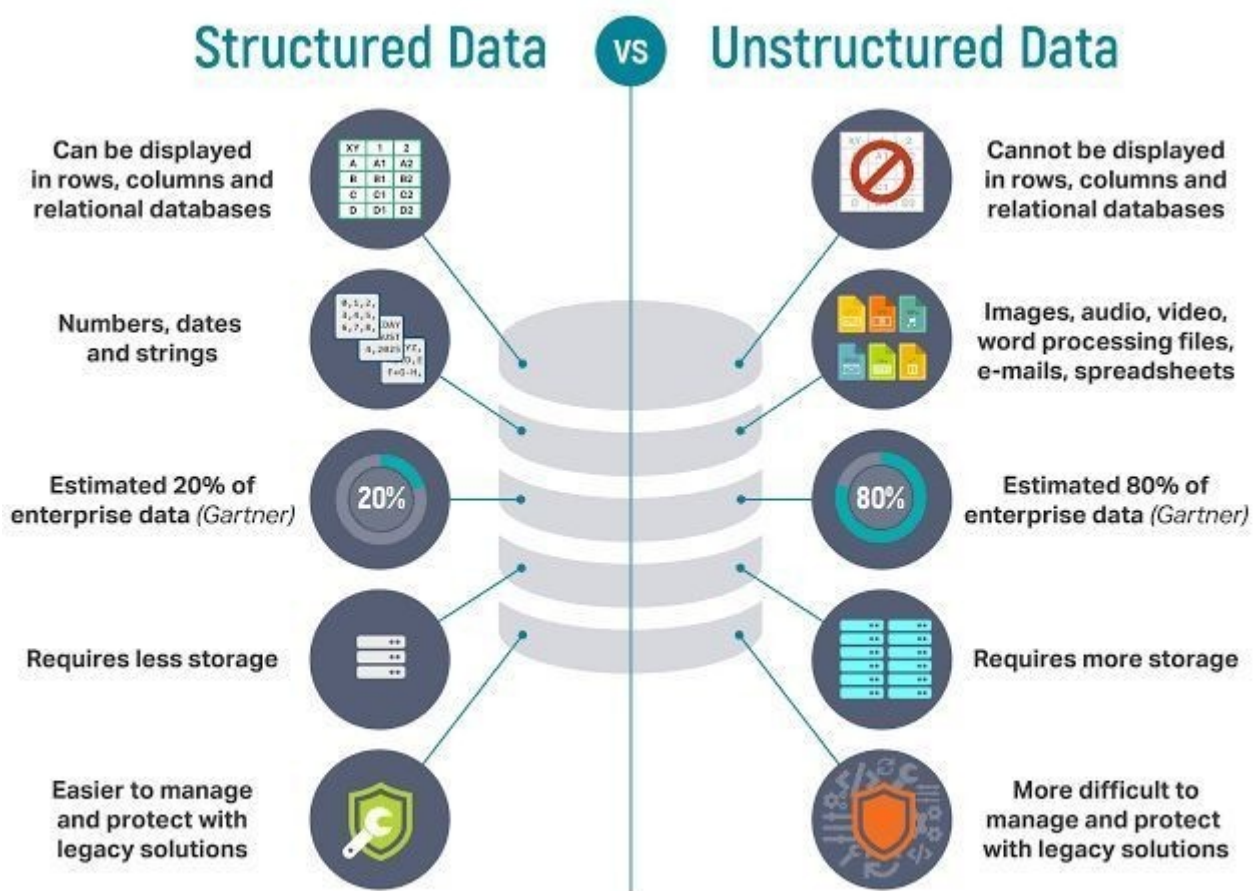
Prílohy

Príloha č. 1 – algoritmy strojového učenia



Kredit: AI4Diversity (<https://twitter.com/AI4Diversity>)

Príloha č. 2 – štruktúrované a neštruktúrované dáta



Kredit: AI4Diversity (<https://twitter.com/AI4Diversity>)

Príloha č. 3 – čínsky systém sociálneho kreditu

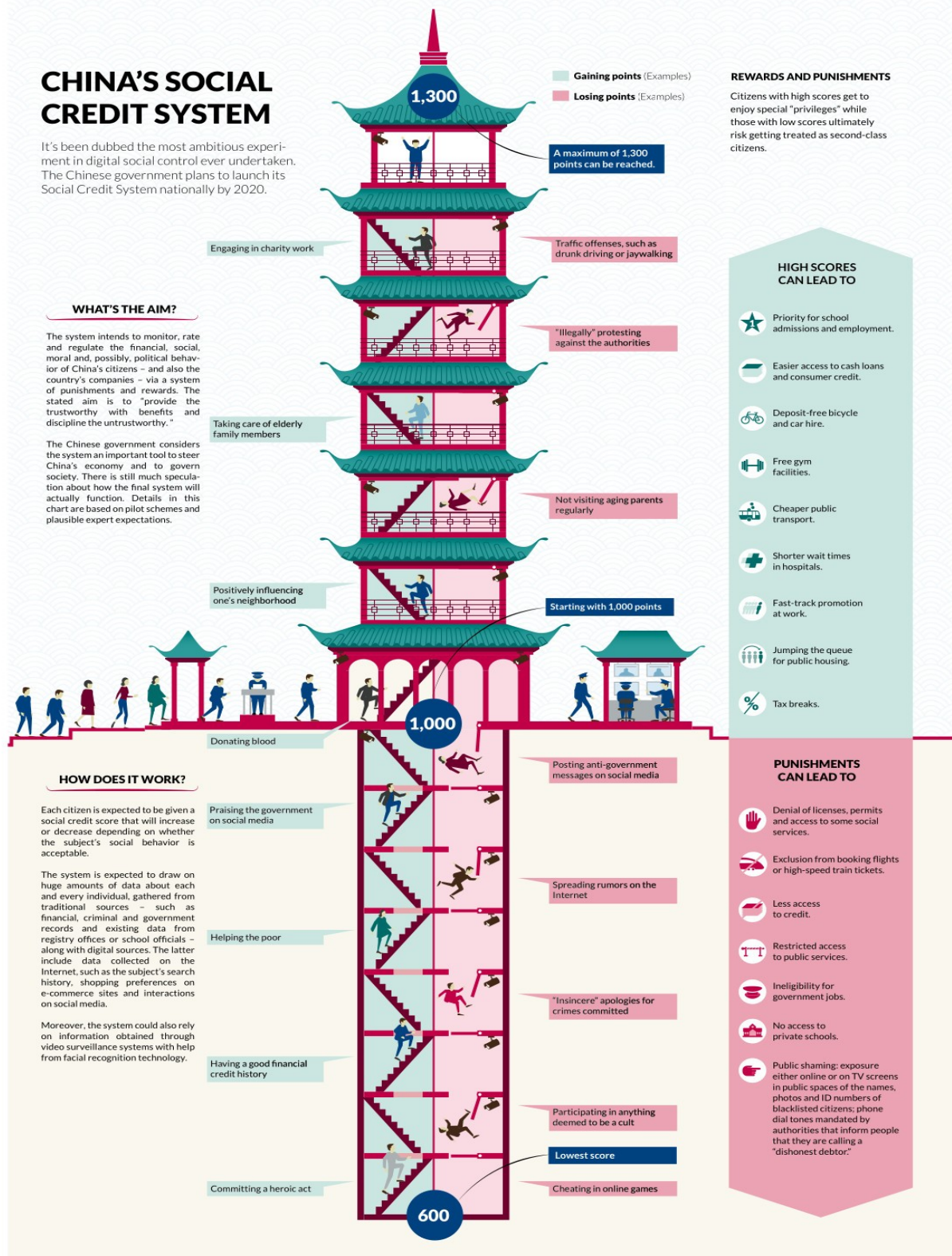
CHINA'S SOCIAL CREDIT SYSTEM

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, moral and, possibly, political behavior of China's citizens – and also the country's companies – via a system of punishments and rewards. The stated aim is to "provide the trustworthy with benefits and discipline the untrustworthy."

