

Univerzita Hradec Králové

Pedagogická fakulta

Katedra kulturních a náboženských studií

## **Bezpečnostní rizika umělé inteligence v současném světě**

Bakalářská práce

Autor: Anna Ryzáková

Studijní program: TK

Studijní obor: Transkulturní komunikace

Vedoucí práce: Mgr. Petr Macek, Ph.D.

Oponent práce: ThLic. Lukáš de la Vega Nosek, Ph.D.



## Zadání bakalářské práce

**Autor:** Anna Ryzáková

**Studium:** P19P0190

**Studijní program:** B6107 Humanitní studia

**Studijní obor:** Transkulturní komunikace

**Název bakalářské práce:** **Bezpečnostní rizika umělé inteligence v současném světě**

**Název bakalářské práce** Security risks of artificial intelligence in today's world  
**AJ:**

### **Cíl, metody, literatura, předpoklady:**

Bakalářská práce se věnuje bezpečnostním otázkám spojeným s umělou inteligencí (AI) a rizikovosti jejího využití. V práci se nejprve seznamujeme se samotným pojmem umělé inteligence, jejím vznikem a následným vývojem a využitím. Zároveň budou vysvětleny základní principy fungování AI spolu s příklady jejího praktického využití. Práce poukazuje především na nepřehlédnutelná rizika spojená s rozšířením AI do našich každodenních životů. V práci tak vystávají některé problematické otázky současnosti a budoucnosti lidské populace spojené s AI. Dnes je umělá inteligence skoro všude kolem nás, denně nám pomáhá v každodenních úkonech, zastupuje člověka v situacích, kde sám nestačí a usnadňuje mu život na každém kroku. Práce chce primárně poukázat na to, jak je tenká hranice mezi pokrokem v AI a jejími nezanedbatelnými přínosy a na druhé straně riziky a potencionální nekontrolovatelností. Práce se věnuje etickým otázkám AI, zabývá se riziky případného zneužití, ale i širším otázkám, které AI klade ve vztahu k lidské existenci a přirozenosti. Také poukazuje na novodobé trendy, které jsou spjaty s příchodem a nepřetržitým vývojem AI a jejich schopností.

KOLOUCH, Jan. CyberCrime. Praha: CZ.NIC, z.s.p.o., 2016. CZ.NIC. ISBN 978-80-88168-15-7.  
LANIER, Jaron. Komu patří budoucnost?: nejste zákazníkem internetových firem: jste jejich produktem. Přeložil Petr HOLČÁK. Praha: Argo, 2016. Zip (Argo: Dokořán). ISBN 978-80-257-1994-7. CHARVÁT, Martin. O nových médiích, modularitě a simulaci. Praha: Togga, 2017. ISBN 978-80-7476-121-8. LEE, Kai-fu. Supervelmoci umělé inteligence: Čína, Silicon Valley a svět v éře AI. Přeložil Petr HOLČÁK. Praha: Argo, 2019. Crossover. ISBN 978-80-2573-050-8. TEGMARK, Max. Život 3.0: člověk v éře umělé inteligence. Přeložil Markéta IVÁNKOVÁ. Praha: Argo, 2020. Zip (Argo: Dokořán). ISBN isbn:978-80-736-3948-8.

**Garantující pracoviště:** Katedra kulturních a náboženských studií,  
Pedagogická fakulta

**Vedoucí práce:** Mgr. Petr Macek, Ph.D.

**Oponent:** ThLic. Mgr. Lukáš de la Vega Nosek, Ph.D.

**Datum zadání závěrečné práce:** 29.3.2017

## **Prohlášení**

Prohlašuji, že jsem tuto bakalářskou práci *Bezpečnostní rizika umělé inteligence v současném světě* vypracovala samostatně, pod vedením vedoucího bakalářské práce a řádně jsem uvedla všechny použité prameny a literaturu.

**V Hradci Králové dne 29. 3. 2022**

**Anna Ryzáková**

## **Poděkování**

Mé poděkování patří Mgr. Petru Mackovi, Ph.D. za odborné vedení, trpělivost a ochotu, kterou mi v průběhu zpracování bakalářské práce věnoval.

## **Anotace**

RYZÁKOVÁ, Anna. *Bezpečnostní rizika umělé inteligence*. Hradec Králové: Pedagogická fakulta Univerzity Hradec Králové, 2022, 46s. Bakalářská práce

Tato bakalářská práce se věnuje bezpečnostním otázkám umělé inteligence a rizikovosti v jejím rozsáhlém využití v životě člověka a její vliv na něj. V práci se nejprve seznamujeme se samotným pojmem umělé inteligence, jejím vznikem a následným vývojem a využitím. Vysvětluje, jak AI funguje a udává některé z příkladů. Práce hlavně poukazuje na nepřehlédnutelná rizika rozšíření AI do našich každodenních životů.

V práci tak vyvstávají některé z až děsivých otázek současnosti a budoucnosti lidské populace. Dnes je umělá inteligence skoro všude kolem nás, denně nám pomáhá v každodenních úkonech, zastupuje člověka v situacích, kde sám nestačí a usnadňuje nám život na každém kroku. Tato práce chce primárně poukázat na to, jak je tato hranice velmi tenká, jak s každým pokrokem ve vývoji AI stále více balancujeme nad propastí bez kontroly.

Práce se věnuje etickým otázkám AI, zabývá se riziky případného zneužití, ale také okrajově existenciálním úvahám AI v našich každodenních životech. Také poukazuje na novodobé trendy, které jsou spjaté s příchodem a nepřetržitým vývojem AI a jejích schopností.

Práce si klade za cíl AI představit a upozornit na některé z rizikových pilířů umělé inteligence, jemně vykreslit šachovnici a na konkrétních příkladech ze současného světa poukázat na hrozbu změny pravidel hry.

**Klíčová slova:** Umělá inteligence, AI, bezpečnostní rizika, etika, člověk, stroj.

## **Annotation**

RYZÁKOVÁ, Anna. *Security risks of artificial intelligence in today`s world*. Hradec Králové: Faculty of Education, University of Hradec Králové, 2022. 46pp. Bachelor thesis

This bachelor thesis explores the security issues of artificial intelligence and the risks involved in its widespread use in human life and its impact on human life. The thesis first introduces the concept of artificial intelligence itself, its origin and subsequent development and use, explains how AI works and gives some examples. The thesis mainly points out the overlooked risks of the spread of AI into our everyday lives. The work thus raises some of the almost frightening questions of the present and future of the human population. Today, artificial intelligence is almost everywhere around us, helping us in everyday tasks, standing in for humans, in situations where they are not sufficient on their own, and making our lives easier at every advance in AI development we are increasingly teetering over the precipice of no control. The thesis explores the ethical issues of AI, addressing the risks of potential misuse, but also, peripherally, the existential considerations of AI in our lives. It also highlights modern trends that are linked to the advent and continues development of AI and its capabilities. The paper aims to introduce AI and highlight some of the risky pillars of AI, to gently draw the chessboard and to highlight the threat of changing the rules of the game using concrete examples from the contemporary world.

**Keywords:** Artificial intelligence, AI, security risks, ethics, human, machine.

## Obsah

|           |                                           |           |
|-----------|-------------------------------------------|-----------|
| <b>1.</b> | <b>ÚVOD .....</b>                         | <b>9</b>  |
| <b>2</b>  | <b>DEFINICE, VZNIK A VYUŽITÍ AI .....</b> | <b>11</b> |
| 2.1.      | DEFINICE .....                            | 11        |
| 2.2.      | VZNIK AI .....                            | 14        |
| 2.3.      | TURINGŮV TEST .....                       | 15        |
| 2.4.      | VÝVOJ AI .....                            | 16        |
| 2.5.      | VYUŽITÍ AI .....                          | 18        |
| 2.5.1.    | <i>AI pro průmyslovou výrobu .....</i>    | <i>18</i> |
| 2.5.2.    | <i>AI pro dopravu.....</i>                | <i>19</i> |
| 2.5.3.    | <i>AI ve zdravotnictví.....</i>           | <i>20</i> |
| 2.5.4.    | <i>AI pro komunikaci.....</i>             | <i>21</i> |
| <b>3.</b> | <b>BEZPEČNOSTNÍ RIZIKA AI .....</b>       | <b>22</b> |
| 3.1.      | KYBERPROSTOR .....                        | 22        |
| 3.2.      | SÍTĚ VIRTUÁLNÍ REALITY .....              | 24        |
| 3.3.      | AI A MOC .....                            | 25        |
| 3.3.1.    | <i>Kyberterorismus .....</i>              | <i>25</i> |
| 3.3.2.    | <i>Boti a trollové.....</i>               | <i>25</i> |
| 3.3.3.    | <i>Autonomní zbraně.....</i>              | <i>27</i> |
| <b>4.</b> | <b>ETICKÁ VÝCHODISKA AI.....</b>          | <b>29</b> |
| 4.1.      | DILEMATA INFORMAČNÍ ETIKY.....            | 29        |
| 4.1.1.    | <i>Svoboda slova a cenzura .....</i>      | <i>30</i> |

|           |                                                   |           |
|-----------|---------------------------------------------------|-----------|
| 4.1.2.    | <i>Dezinformace</i> .....                         | 30        |
| 4.1.3.    | <i>AI a důvěra uživatelů</i> .....                | 32        |
| <b>5</b>  | <b>AI A ČLOVĚK</b> .....                          | <b>34</b> |
| 5.1.      | LIDSKÁ MYSL A POČÍTAČ.....                        | 35        |
| 5.2.      | UMĚLÝ SMYSLUPLNÝ ŽIVOT, UMĚLÉ SMYSLUPLNÉ JÁ ..... | 36        |
| 5.3.      | TRANSHUMANISMUS .....                             | 38        |
| 5.3.1.    | <i>Člověk, stroj a tělesnost</i> .....            | 39        |
| <b>6.</b> | <b>ZÁVĚR</b> .....                                | <b>40</b> |
|           | <b>POUŽITÉ ZDROJE</b> .....                       | <b>41</b> |



## 1. Úvod

Technologie a jejich užití s příchodem umělé inteligence (AI) nabírá nových rozměrů. Systémy umělé inteligence nás dnes obklopují na každém kroku. Denně nám pomáhají a zastupují nás v činnostech, na které už nestačíme a její různé formy nám usnadňují naše životy. Mění a tvoří naše chápání nás samotných a světa kolem. AI už je s našimi životy nenávratně spjata, do tohoto vlaku už jsme nastoupili. V diskurzu o AI by mělo být klíčovou otázkou, jak zajistit, aby se umělá inteligence nestala hrozbou pro lidstvo.