The Chinese government considers the system an important tool to steer China's economy and to govern society. There is still much speculation about how the final system will actually function. Details in this chart are based on pilot schemes and plausible expert expectations.



Kredit: Bertelsmann Stiftung (https://www.bertelsmann-stiftung.de/fileadmin/files/aam/Asia-Book_A_03_China_Social_Credit_System.pdf)

Zoznam použitej literatúry

ADIB-MOGHADDAM, A. *Artificial intelligence must not be allowed to replace the imperfection of human empathy*. [on-line]. [cit. 20. februára 2022].

Dostupné na internete: <<https://theconversation.com/artificial-intelligence-must-not-be-allowed-to-replace-the-imperfection-of-human-empathy-151636>>

AI Principles: Recommendations on the Ethical Use of Artificial Intelligence by the Department of Defense. [on-line]. [cit. 9. marca 2022].

Dostupné na internete: <[https://admin.govexec.com/media/dib_ai_principles_-_supporting_document_-_embargoed_copy_\(oct_2019\).pdf](https://admin.govexec.com/media/dib_ai_principles_-_supporting_document_-_embargoed_copy_(oct_2019).pdf)>

ALFONSEC, M., CEBRIAN, M. et al.: *Superintelligence Cannot be Contained: Lessons from Computability Theory* [on-line]. [cit. 26. januára 2021].

Dostupné na internete: <<https://jair.org/index.php/jair/article/view/12202>>

Algoritmy strojového učenia I. [on-line]. [cit. 3. januára 2022].

Dostupné na internete: <<https://umelainteligencia.sk/algoritmy-strojoveho-ucenia/>>

Algoritmy strojového učenia II. [on-line]. [cit. 3. januára 2022].

Dostupné na internete: <<https://umelainteligencia.sk/algoritmy-strojoveho-ucenia-ii-ucenie-bez-ucitela/>>

Algoritmy strojového učenia III. [on-line]. [cit. 4. januára 2022].

Dostupné na internete: <<https://umelainteligencia.sk/algoritmy-strojoveho-ucenia-iii-ucenie-formou-odmenovania/>>

AMODEI, D., HERNANDEZ, D. *AI and Compute*. [on-line]. [cit. 19. januára 2022].

Dostupné na internete: <<https://openai.com/blog/ai-and-compute/>>

An Open Letter to the United Nations: Convention on Certain Conventional Weapons. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://futureoflife.org/2017/08/20/autonomous-weapons-open-letter-2017/>>

ANDERSON, K., WAXMAN, M. C. *Law and Ethics for Autonomous Weapon Systems: Why a Ban Won't Work and How the Laws of War Can*. In: *Jean Perkins Task Force on National Security and Law Essay Series*. Stanford University: Hoover Institution Press, 2013.

ANDERSON, K., WAXMAN, M. C. *Law and Ethics for Robot Soldiers*. In: *Policy Review*. 2012, 176, 12.

Artificial Intelligence: Principles, laws, and frameworks. OneTrust DataGuidance Limited, 2022. ISSN 2398-9955.

Autonómne vozidlá [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<http://akgunis.sk/autonomne-vozidla/>>

Autonomous Weapons: An Open Letter from AI [Artificial Intelligence] & Robotics Researchers. [on-line]. Future of Life Institute, 2015. [cit. 8. marca 2022].

Dostupné na internete: <<http://futureoflife.org/open-letter-autonomous-weapons/>>

BECK, B. R. et al. *Predicting commercially available antiviral drugs that may act on the novel coronavirus (2019-nCoV), Wuhan, China through a drug-target interaction deep learning model*. In: *bioRxiv*. [online]. 2020, 2020.01.31.929547. [cit. 6. augusta 2020].

DOI: 10.1101/2020.01.31.929547

Dostupné na internete: <<https://www.biorxiv.org/content/10.1101/2020.01.31.929547v1>>

BENGIO, Y. *Machines Dream*. In: BEYER D. ed. *The Future of Machine Intelligence: Perspectives from Leading Practitioners*. Sebastopol, Calif.: O'Reilly Media, 2016.

BOGERT, E., SCHECTER, A., WATSON, R. T. *Humans rely more on algorithms than social influence as a task becomes more difficult*. In: *Sci Rep*. [on-line]. 2021, 11, 8028. [cit. 20. februára 2022].

DOI: 10.1038/s41598-021-87480-9

Dostupné na internete: <<https://doi.org/10.1038/s41598-021-87480-9>>

BOSTROM, N. *A history of transhumanist thought*. In: *Journal of Evolution and Technology*. [on-line]. 2005, roč. 14, vyd. 1. [cit. 8. januára 2022].

Dostupné na internete: <<http://www.nickbostrom.com/papers/history.pdf>>

CAMPBELL, J. *Dvojník*. Fantom Print, 2015. ISBN 978-80-7398-329-7.

CELLAN-JONES, R. *Stephen Hawking Warns Artificial Intelligence Could End Mankind*. In: *BBC News*. [on-line]. 2014, 2. 12. [cit. 5. augusta 2020].

Dostupné na internete: <<https://www.bbc.com/news/technology-30290540>>

CLAPPER, J. R. Jr. et al. *Unmanned Systems Roadmap: 2007-2032*. [on-line].

Washington, DC: Department of Defense [DOD], 2007. [cit. 5. marca 2022].

Dostupné na internete: <http://www.globalsecurity.org/intell/library/reports/2007/dod-unmanned-systems-roadmap_2007-2032.pdf>

CLARK, A. *Being There: Putting Brain, Body, and World Together Again*. Cambridge, Mass.: MIT Press, 1996.

Colonial Pipeline offline due to ransomware attack. [on-line]. [cit. 13. marca 2022].

Dostupné na internete: <<https://www.fortiguard.com/outbreak-alert/darkside>>

Compilation of open letters against autonomous weapons. [on-line]. [cit. 19. augusta 2020].

Dostupné na internete: <<https://autonomousweapons.org/compilation-of-open-letters-against-autonomous-weapons/>>

Conversations That Matter: The Crossroads of Science and Human Dignity Fall 2021. [on-line]. [cit. 28. marca 2022].

Dostupné na internete:

<<https://mcgrath.nd.edu/conferences/academic-pastoral/conversations-that-matter-the-crossroads-of-science-and-human-dignity/conversations-that-matter-the-crossroads-of-science-and-human-dignity-fall-2021/>>

Cybercrime Is Now More Profitable Than The Drug Trade [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <<https://www.tripwire.com/state-of-security/regulatory-compliance/pci/cybercrime-is-now-more-profitable-than-the-drug-trade/>>

Data Mining Techniques [on-line]. [cit. 7. decembra 2015].

Dostupné na internete: <<http://documents.software.dell.com/statistics/textbook/data-mining-techniques>>

Defense Science Board. *Task Force Report: The Role of Autonomy in DoD Systems*. Washington, DC: Office of the Under Secretary of Defense for Acquisition, Technology and Logistics, 2012.

DOBBINS, J., COHEN, R. S., CHANDLER, N. et al. *Overextending and Unbalancing Russia: Assessing the Impact of Cost-Imposing Options*. [on-line]. Santa Monica, CA: RAND Corporation, 2019. [cit. 18. marca 2022].

Dostupné na internete: <https://www.rand.org/pubs/research_briefs/RB10014.html>

DOE/NNSA/LANL/SNL [on-line]. [cit. 3. marca 2022].

Dostupné na internete: <<https://www.top500.org/site/50334/>>

Dualism. [on-line]. [cit. 7. apríla 2022].

Dostupné na internete: <<https://plato.stanford.edu/entries/dualism/>>

Elements of AI. [on-line]. [cit. 30. marca 2022].

Dostupné na internete: <<https://www.elementsofai.com/>>

Ethics guidelines for trustworthy AI. [on-line]. [cit. 28. marca 2022].

Dostupné na internete: <<https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>>

ETZIONI, A., ETZIONI, O. *Keeping AI Legal*. In: *Vanderbilt Journal of Entertainment & Technology Law*. [on-line]. 2016, 19, no. 1, s. 133-146. [cit. 8. marca 2017].

Dostupné na internete:

<http://www.jetlaw.org/wp-content/uploads/2016/12/Etzioni_Final.pdf>

ETZIONI, A., ETZIONI, O. *Pros and Cons of Autonomous Weapons Systems* [on-line]. [cit. 5. marca 2022].

Dostupné na internete:

<<https://www.armyupress.army.mil/Journals/Military-Review/English-Edition-Archives/May-June-2017/Pros-and-Cons-of-Autonomous-Weapons-Systems/>>

FELDSTEIN, S. *The Global Expansion of AI Surveillance*. [on-line]. [cit. 28. februára 2022].

Dostupné na internete: <<https://carnegieendowment.org/2019/09/17/global-expansion-of-ai-surveillance-pub-79847>>

FERRIS, R. *Elon Musk thinks we will have to use AI this way to avoid a catastrophic future*. [on-line]. [cit. 24. januára 2022].

Dostupné na internete: <<https://www.cnbc.com/2017/01/31/elon-musk-thinks-we-will-have-to-use-ai-this-way-to-avoid-a-catastrophic-future.html>>

First Turing Test success marks milestone in computing history [on-line]. [cit. 29. januára 2021].

Dostupné na internete: <<https://phys.org/news/2014-06-turing-success-milestone-history.html>>

CHADHA, S. *“Common Sense” is the Dark Matter of Artificial Intelligence*. [on-line]. [cit. 4. februára 2022].

Dostupné na internete: <<https://hackernoon.com/the-dark-matter-of-ai-common-sense-is-not-so-common>>

GARAMONE, J. *9/11 Drove Change in Intelligence Community, NSA Chief Says*. [on-line]. [cit. 27. februára 2022].

Dostupné na internete:

<<https://www.defense.gov/News/News-Stories/Article/Article/945544/911-drove-change-in-intelligence-community-nsa-chief-says/>>

Gartner *Top Strategic Predictions For 2020 And Beyond*. [on-line]. [cit. 23. marca 2022].

Dostupné na internete: <<https://www.gartner.com/smarterwithgartner/gartner-top-strategic-predictions-for-2020-and-beyond>>

GEIST, E., LOHN, A. J. *How Might Artificial Intelligence Affect the Risk of Nuclear War?* [on-line]. Santa Monica, CA: RAND Corporation, 2018. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.rand.org/pubs/perspectives/PE296.html>>

GIBNEY, A. *Zero Days*. [filmový dokument]. [cit. 10. marca 2022].

Dostupné na internete: <<http://www.zerodaysfilm.com/>>

GIDDA, M. *Edward Snowden and the NSA files – timeline* [on-line]. In: *The Gurdian*. 2013, 23. 6. [cit. 23. februára 2022].

Dostupné na internete: <<http://www.guardian.co.uk/world/2013/jun/23/edward-snowden-nsa-files-timeline>>

GRIFFIN, A. *Facebook's artificial intelligence robots shut down after they start talking to each other in their own language* [on-line]. [cit. 6. augusta 2020].

Dostupné na internete:

<<https://www.independent.co.uk/life-style/gadgets-and-tech/news/facebook-artificial-intelligence-ai-chatbot-new-language-research-openai-google-a7869706.html>>

Google » *Tensorflow: Vulnerability Statistics* [on-line]. [cit. 11. januára 2022].

Dostupné na internete: <https://www.cvedetails.com/product/53738/Google-Tensorflow.html?vendor_id=1224>

Google » *Tensorflow: Security Vulnerabilities* [on-line]. [cit. 11. januára 2022].

Dostupné na internete:

<https://www.cvedetails.com/vulnerability-list/vendor_id-1224/product_id-53738/Google-Tensorflow.html>

GREGOROVÁ, D. *Nedostatek spánku působí na mozek* [on-line]. [cit. 20. januára 2022].

Dostupné na internete: <<https://www.osel.cz/12127-nedostatek-spanku-pusobi-na>>

mozek.html>

HAMILTON, I. A. *Amazon built an AI tool to hire people but had to shut it down because it was discriminating against women*. [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://www.businessinsider.com/amazon-built-ai-to-hire-people-discriminated-against-women-2018-10>>

HE, K., CHEN, X., XIE, S., LI, Y., DOLLÁR, P., GIRSHICK, R.: *Masked Autoencoders Are Scalable Vision Learners* [on-line]. Facebook AI Research (FAIR), 2021. [cit. 18. novembra 2021].

Dostupné na internete: <<https://arxiv.org/abs/2111.06377>>

HE, S. et al. *Learning to predict the cosmological structure formation*. In: *Proceedings of the National Academy of Sciences* [online]. 2019, roč. 116, č. 28, s. 13825. [cit. 6. augusta 2020].

DOI: 10.1073/pnas.1821458116

Dostupné na internete: <<https://www.pnas.org/content/116/28/13825>>

HIGGINS, A. *Stephen Hawking's final warning for humanity: AI is coming for us*. [on-line]. [cit. 13. marca 2022].

Dostupné na internete:

<<https://www.vox.com/future-perfect/2018/10/16/17978596/stephen-hawking-ai-climate-change-robots-future-universe-earth>>

Hi Reddit, I'm Bill Gates and I'm back for my third AMA. Ask me anything. In: *Reddit*. [on-line]. 2015, 28. 1. [cit. 7. augusta 2020].

Dostupné na internete:

<https://www.reddit.com/r/IAmA/comments/2tzjp7/hi_reddit_im_bill_gates_and_im_back_for_my_third/>

HOFSTADTER, D. R. *Analogy as the Core of Cognition*. [on-line]. Presidential Lecture, Stanford University, 2009. [cit. 7. apríla 2022].

Dostupné na internete: <<https://www.youtube.com/watch?v=n8m7IFQ3njk>>

HOFSTADTER, D. R. *Gödel, Escher, Bach: an Eternal Golden Braid*. New York: Basic Books, 1979. ISBN: 978-0-465-02656-2.

HOFSTADTER D. R. *Staring Emmy Straight in the Eye – and Doing My Best Not to Flinch*. In: DARTNELL T. *Creativity, Cognition, and Knowledge*. Westport, Conn.: Praeger, 2002.

HOFSTADTER, D. R, SANDER, E. *Surfaces and Essences*. New York: Basic Books, 2013.

How 9/11 Changed the Course of Personal Data Collection and Surveillance. [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <<https://www.startpage.com/privacy-please/startpage-articles/how-9-11-changed-the-course-of-personal-data-collection-and-surveillance>>

HUMBER, M. *Technology and Workforce: Comparison between the Information Revolution and the Industrial Revolution* [on-line]. Berkeley: University of California, 2007. [cit. 20. augusta 2020].

Dostupné na internete:

<<http://infoscience.epfl.ch/record/146804/files/InformationSchool.pdf>>

CHILD, A. *Origin of the Species*. In: *Apple TV*. [filmový dokument]. 2021. [cit. 25. apríla 2022].

Dostupné na internete:

<<https://tv.apple.com/us/movie/origin-of-the-species/umc.cmc.1vsh8or9ojg824l4kn2kidkuw>>

China wants to make its own rules for AI ethics. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.scmp.com/abacus/tech/article/3029194/china-wants-make-its-own-rules-ai-ethics>>

CHOI, Q. CH. *Superintelligent AI May Be Impossible to Control; That's the Good News* [on-line]. [cit. 26. januára 2021].