Toto téma jsem si vybrala právě proto, že je více nežli aktuální. Vývoj umělé inteligence jde neustále dopředu a provázanost s našimi životy bude do budoucna pouze narůstat. Téma je stále převážně předmětem diskusí a minimálně pro širokou veřejnost je spíše jednou velkou hádankou. Práce si klade za cíl provést analýzu umělé inteligence s pomocí několika příkladů vlivu AI na člověka a snaží se na jejich základě vykreslit alespoň základní bezpečnostní rizika, která můžeme v jejím vývoji a využití zaznamenat v současnosti, nebo historicky.

V první a druhé kapitole se práce věnuje samotné definici umělé inteligence a s ní spojených pojmů. Samotná definice obsahuje mnoho dilemat. Upozorňuje na problematiku samotného pojmu inteligence (praktická a intelektuální) a jejího výkladu, který je důležitý pro následující vývoj AI. Pokládá si otázky, zda by byla AI schopna dosáhnout úrovně inteligence o které mluví například Coreth. Rozmach inteligentního chování do našich životů si zmapujeme několika hlavními mezníky, kterými AI musela projít, aby dostala takovou podobu, jakou známe dnes a představíme si několik odvětví, kde AI lze potkat.

Třetí kapitolou jsou rozpracovaná bezpečnostní rizika, která nejsou pojata jako roboti, kteří se pokouší o zkázu lidstva. Otvírají se nová pole politiky, ekonomiky a sociálně – kulturní povahy. Ta mění světový slet událostí. Práce předkládá rizika jako, neznalost uživatele na druhé straně sítě, pokud tam nějaký je, upozorňuje na tvorbu nových identit a hodnot, které se pohybují až na globální úrovni. A to díky propojenosti sítí ve virtuálním světě a celkové blízkosti světa. Také se snaží ukázat, na jednotlivých příkladech, že uživatel je ten, kdo je nejslabším článkem, ale zároveň tím, kdo dělá z AI zbraň.

Práce se pokouší do diskurzu o AI přiložit i vhled humanitních věd. Proto se poslední dvě kapitoly věnují etickým a filozofickým úvahám. Eticky se pokouší zhodnotit některá současná bezpečnostní rizika a upozorňuje na to, že to není o tom, zda budeme žít po boku s AI,

ale otázka zní jak. Primárně si klade otázku, jak přistupovat k vývoji AI, aby neohrožovala svobodu člověka. A jaké otázky si vlastně vůbec pokládat.

V poslední kapitole se dívá na samotný obraz člověka a umělé inteligence. Oba si jednotlivě prošli řadou různých fází, ale dohromady tvoří do budoucnosti představu tzv. transčlověka. Tuto kapitolu zakládá na otázkách lidské inteligence. Byl by člověk vůbec schopný stvořit sám sebe? Nebo svojí lepší verzi? Fungoval by mozek vylepšený technologickými implantáty stejně jako člověk bez nich? Byl by takový člověk stále člověkem? Vycházel by stále ze své lidské přirozenosti? Práce se nesnaží nutně na všechny odpovědět, ale snaží se poukázat na velmi tenkou hranici, která vzniká nedostatečným reflektováním AI a jejího vývoje. Abychom nebalancovali nad propastí bez kontroly.

## 2 Definice, vznik a využití AI

V úvodní kapitole si vyjasníme důležité pojmy, vznik a historii umělé inteligence, které s tématem úzce souvisí a pomohou nám k lepšímu pochopení tématu této práce. Na samotném konci vývoje AI je představa vyvinuté umělé inteligence, která nebude naprogramována od jednoho bodu k druhému, ale která už bude schopna se sama spravovat a dospět samostatně k naplnění zadaného cíle. V takové fázi ještě nejsme, a i když dnes můžeme mluvit o neuronových sítích či metodách strojového učení, tak se dnes AI využívá spíše ke klasifikaci dat a k věcem, na které člověk sám už nemá kapacitu.

### 2.1. Definice

Definice AI by mohla znít nějak takto: Umělá inteligence je tvořena systémy, které jsou schopny samy napodobit inteligentní lidské chování, samostatně vyhodnocovat a promýšlet jednotlivé možnosti, zvolit je a zhodnotit si následující kroky, které je dovedou nebo zvyšují šanci ke zdárnému dosažení cíle. Tyto systémy, které stojí na umělé inteligenci jsou pouze softwarové a působí ve virtuálním světě. Oblíbená představa tvůrců kultovních sci-fi filmů je obraz umělé inteligence v hardwarových zařízeních jako jsou například roboti, auta či drony. Ale jedná se pouze o její prezentaci, vtělení, zprostředkování skrze tyto prostředky do reálného světa. Tento obraz z filmů předběhl reálný vývoj, ale utvořil jasnou představu pro širokou veřejnost. I když by se tato vize sci-fi filmů nebo příběhů z dávnější historie lidstva dala spíše vysvětlit jako analýza, metafora lidstva a lidských vlastností a tato robotická stvoření vnímat předněji jako představitele utlačovaných menšin.<sup>1</sup>

Definice AI, existuje hned několik, ale ani odborníci se nemohou shodnout na přesné definici. Z toho asi také pramení, proč je široká veřejnost ohledně tohoto tématu plná mystifikací. Definice dle Oxfordského slovníku zní: „Schopnost počítačů nebo jiných strojů vykazovat nebo simulovat inteligentní chování; studijní obor, který se toho týká. Zkráceně AI“<sup>2</sup>. Jednou z těch nejznámějších, je definice Marvinů Minského, který je jeden ze zakladatelů nové vědy: „Umělá

---

<sup>1</sup>Elements of AI: <https://course.elementsofai.com/cs/1/1> [online]. Internet [cit. 2022-03-13]. Dostupné z: <https://course.elementsofai.com/cs/1/1>

<sup>2</sup> Oxford English Dictionary: <https://www.oed.com/viewdictionaryentry/Entry/271625>: Artificial intelligence. *Oxford English Dictionary* [online]. Oxford University press [cit. 2022-03-02]. Dostupné z: <https://www.oed.com/viewdictionaryentry/Entry/271625>

intelligence je věda o vytvoření strojů nebo systémů, které budou při řešení určitého úkolu užívat takového postupu, který – kdyby ho dělal člověk – bychom považovali za projev jeho intelligence.“<sup>3</sup> Pro samotnou definici je ale důležité si vysvětlit několik pojmů.

Zatímco s definicí „umělé“ to je o něco jednodušší, tak s pojmem intelligence to je komplikované. Použiji zde velmi stručnou definici Maxe Tegmarka: „Intelligence, tedy schopnost dosáhnout komplexních cílů<sup>4</sup>. Lidská intelligence je oproti té umělé neobyčejně široká, umíme si osvojovat veliké rozpětí nejrůznějších dovedností. Když tedy srovnáme člověka se strojem, tak stroje zvítězí v úzkých odvětvích daných oblastí. „Mnozí tvrdí, že jde o dar, který nelze vysvětlit či mechanizovat. Tyto pozice lze hájit neustálým smršťováním definic. Také jakmile je nějaký proces mechanizován či jakkoli vysvětlen, tak se nám vše zdá jednoduché a velmi málo inteligentní.“<sup>5</sup> V první řadě je velmi důležité si neplést inteligenci s inteligentním chováním.

„Chceme-li hovořit o inteligenci, jde zásadně jen o tu praktickou, ne o intelektuální poznání ve smyslu teoretického nahlédnutí a obecného postižení smyslu a podstaty“.<sup>6</sup> Takhle Coreth popisuje inteligenci v kontrastu s živočichem. Je důležité rozlišovat lidskou inteligenci od té umělé. Cílem vývoje AI je vytvoření toho nejkomplexnějšího umělého bytí, tedy aby svými schopnostmi dokázal dosáhnout jakéhokoli cíle skoro tak dobře jako člověk. To nazýváme AGI, tedy *Artificial general intelligence*. Ta by měla mít schopnost učit se a plnit zadané úkoly, bez neustálé kontroly uživatele.

Za to pro Coretha, není intelligence pouze schopnost komplexně řešit úkoly, ale schopnost člověka, který díky rozumu může poznat vnitřní, metafyzické principy skutečnosti. Vliv vnějších jevů, tedy něco vidět, slyšet a osahat, vědět, že něco vidíme, slyšíme a hmatatelně poznáváme, námi zároveň prostupuje do vnitřního vědomí. Myslící člověk ve své podstatě přesahuje dimenzi hmotného bytí, má určitou schopnost, která svým bytím už nepatří ke stupni

---

<sup>3</sup> MAŘÍK, V. Umělá intelligence 1. 1. vyd. Praha: Academia, 1993. 264 s. ISBN 80-200-0496-3. Kapitola 1, Některé charakteristiky umělé intelligence, s. 17–18.

<sup>4</sup> TEGMARK, Max. *Život 3.0: člověk v éře umělé intelligence*. Přeložil Markéta IVÁNKOVÁ. Praha: Argo, 2020. Zip (Argo: Dokořán). ISBN 978-80-257-3173-4, s. 31.

<sup>5</sup> NICOLA, Ubaldo. *Obrazové dějiny filozofie*. V Praze: Euromedia Group - Knižní klub, 2006. Universum (Knižní klub). ISBN 80-242-1578-0, s. 560.

<sup>6</sup> CORETH, Emerich. *Co je člověk?: Základy filozofické antropologie*. Přeložil Bohuslav VIK. Praha: Zvon, 1994. ISBN 80-7113-098-2, s. 63-64.

hmotného jsoucna, ale stává se bytím vyššího druhu, touto schopností je myšlen rozum, který je nemateriální a duchovní.<sup>7</sup> Vize takto silné AI, která by byla skutečně inteligentní a uvědomovala by si vůbec sama sebe, je zatím otázkou vzdálené budoucnosti, nicméně je otázkou, zda AI vůbec bude moct dosáhnout úrovně lidské inteligence o které mluví Coreth. My se momentálně nacházíme ve fázi slabé inteligence, která pouze vykazuje prvky inteligentního chování. Dalším důležitým pojmem je úzká neboli narrow AI, která má zadáný jediný cíl. Nelze ji vložit do jiného kontextu.