Dostupné na internete: <<https://spectrum.ieee.org/tech-talk/robotics/artificial-intelligence/super-artificialintelligence>>

IEEE Global Initiative on Ethics of Autonomous and Intelligent Systems. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <http://standards.ieee.org/develop/indconn/ec/ead_v2.pdf>

IEEE Global Initiative on Ethics of Autonomous and Intelligent Systems. [on-line]. [cit. 31. marca 2022].

Dostupné na internete: <<https://standards.ieee.org/industry-connections/ec/autonomous-systems/>>

ITAPA: Umelá inteligencia v bežnom živote. [on-line]. [cit. 24. júna 2021].

Dostupné na internete: <<https://www.itapa.sk/12826-sk/program/>>

It's Getting Harder to Spot a Deep Fake Video. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.youtube.com/watch?v=gLoI9hAX9dw>>

JÁN PAVOL II. *Redemptoris Missio*. Praha: Zvon, 1994.

JOBIN, A., IENCA, M., VAYENA, E. *The global landscape of AI ethics guidelines*. In: *Nat Mach Intell*. [on-line]. 2019, 1, s. 389–399. [cit. 28. marca 2022].

Dostupné na internete: <<https://doi.org/10.1038/s42256-019-0088-2>>

JOBIN, A., IENCA, M., VAYENA, E. *Artificial Intelligence: the global landscape of ethics guidelines*. [on-line]. [cit. 19. augusta 2020].

Dostupné na internete: <<https://arxiv.org/pdf/1906.11668>>

Katechizmus Katolíckej cirkvi. [on-line]. [cit. 10. apríla 2022].

Dostupné na internete: <<https://katechizmus.sk/>>

KHIMJI, I. *Cybercrime Is Now More Profitable Than The Drug Trade*. [on-line]. [cit. 14. februára 2022].

Dostupné na internete: <<https://www.tripwire.com/state-of-security/regulatory-compliance/pci/cybercrime-is-now-more-profitable-than-the-drug-trade/>>

KNIGHT, W. *The Dark Secret at the Heart of AI*. In: *Technology Review*. [on-line]. 2017, 11. 4.. [cit. 10. februára 2022].

Dostupné na internete: <<https://www.technologyreview.com/2017/04/11/5113/the-dark-secret-at-the-heart-of-ai/>>

KORAKOVOUNIS, D. *Spiking Neural Networks: where neuroscience meets artificial intelligence*. [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <<https://theaisummer.com/spiking-neural-networks/>>

KOUSHIK, J. *Understanding Convolutional Neural Networks*. [on-line]. [cit. 31. januára 2022].

Dostupné na internete: <<https://arxiv.org/abs/1605.09081>>

KRAFT, A. *Microsoft shuts down AI chatbot after it turned into a Nazi*. [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://www.cbsnews.com/news/microsoft-shuts-down-ai-chatbot-after-it-turned-into-racist-nazi/>>

LAM, L., CHEN, S. *US spies on Chinese mobile phone companies, steals SMS data: Edward Snowden* In: *South China Morning Post*. [on-line]. 2013, 23. 3. [cit. 23. februára 2022].

Dostupné na internete: <<https://www.scmp.com/news/china/article/1266821/us-hacks-chinese-mobile-phone-companies-steals-sms-data-edward-snowden>>

LECUN, Y., MISRA, I. *Self-supervised learning: The dark matter of intelligence* [on-line]. [cit. 4. februára 2022].

Dostupné na internete: <<https://ai.facebook.com/blog/self-supervised-learning-the-dark-matter-of-intelligence>>

LEE, I. *Equalism: Paradise Regained*. In: LEE N. (ed.). *The Transhumanist Handbook*. Springer, 2019, s. 849 – 863.

LEHMAN, J., CLUNE, J., RISI, S.: *An Anarchy of Methods: Current Trends in How Intelligence Is Abstracted in AI*. In: *IEEE Intelligent Systems*. 2014, 29, č. 6.

MAASS, W. *Networks of spiking neurons: The third generation of neural network models*. [on-line]. [cit. 28. februára 2022].

Dostupné na internete:

<<https://www.sciencedirect.com/science/article/abs/pii/S0893608097000117>>

MARCUS, G. *Deep Learning: A Critical Appraisal*. [on-line]. [cit. 7. apríla 2022].

Dostupné na internete: <<https://arxiv.org/abs/1801.00631>>

MARCHANT, G. E. et al. *International Governance of Autonomous Military Robots*. In: *Columbia Science and Technology Law Review*. [on-line]. 2011, 12. 6., s. 272–276. [cit. 27. marca 2017].

Dostupné na internete: <<http://stlr.org/download/volumes/volume12/marchant.pdf>>

MCCARTHY, J. et al. *Proposal for the Dartmouth Summer Research Project in Artificial Intelligence*. In: *AI Magazine*, [on-line]. 1955, 27(4). [cit. 3. februára 2021].

Dostupné na internete: <<https://doi.org/10.1609/aimag.v27i4.1904>>

MCFARLAND, M. *Elon Musk: „With Artificial Intelligence, We Are Summoning the Demon“*. In: *Washington Post*. 2014, 24. 10..

Measuring trends in Artificial Intelligence – 2021 AI Index Report. [on-line]. [cit. 28. marca 2022].

Dostupné na internete: <<https://aiindex.stanford.edu/ai-index-report-2021/>>

MERCER, C., TROTHEN, T. J. *Religion and Transhumanism: The Unknown Future of Human Enhancement*. Praeger, 2014.

METZ, C. *A New Way for Machines to See, Taking Shape in Toronto*. New York Times, Nov. 28, 2017.

MIHULKA, S. *První varování před koronavirem Wuhan poslala umělá inteligence BlueDot*. [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://www.osel.cz/11002-prvni-varovani-pred-koronavirem-wuhan-poslala-umela-inteligence-bluedot.html>>

MIHULKA, S. *Letecká bojová umělá inteligence si natřela na chleba taktické experty*. [on-line]. [cit. 5. marca 2022].

Dostupné na internete: <<https://www.osel.cz/8903-letecka-bojova-umela-inteligence-si-natrelo-na-chleba-takticke-experty.html>>

MINSKY, M. L. *Computation: Finite and Infinite Machines*. Upper Saddle River, N.J.: Prentice Hall, 1967.

MINSKY, M. L. *The Emotion Machine: Commonsense Thinking, Artificial Intelligence, and the Future of the Human Mind*. New York: Simon & Schuster, 2006.

MINSKY, M. L., PAPERT, S. L. *Perceptrons: An Introduction to Computational Geometry*. Cambridge, Mass: MIT Press, 1969.

MINSKY, M. *Society of Mind*. New York: Simon & Schuster, 1986.

MITCHELL, M. *An Introduction to Genetic Algorithms*. Cambridge, Mas: MIT Press, 1996.

MITCHELL, M. *Artificial Intelligence*. Farrar, Straus and Giroux, 2019, ISBN: 978-0-374-71523-6.

MITCHELL, M. *Conceptual Abstraction and Analogy in Artificial Intelligence*. In: *ALIFE 2020: The 2020 Conference on Artificial Life*. [on-line]. 2020. [cit. 5. apríla 2022].

Dostupné na internete: <https://doi.org/10.1162/isal_a_00354>

MOEWES, C., NÜRNBERGER, A. *Computational Intelligence in Intelligent Data Analysis*. New York: Springer, 2013.

MORAVEC, H. *Mind Children: The Future of Robot and Human Intelligence*. Cambridge, Mass.: Harvard University Press, 1988.

MÜLLER, V. C., BOSTROM, N. *Future Progress in Artificial Intelligence: A Survey of*

Expert Opinion. In: *Fundamental Issues of Artificial Intelligence*. Cham, Switzerland: Springer International, 2016, s. 555-572.

Nariadenie Európskeho parlamentu a Rady (EÚ) 2016/679 z 27. apríla 2016 o ochrane fyzických osôb pri spracúvaní osobných údajov a o voľnom pohybe takýchto údajov, ktorým sa zrušuje smernica 95/46/ES (všeobecné nariadenie o ochrane údajov). [on-line]. [cit. 19. februára 2022].

Dostupné na internete: <<https://eur-lex.europa.eu/legal-content/SK/TXT/?uri=celex%3A32016R0679>>

Nariadenie Európskeho parlamentu a Rady, ktorým sa stanovujú harmonizované pravidlá v oblasti umelej inteligencie (Akt o umelej inteligencii) a menia niektoré legislatívne akty únie. [on-line]. [cit. 24. marca 2022].

Dostupné na internete: <<https://eur-lex.europa.eu/legal-content/SK/TXT/?uri=CELEX:52021PC0206>>

New Deep Learning Method Adds 301 Planets to Kepler's Total Count. [on-line]. [cit. 28. novembra 2021].

Dostupné na internete: <<https://www.jpl.nasa.gov/news/new-deep-learning-method-adds-301-planets-to-keplers-total-count>>

NG, A. *Deep Learning in Practice: Speech Recognition and Beyond*. In: *EmTech Digital*. [on-line]. 2016, 23. 5. [cit. 3. februára 2022].

Dostupné na internete: <<https://events.technologyreview.com/video/watch/andrew-ng-deep-learning/>>

NGUYEN, A., YOSINSKI, J., CLUNE, J. *Deep Neural Networks are Easily Fooled: High Confidence Predictions for Unrecognizable Images*. [on-line]. CVPR, 2015. [cit. 12. februára 2022].

Dostupné na internete: <https://cv-foundation.org/openaccess/content_cvpr_2015/papers/Nguyen_Deep_Neural_Networks_2015_CVPR_paper.pdf>

NIELSEN, M. *Neural Networks and Deep Learning*. [on-line]. [cit. 19. januára 2022].

Dostupné na internete: <<http://neuralnetworksanddeeplearning.com/>>

NILSSON, N. J., MCCARTHY, J. *A Biographical Memoir*. Washington D.C.: National Academy of Sciences, 2012.

One Hundred Year Study on Artificial Intelligence. [on-line]. AI100, 2016. [cit. 14. novembra

2021].

Dostupné na internete: <<https://ai100.stanford.edu/2016-report>>

One Hundred Year Study on Artificial Intelligence. [on-line]. AI100, 2021. [cit. 14. novembra 2021].

Dostupné na internete: <<https://ai100.stanford.edu/2021-report>>

Open Letter on Research Priorities for Robust and Beneficial Artificial Intelligence. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://futureoflife.org/2015/10/27/ai-open-letter/>>

ORLOWSKI, J. *The Social Dilemma*. Netflix, 2020. [filmový dokument]. [cit. 7. decembra 2021].

Dostupné na internete: <<https://www.netflix.com/sk/title/81254224>>

От робототехники до беспилотников: Шойгу рассказал о новинках ОПК на форуме «Армия-2020». [on-line]. [cit. 1. marca 2022].

Dostupné na internete: <<https://youtu.be/U8V4OZyNSac?t=152>>

PAROULKOVÁ, V. *Lidská práva pro roboty? Evropská unie chystá právní statut tzv. umělé osoby* [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://plus.rozhlas.cz/lidska-prava-pro-roboty-evropska-unie-chysta-pravni-statut-tzv-umele-osoby-6598059>>

PATTYNOVÁ, J. *Výzvy a právní aspekty umělé inteligence*. In: *Umělá inteligence 2021*. Praha: TUESDAY Business Network, 2021.

Dostupné na internete: <<https://www.tuesday.cz/akce/umela-inteligence-1/zaznam-akce/>>

PEARL, J., MACKENZIE, D. *The Book of Why: The New Science of Cause and Effect*. New York: Basic Books, 2018.

PFEIFFER, M., PFEIL, T. *Deep Learning With Spiking Neurons: Opportunities and Challenges*. [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <<https://www.frontiersin.org/articles/10.3389/fnins.2018.00774/full>>

POITRAS, L., ROSENBACH, M., SCHMID, F., STARK, H. *NSA horcht EU-Vertretungen mit Wanzen aus* In: *Der Spiegel*. [on-line]. 2013, 29. 6. [cit. 23. februára 2022].

Dostupné na internete: <<http://www.spiegel.de/netzwelt/netzpolitik/nsa-hat-wanzen-in-eu-gebaeuden-installiert-a-908515.html>>

POITRAS, L., ROSENBACH, M., STARK, H. *NSA überwacht 500 Millionen Verbindungen in Deutschland* In: *Der Spiegel*. [on-line]. 2013, 30. 6. [cit. 23. februára 2022].

Dostupné na internete: <<http://www.spiegel.de/netzwelt/netzpolitik/nsa-ueberwacht-500-millionen-verbindungen-in-deutschland-a-908517.html>>

Prečo Industry 4.0 [on-line]. [cit. 20. augusta 2020].

Dostupné na internete: <<http://industry4.sk/o-industry-4-0/co-je-industry-4-0/>>

PRESS, G. *12 Observations About Artificial Intelligence From The O'Reilly AI Conference*. In: *Forbes*. [on-line]. 2016, 31. 10. [cit. 7. augusta 2020].

Dostupné na internete: <<https://www.forbes.com/sites/gilpress/2016/10/31/12-observations-about-artificial-intelligence-from-the-oreilly-ai-conference/>>

Profile: Edward Snowden In: *BBC*. [on-line]. 2013, 24. 6. [cit. 23. februára 2022].

Dostupné na internete: <<http://www.bbc.co.uk/news/world-us-canada-22837100>>

REHÁK, M. *Útoky na systémy umělé inteligence a jejich obrana*. In: *Umělá inteligence 2021*. Praha: TUESDAY Business Network, 2021.

Dostupné na internete: <<https://www.tuesday.cz/akce/umela-inteligence-1/zaznam-akce/>>

Robot Sophia speaks at Saudi Arabia's Future Investment Initiative [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://youtu.be/dMrX08PxUNY>>

Rome Call for AI Ethics. [on-line]. [cit. 19. augusta 2020].

Dostupné na internete: <<http://www.academyforlife.va/content/pav/en/events/intelligenza-artificiale.html>>

Rome Call for AI Ethics (document). [on-line]. [cit. 28. marca 2022].

Dostupné na internete:

<https://www.romecall.org/wp-content/uploads/2022/03/RomeCall_Paper_web.pdf>

ROSENBLATT, F. *The Perceptron: A Probabilistic Model of Information Storage and Organization in the Brain*. In: *Psychological Review*. 1958, 65, č. 6.

ROTA, G.C. *Indiscrete Thoughts*. Boston: Berkhäuser, 1997.

RUMELHART, D. E., MCCLELLAND, J. L. *Parallel Distributed Processing*. Vol 1/2. Bradford Book, 1986.

RUSSELL, S. *Human Compatible*. Penguin Books, 2020, ISBN: 978-0-241-33524-6.

SAMPLE, I. *Thousands of leading AI researchers sign pledge against killer robots*. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.theguardian.com/science/2018/jul/18/thousands-of-scientists-pledge-not-to-help-build-killer-ai-robots>>

SAMUEL, A. L. *Some Studies in Machine Learning Using the Game of Checkers*. In: *IBM Journal of Research and Development*. 1959, č. 3.

Scientists' Call to Ban Autonomous Lethal Robots. ICRC, October 2013. [on-line]. [cit. 8. marca 2022].

Dostupné na internete: <<http://www.icrac.net/>>

SEARLE, R. J. *Minds, Brains, and Programs* In: *The Behavioral and Brain Sciences*. [on-line]. Cambridge University Press, 1980, zv. 3. [cit. 30. januára 2021].