Pro AI jsou charakteristické pojmy – autonomní nebo adaptabilní. Ty se učí ze svých minulých pohybů a chyb a je schopna je vyhodnocovat bez svého uživatele. Pod strojovým učením je pojem, který se jmenuje **Deep Learning** neboli hluboké učení. Jedná se o samostatné učení bez přímého naprogramovaného zadání. Funguje na principu uměle vytvořených neuronových sítí, které se velmi podobají těm, které má mozek lidský. „Zjednodušeně řečeno, aplikace hlubokého učení nejsou programovány, ale jsou cvičeny na skutečných velkých datech, jak se v různých situacích chovat. Ani to ale není jednoduché, protože jsou náchylné na chybovou interpretaci dat, takže se neobejdou bez týmu zkušených specialistů.“<sup>8</sup>

Existují nejméně 4 typické oblasti, kde můžeme hluboké učení potkat. A probíhá v nich patrný souboj špičkových vědeckých pracovišť, skoro každý den lze zaznamenat významný pokrok<sup>9</sup>. První oblastí je **porozumění textu**, kdy už AI pomalu začíná chápat obsah daného textu.<sup>10</sup> A totéž u **mluveného slova**. Další oblastí je **rozpoznání obličeje** – stačí jediná fotka na Facebooku (či jinde), k níž někdo připsal jméno osoby. Facebook zná obličej milionů lidí a disponuje tak, největší podobnou databází na světě. V tomto se mu ani žádná státní bezpečnostní služba nevyrovná.<sup>11</sup>

---

<sup>7</sup> CORETH, Emerich. *Co je člověk?: Základy filozofické antropologie*. Přeložil Bohuslav VIK. Praha: Zvon, 1994. ISBN 80-7113-098-2, s. 83

<sup>8</sup> BRDIČKA, Bořivoj. Co dokážou stroje schopné hlubokého učení. Metodický portál: Spomocník [online]. 25. 04. 2016, [cit. 2022-02-01]. Dostupný z WWW: <https://spomocnik.rvp.cz/clanek/20855/CO-DOKAZI-STROJE-SCHOPNE-HLUBOKEHO-UCENI.html> ISSN 1802-4785.

<sup>9</sup> BRDIČKA, Bořivoj. Co dokážou stroje schopné hlubokého učení. Metodický portál: Spomocník [online]. 25. 04. 2016, [cit. 2022-02-01]. Dostupný z WWW: <https://spomocnik.rvp.cz/clanek/20855/CO-DOKAZI-STROJE-SCHOPNE-HLUBOKEHO-UCENI.html> . ISSN 1802-4785.

<sup>10</sup> LECUN, Yann; BENGIO, Yoshua; HINTON, Geoffrey. Deep learning 2015. [28.1.2022]. Dostupný z <http://www.cs.toronto.edu/~hinton/absps/NatureDeepReview.pdf>

<sup>11</sup>FACEBOOK BUYS ISRAELI FACIAL RECOGNITION FIRM FACE. COM: <https://www.bbc.com/news/technology-18506255> [online]. BBC: Oxford University press, 2012, 19 červen 2012 [cit. 2022-03-23]. Dostupné z: <https://www.bbc.com/news/technology-18506255>

A posledním odvětvím je **simulace inteligentního chování** – nedávno oznámil tým čínské University of Science and Technology spolu s kolegy z výzkumné laboratoře Microsoftu pokoření další mety. Jejich „inteligentní“ agent umí udělat IQ test lépe než průměrný člověk. Podle daných měřítek to tedy znamená, že jeho IQ je větší než 100. Tento agent je specializován jen na tuto jedinou úlohu a nic jiného neumí. Přesto si musí osvojit velké množství schopností.<sup>12</sup>

Krom hlubokého učení je důležitým pojmem tzv. **datová věda**. Pod tímto pojmem si lze představit od statistiky až po některé aspekty počítačových věd. Věnuje se datové vědě, která v sobě zahrnuje AI aspoň částečně. Na vrcholu umělé inteligence stojí pojem robotika, která v sobě kombinuje všechny oblasti AI.

## 2.2. Vznik AI

AI nás dnes doprovází na každém kroku ať už se chceme co nejrychleji a nejpohodlněji dostat na druhou stranu města, přeložit text z jednoho jazyka do druhého, porozumět mu, odemknout si telefon pomocí obličeje, analyzovat akciový trh nebo pomoci při vyhodnocení diagnóz různých nemocí. Počítače musí být stále častěji schopny interakce s okolním prostředím, a to zejména s člověkem. Tyto důvody mají za následek zvětšující se zájem o výzkum na poli počítačového rozpoznávání objektů.<sup>13</sup> Jako obor vznikla AI v polovině 20. století, ale jak nám je známo, tak ta představa je mnohem starší, zaznamenáváme ji hlouběji v historii lidstva. V různých přiběžích, z různých dob a v různém pojetí. Ale jak jsme se do této fáze dostali?

---

<sup>12</sup> BRDIČKA, Bořivoj. Co dokážou stroje schopné hlubokého učení. Metodický portál: Spomocník [online]. 25. 04. 2016, [cit. 2022-02-01]. Dostupný z WWW: <https://spomocnik.rvp.cz/clanek/20855/CO-DOKAZI-STROJE-SCHOPNE-HLUBOKEHO-UCENI.html> ISSN 1802-4785.

<sup>13</sup> VLACH, Jan a Přinosil Jiří. Lokalizace obličeje v obraze s komplexním pozadím. [online]. 11.4.2007 [cit. 28.1. 2022]. Dostupné z: <http://www.elektrorevue.cz/cz/clanky/zpracovani-signalu/0/lokalizace-obliceje-v-obraze-s-komplexnim/>

## 2.3. Turingův test

Velkým krokem pro AI byl tzv. Turingův test, Alan Turing roku 1950 ve své esejí „Computing Machinery and Intelligence“ popsal některé spekulativnější aspekty svých objevů v oblasti AI a jako první myslitel si položil vůbec poprvé otázku, zda počítače mohou myslet.<sup>14</sup> Turingův test funguje tak, že je tazatel a další dva hráči, jeden je počítač a druhý člověk, ti si v chatu vyměňují zprávy, které mají dovést tazatele k odhalení počítače a člověka. Pokud nelze rozpoznat v probíhajícím rozhovoru člověka od počítače, tak stroj dosáhl lidské inteligence. Přičemž klade důraz na otázku, co stroj vůbec je a co znamená samotné myšlení.

Pro Turinga se stroj musí skládat z „paměti (součástí je i kniha pravidel či tabulka instrukcí), řídicí jednotky a kontroly“, takový popis stroje byl na tu dobu možná trochu naivní, ale vedl k tomu, že Turing ve své esejí nabádá k možnosti konstrukce digitálních počítačů, počítačů s neomezenou kapacitou.<sup>15</sup> Turingova definice myšlení je oproti definici stroje mnohem kontroverznější. Pro Turinga je myšlení stejné jako uspět v imitační hře. To vedlo k tomu, že je dnes Turingův test již považován spíše za překonaný. Cílem AI totiž není stroj, který by byl umělým člověkem, ale pouze něco, co by jako člověk mluvilo, aby dokázal zmást ostatní.<sup>16</sup>

Do kontrastu připomínám Corethův popis lidského myšlení. Coreth lidské myšlení chápe v první řadě jako poznávání světa kolem, a to pojmově. Lidskému myšlení je vlastní schopnost tvořit z konkrétní zkušenosti obecné významové obsahy. Myšlení, je možné pouze v prostoru svobody ducha, který se odpoutává od „zde a nyní“ a tvoří svět myšlenky.<sup>17</sup>

Tato evoluce AI neproběhla během jedné noci a nejsme na jejím konci, lidé programují a vylepšují čím dál tím více výkonné algoritmy, které neustále modernizují a vytváří výkonnější verze AI. Vše nějakou dobu trvá a my tuto revoluci umělé inteligence můžeme rozdělit na čtyři vlny vývoje. Každá z těchto vln využívá schopnosti umělé inteligence jinak a každá posune

---

<sup>14</sup> TVRDÝ, Filip. *Turingův test: filozofické aspekty umělé inteligence*. Praha: Togga, 2014. Scholia (Togga). ISBN 978-80-7476-043-3, s. 15.

<sup>15</sup> TVRDÝ, Filip. *Turingův test: filozofické aspekty umělé inteligence*. Praha: Togga, 2014. Scholia (Togga). ISBN 978-80-7476-043-3, s. 26-30

<sup>16</sup> TVRDÝ, Filip. *Turingův test: filozofické aspekty umělé inteligence*. Praha: Togga, 2014. Scholia (Togga). ISBN 978-80-7476-043-3, s. 26-30.

<sup>17</sup> CORETH, Emerich. *Co je člověk?: Základy filozofické antropologie*. Přeložil Bohuslav VIK. Praha: Zvon, 1994. ISBN 80-7113-098-2, s. 78-80

umělou inteligenci hlouběji do našich každodenních životů. Jakou roli v tom vlastně hrají motivy a pohnutky odborníků? Ivan M. Havel uvádí hned čtyři případné typy motivačních zdrojů, které vedly k výzkumu AI: motivační, inženýrská, badatelská, matematická a přirozená lidská hravost.<sup>18</sup>

## 2.4. Vývoj AI

První vlnou AI se nazývá **internetová**. Ta je nám všem dobře známa, pocítujeme ji pokaždé, když otevřeme jakoukoli internetovou platformu nebo když se nás zmocní pocit, že internetové obchody vědí ještě dříve, co chceme koupit, než my sami. Tato AI funguje na principu algoritmů. Ty pozorují naši historii vyhledávání, tedy na co klikáme, co kupujeme, které sledujeme nebo kolik času na daných stránkách strávíme. Na základě těchto osobních preferencí nám vypočítává a předkládá, co sama najde a u toho využívá primárně toho, že jste člověk.<sup>19</sup> Tahle vlna začala zhruba před více jak 15 lety, a nejen že nám vygeneruje to, co vidět chceme, ale také nám ukazuje, v jakém rozsahu je AI schopna se naučit o našich potřebách prostřednictvím nasbíraných dat. Tato data umožňují jejím vlastníkům, tedy obrovským, nadnárodním internetovým společnostem dosahovat obrovských zisků. „Tytéž metody si v odlišném kontextu osvojila i firma Cambridge Analytica, která využila dat z Facebooku k lepšímu pochopení a ovlivňování amerických voličů během kampaně před prezidentskými volbami roku 2016.“<sup>20</sup>

Druhou vlnou je **podniková** AI, ta zase zužitkovává skutečnost, že tradiční společnosti automaticky přiřazují atributy obrovskému množství dat už celá tisíciletí.<sup>21</sup> Funguje na základě probraných dat o našich rozhodnutích v minulosti a tvoří tak typické rysy a smysluplné výsledky. Například, když si v nemocnicích uloží záznamy o svých pacientech a jejich onemocněních, diagnózách, ale i o „mírách přežití“ a následně nabídnou svá doporučení. Tato

---

<sup>18</sup> HAVEL, Ivan M.. Přirozené a umělé myšlení jako filozofický problém [online]. Glosy.info, 3. prosince 2004. [cit. 30.1.2022] Dostupné z: ISSN 1214-8857.