Dostupné na internete:

<<https://web.archive.org/web/20071210043312/http://members.aol.com/NeoNoetics/MindsBrainsPrograms.html>>

SHARKEY, N. *Saying 'No!' to Lethal Autonomous Targeting*. In: *Journal of Military Ethics*. 2010, 9, č. 4, s. 369–383.

Scholastic Method [on-line]. [cit. 6. novembra 2021].

Dostupné na internete: <<https://www.encyclopedia.com/religion/encyclopedias-almanacs-transcripts-and-maps/scholastic-method>>

SIEGLER, M. G. *Eric Schmidt: Every 2 Days We Create As Much Information As We Did Up To 2003* [on-line]. [cit. 12. decembra 2021].

Dostupné na internete: <<http://techcrunch.com/2010/08/04/schmidt-data/>>

SIMON, H. A. *Artificial Intelligence: An Empirical Science*. In: *Artificial Intelligence*. 1955, 77, č. 2.

SIMON, H. A. *The Shape of Automation for Men and Management*. New York: Harper & Row, 1965.

SMITH, R. *Why No Science?* [on-line]. [cit. 6. novembra 2021].

Dostupné na internete: <<https://www.thecatholicthing.org/2021/11/02/why-no-science/>>

SPARROW, R. *Killer Robots*. In: *Journal of Applied Philosophy*. 2007, 24, č. 1, s. 62–77.

STRAHOVNIK, V. *Virtues and transhumanist human enhancement*. In: PETROUŠEK, R.,

ŽALEC, B. *Transhumanism as a Chalange for Ethics and Religion*. 2021, s. 37 – 44.

Stratégia kybernetickej obrany Slovenskej republiky. [on-line]. Ministerstvo obrany SR, Vojenské spravodajstvo. [cit. 17. marca 2022].

Dostupné na internete: <<https://www.slov-lex.sk/legislativne-procesy/-/SK/dokumenty/LP-2022-128>>

SUMMER, T. *The first AI universe sim is fast and accurate—and its creators don't know how it works*. [on-line]. [cit. 6. augusta 2020].

Dostupné na internete: <<https://phys.org/pdf480780725.pdf>>

SZEGEDY, CH. et al. *Intriguing Properties of Neural Networks*. In: *Proceedings of the International Conference on Learning Representations*. 2014.

SZONDY, D. *General Atomics' Gambit autonomous combat drone takes the initiative*. [on-line]. [cit. 7. marca 2022].

Dostupné na internete: <<https://newatlas.com/military/general-atomics-gambit-combat-drone-air-dominance/>>

SZONDY, D. *Skyborg combat drone's "brain" flies for the first time*. [on-line]. [cit. 7. marca 2022].

Dostupné na internete: <<https://newatlas.com/military/skyborg-combat-drones-brain-first-time/>>

ŠANTAVÝ, P. *Analýza virtuálneho sveta v kontexte náboženskej formácie, pôsobenia a života Cirkvi* [licenciátska práca]. [on-line]. Bratislava: RKCMBF UK, 2017. [cit. 19. augusta 2020 – 7. decembra 2021].

Dostupné na internete:

<https://peter.santavy.cloud/docs/Analyza_virtualneho_sveta_v_kontexte_nabozenskej_formacie_posobenia_a_zivota_Cirkvi.pdf>

ŠANTAVÝ, P. *Dejiny spásy: zmluvy Starého zákona a Nová zmluva* [diplomová práca]. [on-line]. Bratislava: RKCMBF UK, 2000. [cit. 6. novembra 2021].

Dostupné na internete: <https://peter.santavy.cloud/docs/Dejiny_spasy-zmluvy_SZ_a_NZ.pdf>

ŠANTAVÝ, P. *Niektoré výzvy informačnej spoločnosti v oblasti morálnej teológie*. In *Doctorandum dies 2018: Varia historia et moralia*. RKCMBF UK, Bratislava 2018.

ŠTARHA, Š., GAŠPAROVIČ, R. *AI z pohľadu práva*. [on-line]. [cit. 29. marca 2022].

Dostupné na internete: <<https://www.epravo.sk/top/clanky/ai-z-pohladu-prava-4483.html>>

TANZ, J. *Soon We Won't Program Computers. We'll Train Them Like Dogs*. Wired, May 17, 2016.

The AI arms race. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.ft.com/content/21eb5996-89a3-11e8-bf9e-8771d5404543>>

The AI Powered State. [on-line]. [cit. 26. marca 2022].

Dostupné na internete: <<https://www.nesta.org.uk/feature/ai-powered-state/>>

The „good“ algorithm. [on-line]. [cit. 28. marca 2022].

Dostupné na internete: <https://www.academyforlife.va/content/dam/pav/documenti/%20pdf/2020/Assemblea/Atti_Assemblea_e_28febbraio/Atti_%20completi_PAV_2020_.pdf>

The Nonhuman Rights Project: Frequently Asked Questions. [on-line]. [cit. 10. apríla 2022].

Dostupné na internete: <<https://www.nonhumanrights.org/frequently-asked-questions/>>

The People's Liberation Army Strategic Support Force. [on-line]. [cit. 13. marca 2022].

Dostupné na internete: <<https://jamestown.org/program/the-peoples-liberation-army-strategic-support-force-update-2019/>>

The Price of Privacy: Re-Evaluating the NSA [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <<https://www.youtube.com/watch?v=kV2HDM86Xgl&t=18m>>

The Turing Digital Archive [on-line]. [cit. 30. januára 2021].

Dostupné na internete: <<http://www.turingarchive.org/>>

THURZO, V. *The Influence of Existentialism and Subjectivism on the Concept of the Human Person*. In: *Spiritual and Social Experience in the Context of Modernism and Postmodernism*. Morrisville, 2021.

Top Strategic Predictions for 2020 and Beyond: Technology Changes the Human Condition. [on-line]. [cit. 24. marca 2022].

Dostupné na internete: <<https://www.gartner.com/document/3970846>>

THOMPSON, N. C., GREENEWALD, K., LEE, K., MANSO, G. F. *Deep learning computational cost*. [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <<https://spectrum.ieee.org/deep-learning-computational-cost>>

THURNHER, J. S. *Legal Implications of Fully Autonomous Targeting*. In: *Joint Force Quarterly*. [on-line]. 2012, 67, 4, s. 83. [cit. 7. marca 2022].

Dostupné na internete: <http://ndupress.ndu.edu/Portals/68/Documents/jfq/jfq-67/JFQ-67_77-84_Thurnher.pdf>

Transhumanist Declaration. [on-line]. [cit. 26. januára 2022].

Dostupné na internete: <https://hpluspedia.org/wiki/Transhumanist_Declaration>

TROTT, D. *The hard things are easy, but the easy things are hard*. [on-line]. [cit. 8. januára 2022].

Dostupné na internete: <<https://www.campaignlive.com/article/hard-things-easy-easy-things-hard/1498154>>

Trustworthy AI is human-centered. [on-line]. [cit. 20. februára 2022].

Dostupné na internete: <<https://www.ibm.com/watson/trustworthy-ai>>

TUCKER, P. *SecDef: China Is Exporting Killer Robots to the Mideast*. [on-line]. [cit. 6. decembra 2021].

Dostupné na internete: <<https://www.defenseone.com/technology/2019/11/secdef-china-exporting-killer-robots-mideast/161100/>>

TUCKER, P. *The Pentagon's AI Ethics Draft Is Actually Pretty Good*. [on-line]. [cit. 9. marca 2022].

Dostupné na internete: <<https://www.defenseone.com/technology/2019/10/pentagons-ai-ethics-draft-actually-pretty-good/161005/>>

Understanding China's AI Strategy. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.cnas.org/publications/reports/understanding-chinas-ai-strategy>>

URBINA, F., LENTZOS, F., INVERNIZZI, C. et al. *Dual use of artificial-intelligence-powered drug discovery*. In: *Nat Mach Intell*. [on-line]. 2022, 4, s. 189–191. [cit. 22. marca 2022].