<sup>19</sup> FORUM Magazín Univerzity Karlovy: Dataholka: Jak sociální sítě zneužívají toho, co nás dělá lidmi. [www.ukforum.cz](http://www.ukforum.cz) [online]. Web Univerzity Kralovy, 2021, 23 listopad 2021 [cit. 2022-03-23]. Dostupné z: <https://www.ukforum.cz/rubriky/alumni/8132-jak-socialni-site-zneuzivaji-toho-co-nas-dela-lidmi>

<sup>20</sup> LEE, Kai-fu. *Supervelmoci umělé inteligence: Čína, Silicon Valley a svět v éře AI*. Přeložil Petr HOLČÁK. Praha: Argo. Crossover. ISBN 978-80-257-3050-8, s. 136-137.

<sup>21</sup> LEE, Kai-fu. *Supervelmoci umělé inteligence: Čína, Silicon Valley a svět v éře AI*. Přeložil Petr HOLČÁK. Praha: Argo. Crossover. ISBN 978-80-257-3050-8, s.140.



podniková AI hledá v těchto databázích vzájemné vztahy, podobnosti, kterých si člověk běžně není schopen všimnout a může tak jít v této vlně i o víc než jen o peníze. Stále však pracují pouze s člověkem zadanými digitálními informacemi.

Třetí vlna posune vývoj AI o několik kroků dál. Na scénu přichází tzv. **strojové učení**, tím se myslí systémy, které zlepšují svůj výkon v dané úloze na základě narůstajících zkušeností nebo objemu dat.<sup>22</sup> Učí se rozeznávat jednotlivá slova a rozklíčovat významy celých vět. Například: pro stroj je fotografie pouze seskupení pixelů. S příchodem AI dokážou rozpoznat jednotlivé objekty ve fotografiích a rozdělit je, fotky s kamarádkou Emou, fotky se svým domácím mazlíčkem, fotky mé osoby, a to dokáže nejen u fotografií.

Na základě toho mohla vzniknout například ruská aplikace FindFace, která rozpozná obličej daného člověka a hned nám poskytne jeho sociální síť a ostatní informace o něm, pokud se jedná o hledaného zločince, tak okamžitě informuje příslušné osoby. Dnes se tato aplikace používá skrze tisíce kamer, využívá se na hledání pohřešovaných lidí nebo zločinců, jak ve venkovních, tak vnitřních prostorech. Stejný koncept AI využívají i v Číně: „čínská komunistická vláda k detailnímu sledování své populace a postihování lidí, které režim pokládá za nepřátele. Firma produkuje i software, který je schopný na ulici identifikovat a pro tamní policii označit příslušníky etnické menšiny Ujgurů podle tvaru jejich obličeje.“<sup>23</sup> AI tak například umožňuje a podporuje genocidu Ujgurů.

Jiné aplikace nám zase poví, na jaký druh houby či hada se s největší pravděpodobností právě díváme. S touto vlnou přichází velká expanze AI do našich každodenních životů, svět se začíná rychleji digitalizovat a prostředí se mění na digitální data a smazává hranice mezi online a off-line světem.<sup>24</sup>

---

<sup>22</sup> ELEMENTS OF AI: Související oblasti. <https://www.elementsofai.cz> [online]. [cit. 2022-03-14]. Dostupné z: <https://course.elementsofai.com/cs/1/2>

<sup>23</sup> ČÍNSKÉ KAMERY SE ZDOKONALILY NA ZADRŽENÝCH UJGURECH. Čechy snímají na ulici i úřadech. <https://www.aktualne.cz> [online]. online, 2021, 24.8. 2021 [cit. 2022-03-18]. Dostupné z: <https://zpravy.aktualne.cz/domaci/armada-skoly-i-urad-vlady-cesko-monitoruji-cinske-kamery-pod/r~59ac5ce0da6d11eb94d2ac1f6b220ee8/>

<sup>24</sup> LEE, Kai-fu. *Supervelmoci umělé inteligence: Čína, Silicon Valley a svět v éře AI*. Přeložil Petr HOLČÁK. Praha: Argo. Crossover. ISBN 978-80-257-3050-8, s.147.

Poslední vlnou je **autonomní** AI, jakmile budou stroje schopny slyšet a vidět, budou také schopny se v něm už bezpečně pohybovat a produktivně v něm pracovat. Vzniknou stroje, které budou schopny světu kolem nejen rozumět, ale budou ho i utvářet. Budou schopny se rozhodovat a při jakékoli odchylce, dokáží improvizovaně reagovat. To jim umožňuje strojové hluboké učení se z vlastních zkušeností na základě sesbíraných dat.

## 2.5. Využití AI

Kde všude můžeme umělou inteligenci potkat? Je v mém kávovaru AI? Odpověď není tak jednoduchá, jak by se na první pohled mohlo zdát. V následujících podkapitolách uvedu pouze několik příkladů, kde se můžeme s AI setkat.

### 2.5.1. AI pro průmyslovou výrobu

AI má v průmyslové výrobě obrovský potenciál. Zvyšuje efektivitu a přesnost produkce. Existuje stroj na sklizení jahod, který byl stvořen z důvodu nedostatku pracovníků, přičemž pak pěstitelé museli přihlížet, jak jim veliké množství jahod na poli hnije. Stroje nabývají schopnosti vnímání, která je klíčová při strojovém rozhodování, díky níž můžeme stroji nahradit lidskou práci.

Stroj dnes umí například rozpoznat zralé ovoce od nezralého a podle toho vyhodnotit, zda bude sklizeno nebo necháno na dozrání. Také dokážou rozeznat plevel od obilí, aby mohl postříkat pesticidy jen plevelné rostliny, nebo naopak dodat plodině další živiny. Kamerové a měřicí systémy jsou schopné identifikovat i nemocná zvířata, a to bez lidské kontroly<sup>25</sup> Dnes, v době, kdy 3D tiskárny dovedou vyrobit prototypy čehokoli od kancelářských budov až po mikro mechanické objekty, které jsou ve velikosti zrnka soli. Nebo obrovské průmyslové stroje dokáží montovat součástky do aut a letadel. Čím více strojů nás obklopuje tím větší výzvou však je otázka splnění a správnosti specifikací.<sup>26</sup>

---

<sup>25</sup> BOROVIČKA, Tomáš. Umělá inteligence přispívá k větší udržitelnosti zemědělství. <https://prg.ai/> [online]. 2021, 18.11.2021 [cit. 2022-03-25]. Dostupné z: <https://prg.ai/tomas-borovicka-datamole-ai-umela-inteligence-zemedelstvi/>

<sup>26</sup> TEGMARK, Max. *Život 3.0: člověk v éře umělé inteligence*. Přeložil Markéta IVÁNKOVÁ. Praha: Argo, 2020. Zip (Argo: Dokořán). ISBN 978-80-257-3173-4, s. 86.

### 2.5.2. AI pro dopravu

Jedním z největších témat ohledně AI a dopravy jsou tzv. autonomní vozidla. Jelikož jsou skoro všechny autonehody zapříčiněny chybou lidského faktoru, tak se předpokládá, že samořídící auta ovládaná umělou inteligencí mohou eliminovat alespoň 90% úmrtí na silnicích. Podle Elona Muska budou samořídící automobily nejen bezpečnější, ale navíc mohou konkurovat třeba Uberu a Lyftu. Případných nehod a problémů není zatím známo mnoho, proto to vypadá, že jsou skutečně lepší než lidští řidiči. Když už se něco přihodí, tak je to v důsledku verifikace (ověření na základě modelů) a validace (ověření na základě reálných vypočítaných výsledků). Aby k nehodám vůbec nedocházelo, je potřeba velmi dobrá kontrola a jasná komunikace mezi člověkem a strojem, což je faktor, který je velmi důležitý a jednou z největších otázek pro vývoj autonomních zbraní.

K první smrtelné nehodě způsobené samo řídicím vozem Tesla v roce 2016, došlo v důsledku dvou chybných předpokladů. Stroj špatně vyhodnotil překážku před ním a za druhé, předpokládal, že jeho řidič dává pozor a v případě potřeby bude schopen zasáhnout.<sup>27</sup> Co se týká autonomních vozidel a dvou hlavních lídrů, kterými jsou Google a Tesla, tak každý z nich symbolizuje a nese dvě naprosto odlišné filozofie. Sestavují týmy, vyvíjejí, testují a sbírají data potřebná pro nasazení autonomních vozů, která mají vyřadit člověka jako řidiče.<sup>28</sup>

Když probíhaly olympijské hry v Tokiu, Toyota měla zakázku na pohyb účastníků mezi místy v olympijské vesnici. Tam se ukázalo, že to, co neumí auto, musí umět okolí. Když auto nerozezná semafor, může on z křižovatky vysílat elektromagnetický nebo zvukový signál, který řekne, zda je zelená, či červená. Můžeme to nazývat jištěním, protože pravděpodobnost, že selžou oči i sluch je malá.<sup>29</sup> Tento příklad autonomních vozidel vykresluje i to, jaká je potřeba kombinace zmíněných technik AI. Nalezení a naplánování nejjednodušší cesty do zadaného

---

<sup>27</sup> TEGMARK, Max. *Život 3.0: člověk v éře umělé inteligence*. Přeložil Markéta IVÁNKOVÁ. Praha: Argo, 2020. Zip (Argo: Dokořán). ISBN 978-80-257-3173-4, s. 87-88

<sup>28</sup> LEE, Kai-fu. *Supervelmoci umělé inteligence: Čína, Silicon Valley a svět v éře AI*. Přeložil Petr HOLČÁK. Praha: Argo. Crossover. ISBN 978-80-257-3050-8, s