Dostupné na internete: <<https://doi.org/10.1038/s42256-022-00465-9>>

Uznesenie Európskeho parlamentu zo dňa 16. 2. 2017 obsahujúce odporúčania pre Komisiu k normám občianskeho práva v oblasti robotiky (2015/2103(INL)). [on-line]. [cit. 29. marca 2022].

Dostupné na internete: <<https://eur-lex.europa.eu/legal-content/SK/TXT/PDF/?uri=CELEX:52017IP0051&from=EN>>

Vatican Hackathon – harnessing youth, technology to serve common good. [on-line]. [cit. 28. marca 2022].

Dostupné na internete:

<<https://www.vaticannews.va/en/vatican-city/news/2018-03/vatican-hackathon--.html>>

VIGNARD, K. *Manifestos and open letters: Back to the future?* [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://thebulletin.org/2018/04/manifestos-and-open-letters-back-to-the-future/>>

Viktor Frankl citáty. [on-line]. [cit. 25. apríla 2022].

Dostupné na internete: <<https://citaty-slavných.sk/autori/viktor-frankl/>>

‘Wake up’: Ex-Air Force software officer warns China is winning AI battle. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.washingtontimes.com/news/2021/oct/11/nicolas-chaillan-warns-china-winning-ai-battle/>>

Why Advanced Ransomware Is Cybercrime's Most Profitable Business Model [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <<https://blog.knowbe4.com/why-advanced-ransomware-is-cybercrimes-most-profitable-business-model>>

WikiLeaks Releases Trove of Alleged C.I.A. Hacking Documents. [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <<https://www.nytimes.com/2017/03/07/world/europe/wikileaks-cia-hacking.html>>

WOODS, E. T. Jr. *Ako Katolícka cirkev budovala západnú civilizáciu*. Bratislava: Redemptoristi - Slovo medzi nami, 2010.

ZUBOFF, SH. *The real reason why Facebook and Google won't change*. [on-line]. [cit. 21. marca 2022].

Dostupné na internete: <<https://www.theguardian.com/science/2018/jul/18/thousands-of-scientists-pledge-not-to-help-build-killer-ai-robots>>

Twitter:

BEYER, L. *The return of patch-based self-supervision...* [on-line]. [cit. 18. novembra 2021].

Dostupné na internete:

<<https://twitter.com/giffmana/status/1459092079020285976>>

KARPATHY, A. *Computer vision research feels a bit stagnating in a local minimum of 2D texture recognition on ImageNet...* [on-line]. [cit. 10. februára 2022].

Dostupné na internete:

<<https://twitter.com/karpathy/status/1491452689825165314>>

Wikipedia.org:

5-HT_{2A} receptor. [on-line]. [cit. 20. januára 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/5-HT2A_receptor>

Action potential [on-line]. [cit. 27. februára 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Action_potential>

Alan Turing. [on-line]. [cit. 30. januára 2021].

Dostupné na internete:

<https://en.wikipedia.org/wiki/Alan_Turing#Pattern_formation_and_mathematical_biology>

AlphaGo. [on-line]. [cit. 30. júla 2020].

Dostupné na internete: <<https://en.wikipedia.org/wiki/AlphaGo>>

Carnivore [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <[https://en.wikipedia.org/wiki/Carnivore_\(software\)](https://en.wikipedia.org/wiki/Carnivore_(software))>

Common Vulnerabilities and Exposures. [on-line]. [cit. 15. februára 2022].

Dostupné na internete:

<https://en.wikipedia.org/wiki/Common_Vulnerabilities_and_Exposures>

Darknet [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <<https://en.wikipedia.org/wiki/Darknet>>

Dartmouth workshop. [on-line]. [cit. 1. februára 2021].

Dostupné na internete: <https://en.wikipedia.org/wiki/Dartmouth_workshop>

Douglas Hofstadter. [on-line]. [cit. 30. júla 2020].

Dostupné na internete: <https://en.wikipedia.org/wiki/Douglas_Hofstadter>

Deep Blue versus Garry Kasparov. [on-line]. [cit. 30. júla 2020].

Dostupné na internete:

<https://en.wikipedia.org/wiki/Deep_Blue_versus_Garry_Kasparov>

Echelon [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <<https://en.wikipedia.org/wiki/ECHELON>>

Equation Group. [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Equation_Group>

Fourth-generation warfare. [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Fourth-generation_warfare>

Global surveillance [on-line]. [cit. 23. februára 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Global_surveillance>

HermeticWiper: New data-wiping malware hits Ukraine. [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <<https://www.welivesecurity.com/2022/02/24/hermeticwiper-new-data-wiping-malware-hits-ukraine/>>

IsaacWiper and HermeticWizard: New wiper and worm targeting Ukraine. [on-line]. [cit. 10. marca 2022].

Dostupné na internete: <<https://www.welivesecurity.com/2022/03/01/isaacwiper-hermeticwizard-wiper-worm-targeting-ukraine/>>

Kybernetika. [on-line]. [cit. 3. februára 2021].

Dostupné na internete: <<https://sk.wikipedia.org/wiki/Kybernetika>>

Lethal autonomous weapon. [on-line]. [cit. 5. marca 2022].

Dostupné na internete: <https://en.wikipedia.org/wiki/Lethal_autonomous_weapon>

Missionaries and cannibals problem. [on-line]. [cit. 18. novembra 2021].

Dostupné na internete:

<https://en.wikipedia.org/wiki/Missionaries_and_cannibals_problem>

Neuralink. [on-line]. [cit. 5. augusta 2020].

Dostupné na internete: <<https://en.wikipedia.org/wiki/Neuralink>>

Occam's Razor. [on-line]. [cit. 6. novembra 2021].
Dostupné na internete: <https://en.wikipedia.org/wiki/Occam%27s_razor>

Priemyselná revolúcia. [on-line]. [cit. 20. augusta 2020].
Dostupné na internete: <https://sk.wikipedia.org/wiki/Priemyseln%C3%A1_revol%C3%Bacia>

PRISM (NSA) [on-line]. [cit. 23. februára 2022].
Dostupné na internete: <[https://sk.wikipedia.org/wiki/PRISM_\(NSA\)](https://sk.wikipedia.org/wiki/PRISM_(NSA))>

Risk management [on-line]. [cit. 3. marca 2022].
Dostupné na internete: <https://en.wikipedia.org/wiki/Risk_management>

Singularita. [on-line]. [cit. 3. augusta 2020].
Dostupné na internete: <<https://sk.wikipedia.org/wiki/Singularita>>

Social Credit System. [on-line]. [cit. 28. februára 2022].
Dostupné na internete: <https://en.wikipedia.org/wiki/Social_Credit_System>

Stuxnet. [on-line]. [cit. 10. marca 2022].
Dostupné na internete: <<https://en.wikipedia.org/wiki/Stuxnet>>

Tempora [on-line]. [cit. 23. februára 2022].
Dostupné na internete: <<https://sk.wikipedia.org/wiki/Tempora>>

The Shadow Brokers. [on-line]. [cit. 10. marca 2022].
Dostupné na internete: <https://en.wikipedia.org/wiki/The_Shadow_Brokers>

Turek (stroj). [on-line]. [cit. 20. augusta 2020].
Dostupné na internete: <[https://sk.wikipedia.org/wiki/Turek_\(stroj\)](https://sk.wikipedia.org/wiki/Turek_(stroj))>

Turing machine. [on-line]. [cit. 7. apríla 2022].
Dostupné na internete: <https://en.wikipedia.org/wiki/Turing_machine>

Turingov test. [on-line]. [cit. 29. januára 2021].
Dostupné na internete: <https://sk.wikipedia.org/wiki/Turingov_test>

Virtual assistant. [on-line]. [cit. 20. augusta 2020].
Dostupné na internete: <https://en.wikipedia.org/wiki/Virtual_assistant>

Whistleblowing [on-line]. [cit. 23. februára 2022].
Dostupné na internete: <<https://cs.wikipedia.org/wiki/Whistleblowing>>

Wolfgang Kempelen. [on-line]. [cit. 20. augusta 2020].