<sup>29</sup> MATAS, Jiří. Lídři to s předpovědí nástupu autonomních aut přehnali, říká odborník na umělou inteligenci [online]. 12.12.2021 [28.1. 2022]. Dostupné z: <https://www.e15.cz/rozhovory/lidri-to-s-predpovedi-nastupu-autonomnich-aut-prehnali-rika-odbornik-na-umelou-inteligenci-1386151>

cíle, rozeznání překážek a schopnost vozidla pracovat s dynamikou svého prostředí. Tyto kombinace musí všechny fungovat okamžitě aby nedocházelo k nehodám.<sup>30</sup>

### 2.5.3. AI ve zdravotnictví

AI má obrovský potenciál zlepšit zdravotní péči. Dokumentace v digitální formě usnadnila lékařům jednat rychleji a efektivněji. Lékaři mohou získat okamžitou pomoc od svých kolegů z celého světa. Už známe několik příkladů, které nás přesvědčili o tom, že AI dokáže v některých případech určit diagnózu stejně tak dobře, jako lidský lékař a v některých případech dokonce ještě lépe. V posledních letech proběhla řada operací prováděných roboty, při nich se prokázal potenciál strojů být přesnějšími a spolehlivějšími chirurgy, než je sám člověk.<sup>31</sup>

AI také pomáhá při komunikaci mezi doktorem a těžce nemocným člověkem. Velké množství lékařů se cítí nedostatečně připraveno a proškolen na tyto těžké rozhovory s pacienty. Přitom by měli umět zvážit, kdy, jak, a za jakých okolností je prioritou pacientovi sdělit jeho vážný zdravotní stav. Tato komunikace s těžce nemocnými se nazývá zkratkou SIC. Lékaři často SIC aplikují později, než je potřeba, protože nejsou schopni správně vyhodnotit situaci. „Současný pracovní postup se při identifikaci pacientů způsobilých pro SIC opírá o lidský úsudek a manuální úsilí o zahájení SIC, zdokumentování SIC a vyhledání dokumentace SIC v elektronickém zdravotním záznamu (EHR).“<sup>32</sup>

Hybridní pracovní postup umělé inteligence a člověka by využil umělou inteligenci k přesnější identifikaci pacientů způsobilých pro SIC a zefektivnil pracovní postup tím, že by pomohl plnit základní podřadné úkoly, čímž by zajistil, že vážněji nemocní pacienti dostanou včasné SIC, a poskytne lékařům více času a energie, aby se mohli soustředit na kognitivní a emocionální úkoly vyššího řádu, včetně řešení samotného problému.“<sup>33</sup>

---

<sup>30</sup> ELEMENTS OF AI. <https://course.elementsofai.com/cs/1/1>: Jak definovat umělou inteligenci [online]. Internet [cit. 2022-02-25]. Dostupné z: <https://course.elementsofai.com/cs/1/1>

<sup>31</sup> TEGMARK, Max. *Život 3.0: člověk v éře umělé inteligence*. Přeložil Markéta IVÁNKOVÁ. Praha: Argo, 2020. Zip (Argo: Dokořán). ISBN 978-80-257-3173-4, s.89-90.

<sup>32</sup> CHUA, IS, Ritchie, CS & Bates, DW Zlepšení komunikace s vážnými nemocemi pomocí umělé inteligence. *npj číslice. Med.* 5, 14 (2022). <https://doi.org/10.1038/s41746-022-00556-2>

<sup>33</sup> CHUA, IS, Ritchie, CS & Bates, DW Zlepšení komunikace s vážnými nemocemi pomocí umělé inteligence. *npj číslice. Med.* 5, 14 (2022). <https://doi.org/10.1038/s41746-022-00556-2>,

## 2.5.4. AI pro komunikaci

Vliv technologií na komunikaci pocítujeme asi nejvýrazněji ze všech. Ovlivňuje náš svět a způsob jakým v něm interagujeme. Informace ze všech koutů světa máme na dosah a k dispozici. Miliardy lidí se připojují on-line a povídají si napříč hranicemi, mají přístup ke zprávám, kinematografii, literatuře z celého světa. Internet přináší pohodlí, účinnost, přesnost, ekonomický přínos a propojenost světa. Rychlý přesun částí našich životů do uměle vytvořeného prostředí, nemá obdoby. To se promítá do nás a našich životů ať už pozitivním nebo negativním způsobem.

O jejich dopadu už existuje mnoho odborných článků, knih, dokumentů a jiných pramenů. “Digitální technologie narušují vazbu člověka a světa, který zabydluje a který jej drží při životě. Naše spojení se světem je slabší, nikoli pevnější, v důsledku používání takzvaných „komunikačních technologií.“<sup>34</sup> Technologie se začínají zasazovat i do samotného vzdělávacího systému. Děti s nimi přijdou do styku velmi brzy. Čemu se neklade taková pozornost, je dostatečná prevence, informovanost o rizicích jejich užívání. To je důležité proto, aby technologie ve výuce či výchově sloužili primárně jako pomůcka. Tomuto tématu se například věnuje Manfred Spitzer..<sup>35</sup>

---

<sup>34</sup> SPITZER, Manfred, Václav BĚLOHRADSKÝ, Martin LESKOVJAN, Jeremiáš HAVRANU, Jana KARLOVÁ, Martin SOUKUP a Jan D. BLÁHA, SOKOLÍČKOVÁ, Zdenka, ed. Kyber a eko: digitální technologie v enviromentálních souvislostech. Brno: Host, 2019.s.13

<sup>35</sup> SPITZER, Manfred. *Kybernemoc!: jak nám digitalizovaný život ničí zdraví*. Přeložil Iva KRATOCHVÍLOVÁ. Brno: Host - vydavatelství, 2016. ISBN 978-80-7491-792-9. a již zmiňovaná SPITZER, Manfred, Václav BĚLOHRADSKÝ, Martin LESKOVJAN, Jeremiáš HAVRANU, Jana KARLOVÁ, Martin SOUKUP a Jan D. BLÁHA, SOKOLÍČKOVÁ, Zdenka, ed. Kyber a eko: digitální technologie v enviromentálních souvislostech. Brno: Host, 2019.

### 3. Bezpečnostní rizika AI

Už nyní nám je jasné, že v budoucích případných konfliktech mocností, bude rozhodující právě rozvoj a využití technologií jako takových. Důležitou roli v tomto oboru bude hrát právě AI, která bude do budoucna mávat se světovým děním. A to právě díky jejímu vlivu na psychologickou rovinu člověka, ale i svým vojenským využitím. Některé z těchto aspektů můžeme pocítovat už dnes. V následující kapitole si je ve zkratce představíme. Pro usnadnění problematiky a jejího představení použiji ve zkratce alespoň jednu z možností popisu, jak si takový kyberprostor vůbec představit.

#### 3.1. Kyberprostor

Lze k tomu využít představu ledovce, viditelná špička ledovce se nazývá jako tzv. Surface web, kde se nachází prostor pro pohyb běžného uživatele. Zbylé dvě části se nachází pod vodou, nazývají se Deep Web (sociální média, akademické publikace) a Dark Web – dohromady tyto tři části tvoří celý skutečný kyberprostor. Stejně jako s AI je asi důležité si připomenout, že řada nástrojů a prostředků, pomocí kterých mohu páchat kybernetickou či jinou trestnou činnost, mohu zcela legálně páchat i v rámci Surface Webu.<sup>36</sup> Kyberprostor není tvořen jako jeden velký celek, ale je poskládán z komplexních jednotlivých menších sítí. Ty jsou navzájem spojeny pomocí protokolů IP, je tím tedy vytvořen dynamický a neustále proměnný, vyvíjející se systém, který je sice vázán na hardware, ale zároveň vytváří těžko definovatelný a neomezený kyberprostor.

Lídři, kterými jsou beze sporu čínské a americké internetové firmy mají k expanzi na zahraničních trzích diametrálně rozdílné přístupy. V době, kdy budou služby založené na umělé inteligenci pronikat do každého koutu zeměkoule, mohou tyto firmy svádět zástupné konkurenční boje v zemích jako Indie, Indonésie, v částech Blízkého východu a Afriky.<sup>37</sup> Vzhledem k době, kdy tato bakalářská práce vzniká, tedy v době probíhající války o svobodu

---

<sup>36</sup> KOLOUCH, Jan. *CyberCrime*. Praha: CZ.NIC, z.s.p.o., 2016. CZ.NIC. ISBN 978-80-88168-15-7, s. 46-49

<sup>37</sup> LEE, Kai-fu. *Supervelmoci umělé inteligence: Čína, Silicon Valley a svět v éře AI*. Přeložil Petr HOLČÁK. Praha: Argo. Crossover. ISBN 978-80-257-3050-8, s. 170

Ukrajiny potažmo budoucnosti samotné Evropy si můžeme představit příklady velmi aktuální. Právě probíhající konflikt nám ukázal, že se otevřela úplně nová bitevní pole, která jsou použita v historii lidstva úplně poprvé. Nemyslím tím, že by kybernetické útoky neprobíhaly už předtím, ale poprvé je vidíme přímo ve válečném konfliktu.

Bezpečnostní rizika v této práci nejsou založené na představě zlých robotů, kteří nás jednou přemohou, jak si často představuje široká veřejnost – humanoidní robot Sophia<sup>38</sup>. Ale snaží se předložit jiná, aktuálnější, podstatnější a reálnější roviny těchto rizik. Máme nová pole v našich každodenních životech, která hýbají s pojetím světa a člověka v jeho mnoha vrstvách, ale také se pohybujeme na zcela nově vytvořených polích v oblastech politiky – bitevních, ekonomických a sociálně kulturních, která se výrazně mění a spolu s nimi i světový slet událostí.

Tato rovina možná není na první pohled tak úplně vnímána jako ty ostatní, ale je realitou, která znatelně hýbe s děním ve světě a otázky z ní plynoucí jsou více nežli aktuální. Bylo by mylné si ji představovat jako o něco méně nebezpečnou nebo jí neklást takový důraz v diskurzu o AI. To, že postižené třeba přímo nevidíme, nestojí před námi, neznamená, že škoda bude o něco menší.

Timothy Snyder ve své knize *Tyranie* použije příklad, kdy si představíme, že jsme řidiči auta. Je jasné, že nemáme do protijedoucího auta najet, víme, že škoda by byla vzájemná, a to druhého řidiče také nevidíme přímo.<sup>39</sup> Ve virtuálním světě by takový řidič mohl mít pocit, že se může vydat jakýmkoli směrem bude chtít, že své rozhodnutí bude moct kdykoli změnit.<sup>40</sup> Vítejme ve válce těch, kteří mají informace, vítejme ve válce těch, kteří budou prozíravě investovat do vyvinuté technologie, vítejme ve válce, která zatím není vnímána válkou v pravém slova smyslu. Historicky známe případy, kdy lidé vedli útoky do lidí, a to sobě stáli tváří v tvář. Kam až by mohl zajít útok při kterém nevidíme svého protivníka? A jaké budou následky?

---

<sup>38</sup> CNBC. Interview With The Lifelike Hot Robot Named Sophia (Full) | CNBC. YouTube [online]. 2017 [cit. 2022-3-17]. Dostupné z: <https://www.youtube.com/watch?v=S5t6K9iwcdw&t=289s>

<sup>39</sup> SNYDER, Timothy. *Tyranie: 20 lekcí z 20. století*. Přeložil Martin POKORNÝ. V Praze: Paseka, 2017. ISBN 978-80-7432-838-1, s. 70.

<sup>40</sup> CHARVÁT, Martin. *O nových médiích, modularitě a simulaci*. Praha: Togga, 2017. ISBN 978-80-7476-121-8, s. 90.

### 3.2. Síť virtuální reality

Přesouváme naše životy do uměle vytvořené reality, která se těžko definuje a jakkoli ukotvuje. Svět je propojen variací neviditelných **sítí**, které utváří prostor pro nová společenství a identity, jiného charakteru. Síť, jako nervový systém společnosti.<sup>41</sup> Lidé se každý den dostávají do interakce s druhými lidmi napříč hranicemi a kulturami. Hlavním propojovacím prostředníkem jsou **média**, skrze která probíhá tzv. proces mediatizace. S tímto procesem mediatizace a globalizace jsou spojena další rizika a dilemata. Utváří struktury sociokulturní povahy světa, jedince a tak i celých skupin. Tyto světy interagují vzájemně i s fyzickým světem, jejich vliv tedy neprobíhá pouze ve virtuálním světě.

Skupina s jednou ideou může skrze tento přenosný neviditelný svět komunikovat, propojovat či rozdělovat i na druhé straně zeměkoule nebo hned na několika místech najednou.<sup>42</sup> Zároveň ale, tvoří nová pole veřejného prostoru, **transkulturní povahy**, kde dochází k setkávání kultur. To značně mává s pojetím časoprostoru, vědomí virtuální reality se formuluje odlišným způsobem než v každodenním životě. Má vlastní určující zákonitosti, které definuje softwarem. To ale není jedinou hranicí se kterou tato vzájemná blízkost hýbe. Stejně tak i hranice mezi soukromou a veřejnou sférou se stává problematickou. Individuálnímu konzumentovi dává iluzi automatizované autonomie.<sup>43</sup> Také tu lze hovořit o rovině globálních mediálních událostí, a ty nás mohou podporovat ke vzniku určitých společných hodnot.<sup>44</sup> S možnou narůstající potřebou **globálních institucí**, které by aspoň částečně kontrolovaly obsahy v mediálním prostoru.

Vzájemná interakce je spojena i s narůstající potřebou sebeurčení a touze po uznání. To můžeme vidět například na rychlosti za jaké se rozšířila jednotlivá náboženství do všech koutů světa a je schopna jej zmobilizovat. Stejně jako v každodenním životě, tak i v tom virtuálním prostředí se také setkáváme s reprezentacemi ostatních živých bytostí. Ty na sebe mohou brát

---

<sup>41</sup> CHARVÁT, Martin. *O nových médiích, modularitě a simulaci*. Praha: Togga, 2017. ISBN 978-80-7476-121-8, s. 88.

<sup>43</sup> CHARVÁT, Martin. *O nových médiích, modularitě a simulaci*. Praha: Togga, 2017. ISBN 978-80-7476-121-8, s. 90.

<sup>44</sup> BURDA, František. *Za hranice kultur: transkulturní perspektiva*. Brno: Centrum pro studium demokracie a kultury, 2016. ISBN 978-80-7325-402-5, s. 70-80.



určitou roli. František Burda se ve své knize *Za hranice kultur* věnuje například tzv. klamným identitám. „Člověk, který je na cestě uskutečňování své identity v souladu se svou důstojností, může jednoduše uskutečňovat mylný nebo deformovaný obraz své identity a důstojnosti a stejně tak i u druhých“.<sup>45</sup>

### 3.3. AI a moc

Dlouhou dobu bylo trendem získat přístup do co nejvíce systémů, postupně stahovat menší množství dat a v nejvýhodnější moment udeřit. Nejslabším článkem ve světě AI je její uživatel. Vždy je nejjednodušší dostat pomocí umělé inteligence informace z uživatele, ale zároveň je uživatel tím, kdo z AI udělá zbraň. Rozdíl mezi kinetickou válkou (ve fyzickém prostoru – **boots on the ground**) a kybernetickou je paradoxně v jejich přípravě a organizaci. Taková příprava útoků v kybernetickém prostoru trvají týdny, měsíce někdy i roky. Kybernetické útoky mají dvě základní povahy: informační války a útoky na servery (útoky na infrastrukturu atd.). V podkapitolách níže si ve zkratce představíme několik příkladů, které udělali z AI hrozbu.

#### 3.3.1. Kyberterrorismus

V prostředí internetu se jedná o běžnější neletální formy terorismu (lze i v kombinaci s letálními prostředky). Kyberterrorismus je určitě jeden z velkých strašáků 21. století. Je předem plánován a motivován zpravidla politicky či nábožensky, realizován je spíše malými, ne vojensky organizovanými strukturami. Cílem je ovlivnit, co nejvíce, veřejné mínění, což při tak rychlém šíření informačních a komunikačních technologií po celém světě je skutečnou hrozbou.<sup>46</sup>

#### 3.3.2. Boti a trollové

**Boti**, tedy programy, které pro svého správce plní nejrůznější (sběr dat, zpracovává požadavky, rozesílá zprávy, komunikuje s řídicím prvkem) opakované akce. U internetových botů je výhodou jejich rychlost. Jsou mnohem rychlejší než člověk. Boti nevznikly samozřejmě proto, stejně jako AI, aby bylo konáno zlo. Bot legálně funguje jako průzkumník, hesluje internetový

---

<sup>45</sup> BURDA, František. *Za hranice kultur: transkulturní perspektiva*. Brno: Centrum pro studium demokracie a kultury, 2016. ISBN 978-80-7325-402-5, s. 214

<sup>46</sup> KOLOUCH, Jan. *CyberCrime*. Praha: CZ.NIC, z.s.p.o., 2016. CZ.NIC. ISBN 978-80-88168-15-7, s.323

obsah v klasických prohlížečích. Dokáží vyhledat chyby v aplikacích, svou rychlostí dokáže zahltnout servery velkým množstvím požadavků a samozřejmě jsou také jedni ze spolehlivých nositelů *fake news*. To vše zvládne útočník ze svého počítače, ten infikuje druhý počítač virem. Nejčastěji typu worm a trojský kůň.<sup>47</sup> To, co páchají Bot, často nemá hranice. Hackeři díky botům proniknou přes obranu společností, a to i výzkumných a zdravotnických organizací. „Hackeři například nedávno začali používat boty pro útoky na volná očkovací místa pro covid-19 ve stylu prodeje lístků“<sup>48</sup>

**Troll** je oproti tomu člověk, který sedí na jedné straně sítě a schválně píše a rozněcuje debaty na různých platformách. Reaguje na citlivá témata a snaží se tak vyprovokovat emotivní neshody mezi ostatními uživateli. Takový „trolling“ vznikl z mnoha různých důvodů, problémem je, že ho využili i sociotechnici (není přímo kybernetickým útokem, ale je důvodem, proč je taková řada útoků úspěšná<sup>49</sup>). To znamená, že jsou lidé, kteří jsou za trollování placeni a tvoří organizované uskupení, které jsou známé jako tzv. trollí farmy. O jejich vlivu je už mnoho záznamů, ku příkladu finská novinářka je z vlastní zkušenosti ve své knize *Putinovi trollové* popisuje takto: „Trollové dostávali měsíční mzdu za to, že na sociálních sítích vystupovali jako skuteční lidé s názorem, spravovali a aktualizovali své falešné profily, také opěvovali chválou na ruského prezidenta Vladimira Putina a naopak shazovali opozici a USA“<sup>50</sup>.

Nejznámějším příkladem jsou právě **trollí farmy**, které tahají nitkami z Kremlu, právě skrze sociální média představují bezpečnostní hrozbu pro ostatní národy. Rusku se díky tomu daří rozšiřovat metody nezákonného ovlivňování, omezuje svobodu slova a radikalizuje lidi mimo své území. Pro představu si ve zkratce můžeme připomenout americké volby prezidenta v roce 2016. Pozdější výzkumy ukázaly, že si ruská agentura pro výzkum internetu koupila na Facebooku reklamy, ve kterých očerňovala demokratickou kandidátku Hillary Clintonovou.

---

<sup>47</sup> KOLOUCH, Jan. *CyberCrime*. Praha: CZ.NIC, z.s.p.o., 2016. CZ.NIC. ISBN 978-80-88168-15-7, s. 193.

<sup>48</sup> KOROLOV, Maria, Miroslav NEČAS a Saniye ALAYBEYI. Umělá inteligence pro podnikové nasazení. Computerworld [online]. 2021, červenec 2021, 2021(07-08), 8 [cit. 2022-03-19]. Dostupné z: <https://i.info.cz/files/computerworld/6/umela-inteligence-1.pdf>

<sup>49</sup> KOLOUCH, Jan. *CyberCrime*. Praha: CZ.NIC, z.s.p.o., 2016. CZ.NIC. ISBN 978-80-88168-15-7, s. 186.

<sup>50</sup> ARO, Jessikka. *Putinovi trollové: skutečné příběhy z frontových linií ruské informační války*. Přeložil Lenka FÁROVÁ. Praha: N Media, 2020. Edice N. ISBN 978-80-907922-1-0, s.18.