Dostupné na internete: <https://sk.wikipedia.org/wiki/Wolfgang_Kempelen>

Slovník termínov

Adhocracia	spôsob organizácie spoločenstva ľudí. V <i>ad hoc</i> (k určitému účelu vytvorenej) organizácii je každá organizačná zložka modulárna a disponibilná, každá laterálne interaguje s mnohými inými jednotkami. Rozhodnutia sú prijímané podľa okolností, nie štandardným spôsobom.
AGI	skutočná umelá inteligencia, ktorá je všeobecná (general) a silná (strong). Všeobecná, keďže dokáže zvládnuť akúkoľvek intelektuálnu úlohu a má schopnosť generalizovať, t. j. zovšeobecňovať a prenášať, či adaptovať naučené schopnosti na iné úlohy. Silná, pretože aj skutočne rozumie tomu, či rieši a vykonáva.
Algokracia	skrátенý názov pre vládu podľa algoritmov, algoritmický právny poriadok, algoritmické vládnutie a pod. Ide o alternatívnu formu vlády, resp. spoločenského usporiadania, pri ktorom sa na reguláciu, presadzovanie práva a všeobecne na akýkoľvek aspekt každodenného života, ako je doprava alebo registrácia pozemkov, používajú počítačové algoritmy, najmä umelá inteligencia a blockchain.
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 inteligentné 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.
ASI	Artificial Super Intelligence, t.j. umelá inteligencia, ktorá by bola naprieč všetkými oblasťami inteligentnejšia ako človek. Úvahy o ASI sa mnohokrát spájajú s konceptom singularity v oblasti AI, v ktorej umelá inteligencia so schopnosťou učiť sa, zlepšovať sa a samostatne sa vyvíjať rýchlo dosiahne a následne prevýši

inteligenciu človeka. Vo väčšine diskusií sa ASI nespomína samostatne, je vnímaná ako podmnožina AGI.

Cyberspace	virtuálny životný priestor, často nazývaný aj kybernetický priestor, či virtuálny svet, je fiktívny priestor, v ktorom prebieha komunikácia medzi jednotlivými subjektami, pričom tieto sú celkom reálne.
Data mining	(dolovanie dát) je analytický proces navrhnutý na skúmanie veľkého množstva dát podľa konkrétnych vzorov a väzieb. Výsledkom sú množiny dát, splňujúcich žiadané podmienky, prípadne predikcia ďalšieho vývoja, správania, atď.
Digital Age	digitálny vek – obdobie rozvoja spoločnosti po druhej priemyselnej revolúcii, kedy sa na základe využívania IKT spoločnosť mení na spoločnosť založenú na vedomostiach.
Digital divide	digitálne rozdelenie ako dôsledok informačnej nerovnosti, ktorá prerastá do digitálnej chudoby a má evidentné dôsledky na život ľudí.
Digitálna identita	je súbor informácií v elektronickej forme, ktoré reprezentujú subjekt, ktorým môže byť osoba, organizácia alebo objekt reálneho sveta.
eGovernment	elektronická forma výkonu verejnej správy prostredníctvom informačno-komunikačných technológií (IKT). eGovernment je využitie prostriedkov a nástrojov informačných technológií (najmä siete Internet) na skvalitnenie verejných služieb pre občanov, podnikateľov a celú spoločnosť.
eHealth	zdravotnícka starostlivosť podporovaná elektronickými procesmi, informatizáciou a prostriedkami IKT.
Global village	vnímanie spoločnosti ako spoločenského priestoru prirovnávaného celosvetovej dedine. Ide o dôsledok enormného rozvoja využívania prostriedkov IKT, ktoré sú masovo prijímané populáciou, presahujú geografické i kultúrne vzdialenosti a pretvárajú vnímanie spoločnosti, resp. sveta ako jedného spoločného celku.
Haktivizmus	označuje používanie prostriedkov IKT na podporu politickej, či ideologickej agendy. I keď mnohokrát ide o informačnú analógiu protestných akcií, či demonštrácií z reálneho sveta, existujú i početné

prípady podvratného, či neetického využitia IKT v presadzovaní politickej agendy. Cieľom hacktivizmu často krát býva reálna, či falošná sloboda slova, ľudské práva a sloboda informácií.

Internet

je celosvetový systém prepojených počítačových sietí na základe štandardizovaných pravidiel (súbor protokolov TCP/IP,...). Internet v súčasnosti navzájom prepája milióny rôznych počítačových sietí a v nich prakticky miliardy zariadení, pričom zastrešuje neustále rastúci počet informačných zdrojov a služieb.

Sieťová neutralita

je princíp, podľa ktorého by poskytovatelia internetových služieb a vlády mali zaobchádzať so všetkými dátami na internete rovnako, nediskriminujúc, či nepreferujúc podľa užívateľa, obsahu, webovej stránky, platformy, aplikácie, typu pripojeného zariadenia, alebo spôsob komunikácie.

Singularita

pojmem singularita sa v mnohých oblastiach matematiky a fyziky používa na vyjadrenie stavu, ktorý je zvláštny, nedefinovaný, odchyľujúci sa z trendu, či množiny, stav prekračujúci očakávané parametre, nekonečný, či v daných podmienkach neriešiteľný (napr. aktuálne fyzikálne modely, ktorými nedokážeme popísať singulárne body, akými sú napr. Veľký tresk a čierne diery).

Pod termínom technologická singularita sa myslí teoretický bod vo vývoji vedeckej civilizácie (znalostnej spoločnosti), v ktorom sa technologický pokrok zrýchli do nekonečna a prevýši všetky predpovede.

Singularitou v umelej inteligencii je myslený stav, ktorý nastane, ak umelá inteligencia so schopnosťou učiť sa, zlepšovať sa a samostatne sa vyvíjať rýchlo dosiahne a následne prevýši inteligenciu človeka. Teda situácia, keď sa počítačové systémy stanú inteligentnejšími než ľudia.

Šifrovanie

je proces, ktorým sa nezabezpečené elektronické dáta prevádzajú pomocou kryptografie na dáta šifrované, čitateľná len tým, kto pozná dešifrovací kľúč, či pravidlo. Šifrovanie dát slúži k ich ochrane

pred cudzími osobami a používa sa pri ich ukladaní i prenose prostredníctvom telekomunikačných a informačných kanálov.

World Wide Web

(web 1.0) je distribuovaný informačný systém navzájom prepojených hypertextových dokumentov, ktoré sú dostupné prostredníctvom internetu. Distribuovaný znamená, že jednotlivé dokumenty sa môžu nachádzať kdekoľvek na internete. Hypertextové dokumenty sú dokumenty, ktoré obsahujú časti textu s odkazmi (referenciami) na iné dokumenty, umožňujúc tak priamy prístup k odkazovaným informáciám.

Web 2.0 je názov pre súčasné webové technológie charakterizované používateľským vytváraním obsahu a sociálnymi sieťami, širokou technologickou dostupnosťou a webovými aplikáciami, dynamickým i mediálnym obsahom a využívaním v rámci fungovania informačnej spoločnosti.

Web 3.0 by mal byť virtuálnym svetom informácií s vysoko interaktívnym a personalizovaným využitím, ktorý bude okrem iných sofistikovaných technológií (napr. tzv sémantický web, virtuálnu realitu a pod.) v plnej miere využívať aj algoritmy umelej inteligencie.