Trollí farma mimo jiné organizovala skupiny, kde se to hýřilo nenávisťnými debatami vůči cizincům a muslimům. Běžný uživatel ani nepozná, že se stal součástí jedné velké trollí hry.<sup>51</sup>

Díky nim se daří v takovém měřítku rozšiřovat další hoaxy (princip pošli to dál), dezinformace a další projevy kyber kriminality, kterých tu je celá řada. Celá tato účinná strategie stojí na tom, že nejslabším článkem celého bezpečnostního systému je a vždy bude člověk, uživatel. Nikdy nebude existovat počítačový systém, který nebude aspoň částečně závislý na svém uživateli. Proto je nejjednodušší získat potřebné informace právě od člověka.<sup>52</sup>

### 3.3.3. Autonomní zbraně

Stroj, který se bude schopen rozhodovat sám za sebe, učinit rozhodnutí bez lidského faktoru. V této podkapitole si ve zkratce nastíníme autonomizaci zbrojního systému. Co se týče autonomizace jakéhokoli stroje je hybným prvkem AI. Výše jsem popsala fungování hlubokého a strojového učení, které umožnilo nový způsob strojového poznání. AI přináší v oblasti autonomních strojů určitě mnoho výhod. Jako je zvýšená úspěšnost operací, zvětšení bitevních polí, díky její schopnosti rychleji, vytrvaleji a pohyblivěji jednat a reagovat.

Nicméně je otázka vývoje autonomních zbraní stále velmi diskutabilní. A to hlavně proto, že následné činy nejsou jakkoli zákoně přímo ukotveny. Měl by být stroj tím, kdo bude rozhodovat o tom, kdo zemře, a za jakých okolností? Proto se několik vědců a organizací rozhodlo zavázat se ke slibu, že se nebudou podílet na vývoji a výrobě strojů, které jsou schopny útočit bez lidského operátora.<sup>53</sup>

Toto gesto je velmi klíčové alespoň do té doby, dokud se vývoj autonomních zbraní jakkoli zákoně neukotví či nepříjde v mezinárodní normu o nepřijatelnosti jejich užívání. Společnosti

---

<sup>51</sup> ARO, Jessikka. *Putinovi trollové: skutečné příběhy z frontových linií ruské informační války*. Přeložil Lenka FÁROVÁ. Praha: N Media, 2020. Edice N. ISBN 978-80-907922-1-0, s. 297-300

<sup>52</sup> KOLOUCH, Jan. *CyberCrime*. Praha: CZ.NIC, z.s.p.o., 2016. CZ.NIC. ISBN 978-80-88168-15-7, s. 186.

<sup>53</sup> THOUSANDS OF LEADING AI RESEARCHERS SIGN PLDGDE AGAINST KILLER ROBOTS. *Support the Guardian* [online]. 2018, 18.7. 2018 [cit. 2022-04-01]. Dostupné z: <https://www.theguardian.com/science/2018/jul/18/thousands-of-scientists-pledge-not-to-help-build-killer-ai-robots>

tím varují před užíváním smrtících autonomních zbraní, a takový krok od vědců a společnosti může pomoci při případném nepodepsání dohody ze strany samotných států.<sup>54</sup>

---

<sup>54</sup> THOUSANDS OF LEADING AI RESEARCHERS SIGN PLEDGE AGAINST KILLER ROBOTS. *Support the Guardian* [online]. 2018, 18.7. 2018 [cit. 2022-04-01]. Dostupné z: <https://www.theguardian.com/science/2018/jul/18/thousands-of-scientists-pledge-not-to-help-build-killer-ai-robots>

## 4. Etická východiska AI

Technologie a umělá inteligence je vlak, do kterého jsme již nastoupili. V Současném světě žijí různé kultury vedle sebe, a otázkou je jakým způsobem. Totéž lze aplikovat i na umělou inteligenci a technologie jako takové. Hlavní otázkou už není, zda umělou inteligenci vůbec uplatňovat v našich každodenních životech (v takovém množství), ale jak jejím vývojem a uplatněním neohrozit svobodu člověka.

Jak do umělé inteligence při jejím vývoji zabudovat respekt k základním etickým hodnotám lidské společnosti? Jak moc je potřeba umělou inteligenci omezovat? Jak minimalizovat zneužití technologií a systému AI, které vede k ohrožení člověka, světa a jeho svobody, podstaty a přirozenosti v něm. „A to, co nemůže šmírovat stát, pro něj zařídí soukromé firmy – od nákupních řetězců nebo internetových obchodů přes vyhledávače, počítačové firmy a sociální sítě až po pojišťovny a banky.“<sup>55</sup>

### 4.1. Dilemata informační etiky

Možnou obranou proti dezinformačním kampaním, které jsou iniciovány z alternativních médií by bylo jejich případné zrušení. To je ale v rozporu s **právem svobody slova** a pohybuje se na hranici cílené cenzury. Toto téma je obecně velmi problematické, protože média často nabízí svému uživateli informace, které nejsou založené na pravdivých faktech (dezinformace – nepravdivá zpráva). A připravují tak uživatele o **právo na informace**. Dalším důležitým právem je **oprávněnost cenzury**, které je komplikované v případě dezinformací a v tomto pojetí svobody slova.<sup>56</sup>

---

<sup>55</sup> SPITZER, Manfred. *Kybernemoc!: jak nám digitalizovaný život ničí zdraví*. Přeložil Iva KRATOCHVÍLOVÁ. Brno: Host - vydavatelství, 2016. ISBN 978-80-7491-792-9, s. 122

<sup>56</sup> STODOLA, Jiří. PRINCIPY A DILEMATA INFORMAČNÍ ETIKY: SVOBODA SLOVA, PRÁVO NA INFORMACE ACENZURA VKONTEXTU PROBLÉMU DEZINFORMACE. *Proinflow: Časopis pro informační vědy*. Katedra informačních studií a knihovnictví, 2019, 2019(11), 9. ISSN 1804-2406. Dostupné z: [doi:https://doi.org/10.5817/ProIn2019-1-6CCBY](https://doi.org/10.5817/ProIn2019-1-6CCBY) 3.0 CZ

#### 4.1.1. Svoboda slova a cenzura

Každý má právo na svobodu projevu, svobodně zastávat názory, přijímat a rozšiřovat informace nebo myšlenky bez zasahování veřejné moci a bez ohledu na hranice. Zároveň svoboda a pluralita sdělovacích prostředků musí být respektována.<sup>57</sup> Jedná se o existující přirozené právo každého jedince, které je součástí osobnosti každého člověka.<sup>58</sup> Toto pojetí **svobody slova** je v oblasti informačních válek problematické pro jejich minimalizaci. Svoboda slova je sice nebezpečná, ale cenzura je ještě horší.<sup>59</sup>

**Cenzura** je ve své podstatě chápána jako přerušení libovolného toku informací. Tento proces lze popsat jako omezení či úplné odstranění k přístupu projevů (celým či jen částem), které jsou uznány jako nežádoucí.<sup>60</sup> Totéž probíhá i v rámci cenzury na internetu, který se stává stále častějším informačním zdrojem.<sup>61</sup> A to právě skrze algoritmy na sociálních sítích, které svobodu slova ovlivňují. Důvody k cenzuře předběžné nebo průběžné jsou z různých pohnutek.

#### 4.1.2. Dezinformace

Informace různého charakteru se dnes šíří neuvěřitelnou rychlostí. S tím související vzestup dezinformací má za následek také ztrátu důvěry v úřady a média obecně. S rozvojem internetových technologií se může kdokoli připojit, být autorem *fake news* či přispět k jejich šíření.<sup>62</sup> Samotná definice dezinformací je velmi problematická. Přitom pro lidstvo není žádnou novinkou. S vynálezem knihtisku, šíření gramotnosti a s příchodem masmédií, televize a rozhlasu bylo umožněno rychlejší šíření informací a *fake news* samotných.

---

<sup>57</sup> Čl. 11 Evropské Úmluvy o ochraně základních práv a svobod.

<sup>58</sup> BARTOŇ, Michal. *Svoboda projevu: principy, garance, meze*. Praha: Leges, 2010. Teoretik. ISBN 978-80-87212-42-4, s. 26-27

<sup>59</sup> POŽÁR, Rudolf. Svoboda slova je nebezpečná, cenzura je ale mnohem horší, tvrdí expertka z newyorské univerzity. <https://ct24.ceska televize.cz/> [online]. Web, 2021, 19.11.2021 [cit. 2022-04-01]. Dostupné z: <https://ct24.ceska televize.cz/svet/3402340-svoboda-slova-je-nebezpecna-cenzura-je-ale-mnohem-horsi-tvrdi-expertka-z-newyorske>

<sup>60</sup> HIMMA, Kenneth E.; TAVANI, Herman T. (ed.). *The handbook of information and computer ethics*. John Wiley & Sons, 2008. ISBN 978-0-471-79959-7, s. 573 a dále.

<sup>61</sup> WATSON, A., Které zdroje jste použili k získání informací/zpráv o vypuknutí koronaviru v posledním slabém období? 2020. <https://www.statista.com/statistics/1111783/coronavirus-news-sources-by-age-skupina-uk/> Citováno 2. dubna 2022.

<sup>62</sup> GERBINA, TV Science Dezinformace: K problému falešných zpráv. *Sci Tech. Inf. Proč.* 48, 290298(2021). <https://doi.org/10.3103/S0147688221040092>

Co je vlastně dezinformace? Dezinformace je vědomě vytvořená informace založená z části na nepravdě. K vytvoření dezinformační zprávy dochází spojením pravdivé a lživé informace. Celek je pak brán jako pravdivý. Dále lze odvrátit pozornost od pravdivé informace a přesměrovat tak pozornost na lživou část. Nebo když autor vědomě zatají některou část a zvolí si jinou. Dezinformace je informace, která není zcela pravdivá a ovlivňuje tak i jednotlivce či celou skupinu lidí. Samotná snaha o uchopení pojmu dezinformace je problematická. To si lze ukázat třeba na pornografii. Zcela určitě pornografie u mladistvých zkresluje představy o sexuálním životě. Neukazuje svému divákovi realitu, ale předkládá mu upravenou, vyumělkovanou verzi reality. Je tedy založena na pravdě, ale pravdivá není.

Jejich dopad může být opravdu vážný, a to na celou společnost. V průběhu posledních let byly oblíbeným terčem dezinformačních zpráv například politické volby nebo popírání negativního dopadu lidského chování na klima<sup>63</sup>. Nicméně připomeňme si ještě vymyšlený článek, který tvrdil, že vakcíny proti spalničkám, příušnicím a zarděnkám způsobují autismus.<sup>64</sup> Zapříčinil vzrůst těchto podobných falešných zpráv, a to převážně na sociálních sítích. Ale jeho dopad vedl k rekordnímu výskytu spalniček v roce 2018 a zvyšující se počet tzv. Anti-vaxerů.<sup>65</sup>

Jak se bránit proti dezinformacím? Asi nejdůležitější a nejefektivnější je schopnost si ověřovat fakta, výroky a zdroje. A podpora takových zdrojů, vědy a výzkumů. A zařadit je do samotného vzdělávání. Už ve vzdělávacím systému by se měl klást důraz na prevenci a schopnost kritického zhodnocení znalostí a dovedností v hodnocení informací a vědecké gramotnosti společnosti.<sup>66</sup>

Prevenčí proti takovému masovému šíření dezinformačních zpráv je schopnost hodnotit své zdroje a jejich relevantnost. A to například dle základních zákonitostí jako je spolehlivost, úplnost, objektivnost, přesnost a jejich hodnotnost. I v tomto může být AI nápomocná,

---

<sup>63</sup> WATTS, J., Máme 12 let na to, abychom omezili katastrofu způsobenou změnou klimatu, varuje OSN, *The Guardian*. <https://www.theguardian.com/environment/2018/oct/08/global-warming-must-not-exceed-15c-warns-landmark-un-report> Citováno 2. dubna 2022.

<sup>64</sup> POČET PŘÍPADŮ SPALNIČEK DOSÁHL REKORDNÍHO POČTU V EVROPSKÉM REGIONU, evropském regionu WHO. <http://www.euro.who.int/en/media-centre/sections/press-releases/2018/measles-cases-hit-record-high-in-the-european-region> Citováno 2. dubna 2022.

<sup>65</sup> GERBINA, TV Science Dezinformace: K problému falešných zpráv. *Sci. Tech. Inf. Proč.* 48, 290298(2021). <https://doi.org/10.3103/S0147688221040092>

<sup>66</sup> GERBINA TV Science Dezinformace: K problému falešných zpráv. *Sci. Tech. Inf. Proč.* 48, 290298(2021). <https://doi.org/10.3103/S0147688221040092>

například v letošním roce vznikla snaha tří českých studentů o vytvoření aplikace „Verifée“, která by měla fungovat něco jako antivir. Bude čtenáře varovat před problematickými aspekty článku, který vyhodnotí. Takže i když si člověk článek přečte bude předem varován a může tak být k němu kritičtější.<sup>67</sup>

#### 4.1.3. AI a důvěra uživatelů

Etická dilemata umělé inteligence začínají už při jejím samotném programování. Ať už samotná Evropská unie nebo do budoucnosti společnost v globálním měřítku, stojí na souboru základních hodnot. Na základě, kterých je zapotřebí utvořit etický rámec pro vývoj a užívání umělé inteligence. Pokud je AI vyvinuta a uplatňována způsobem, který respektuje sdílené etické hodnoty, můžeme ji považovat za důvěryhodnou. Většina těchto aspektů je zatím ve svých zárodcích a předmětem diskuzí. Jejichž cílem je vyvíjet a užívat AI zaměřenou na člověka a zakládající se na důvěře a neustálé interakci na úrovni člověka a stroje.<sup>68</sup>

**Výběr cílů** umělé inteligence by měl být postaven na předpokladech základních právech, lidské autonomii a být v souladu se vzrůstající a spravedlivou společností. Vývojáři si musí pořádně promyslet jaká data systému poskytnou. Také sledovat jejich dopad na jednotlivce a společnosti jako celku. To by měla hlídat řada kontrolních opatření, mechanismů řízení, jako je tzv. člověk ve smyčce nebo člověk ve velení. Mělo by být zajištěno, aby veřejné orgány měly možnost a byly kompetentní vykonávat dohled. Pokud v některých systémech člověk nemůže vykonávat dostatečný dohled, tím větší, rozsáhlejší a přísnější procesy musí probíhat před uvedením do praxe. A to platí i pokud bude člověk nahrazen v jeho pracovních činnostech, měl by tak dostat náhradní místo, příležitost.

**Digitální databáze** informací, kterou už teď řada společností vlastní, může v člověku vzbuzovat dojem, že nemá nad vlastními údaji žádnou kontrolu a mohou být vždy použity proti

---

<sup>67</sup> HRONOVÁ, Zuzana. Čeští studenti vyvíjí antivir na dezinformace. Nechceme, aby pravda prohrála, říkají. [www.aktualne.cz](https://zpravy.aktualne.cz/domaci/vyvijime-antivir-na-dezinformace-tri-cesti-kluci-vedi-jak-na/r~68363d08367a11ec98380cc47ab5f122/) [online]. 2.11.2021 [cit. 2022-04-02]. Dostupné z: <https://zpravy.aktualne.cz/domaci/vyvijime-antivir-na-dezinformace-tri-cesti-kluci-vedi-jak-na/r~68363d08367a11ec98380cc47ab5f122/>

<sup>68</sup> BUILDING TRUST IN HUMAN-CENTRIC ARTIFICIAL INTELLIGENCE: *communication from the comission to the European parliament, the council, the European economic and social comittee and the comitee of the region*[online]. 8.4. 2019, 11 [cit. 2022-03-24]. Dostupné z: <https://digital-strategy.ec.europa.eu/en/library/communication-building-trust-human-centric-artificial-intelligence>



němu. Zároveň by se mělo již při programování regulovat, kontrolovat, zda systémy AI nešíří dále zajeté předsudky a stereotypy.<sup>69</sup> Které se mohou nechat i následně zneužít, například k diskriminaci. Měla by se tedy soustředit na kvalitu datových souborů, které třeba přejímají jinde. Zároveň samotný přístup k informacím by měl být úměrně kontrolován a řízen.<sup>70</sup>

„Model trénovaný k detekci korelací by neměl být používán k vyvozování kauzálních závěrů nebo naznačování, že může. Váš model se například může dozvědět, že lidé, kteří si kupují basketbalové boty, jsou v průměru vyšší, ale to neznamená, že uživatel, který si kupuje basketbalové boty, bude v důsledku toho vyšší“<sup>71</sup>

Pozornost by měla být věnována **informovanosti všech zúčastněných stran**. Jak o možnostech a omezeních užívání AI, tak by měl být systém rozpoznatelný sám o sobě. Uživatelé tak budou vědět, že jsou v interakci se systémem AI a případně kdo za tím stojí.<sup>72</sup> I když po sobě v kyberprostoru vždy zanechají nějaké stopy, tak co se týče následné odpovědnosti, je to se systémy umělé inteligence složitější. Při negativních dopadech umělé inteligence jako její zneužívání pro podporu a tvorbu konvenčních válek, ale i těch kybernetických, by se měly zhodnotit a do budoucna minimalizovat.

Fukuyama upozorňuje na rozdílná práva mezi lidmi a jako příklad uvádí známou odlišnost pohlaví nebo dělení dle barvy kůže. Na základě toho odvozuje možné problémy v budoucnosti mezi právy u lidí bez implantátů a translidmi (lidmi s implantáty). Systémy AI by měli být ve svém fungování transparentní. Měla by disponovat etickou černou skříňkou včetně rozhodnutí a dilemat plus informací, dat o jeho systému.<sup>73</sup>

---

<sup>69</sup> V AKADEMII VĚD ČR SE DISKUTUJE O PRÁVU A ETICE UMĚLÉ INTELIGENCE. *Akademie věd České republiky* [online]. 12.9.2019 [cit. 2022-04-01]. Dostupné z: <https://www.avcr.cz/cs/veda-a-vyzkum/humanitni-a-filologicke-vedy/V-Akademii-ved-CR-se-diskutuje-o-pravu-a-etice-umele-inteligence/>

<sup>70</sup> BUILDING TRUST IN HUMAN-CENTRIC ARTIFICIAL INTELLIGENCE: *communication from the comission to the European parliament, the council, the European economic and social comittee and the commitee of the region*[online]. 8.4. 2019, 11 [cit. 2022-03-24]. Dostupné z:<https://digital-strategy.ec.europa.eu/en/library/communication-building-trust-human-centric-artificial-intelligence>

<sup>71</sup> RESPONSIBLE AI PRACTICES. *Google* [online]. [cit. 2022-03-22]. Dostupné z: <https://ai.google/responsibilities/responsible-ai-practices/>

<sup>72</sup> Tamtéž.57

<sup>73</sup> FOKUS VÁCLAVA MORAVCE. Umělá Inteligence. Část: Myslet jako člověk: Karel Oliva a Tomáš Mikolov I. [epizoda televizního pořadu]. ČT1 Ivysílání. URL: <https://www.ceskatelevize.cz/porady/11054978064-fokus-vaclava-moravce/219411030530003/cast/682788/>

## 5 AI a člověk

Vznik Nietzscheho termínu **nadčlověk** byl důsledkem hodnotové krize moderní civilizace, měl být jako dalším vývojovým stádiem člověka. Proto byla myšlenka budoucího možného stvoření nadčlověka, který bude disponovat novými a zatím neznámými tvůrčími schopnostmi a novým svobodným prostorem k uplatnění panství nad sebou samým a nad světem.<sup>74</sup> Tedy, nadčlověka, který bude schopný překonat, to typicky lidské. „Člověk se stává mostem k nadčlověku.“<sup>75</sup>

Egon Bondy už přichází s další vývojovou fází člověka jako tzv. **umělých bytostí**. Artificiální bytosti, které se nenechají strhnout svými emocemi, pudy, vášní, budou spíše čistým intelektem, rozumem. Budou sami hledat význam, a hodnoty bez zadání nás samotných, tzv. „najdi hodnoty sama“.

Když upustíme od opravdového významu Nietzscheho nadčlověka, shodneme se na tom, že představa vylepšeného člověka není ničím, co by přišlo s 20.stoletím. Klasičtí hrdinové s nadpřirozenými schopnostmi jsou stále vracející se klasikou, ale sci-fi průmysl ovládla vlna vylepšených robotů. Robotů, kteří samostatně myslí, rozhodují a na konci, nás tato vylepšená verze nás samých přivede ke zkáze. Byl by člověk schopen stvořit sám sebe? Je schopen stvořit svojí lepší verzi?

„Naše mysl je složitý výtvar, spředený z mnoha různých pramenů a zahrnující mnoho různých plánů. Některé prvky jsou staré jako život sám a některé jsou nové jako současná technologie.“<sup>76</sup> Dostane se AI do fáze, kdy bude schopna samostatně myslet? Jak vlastně funguje lidská mysl? Jsou teorie, které si vykládají mysl jako fungující stroj, funguje mozek mechanicky? Pokud ano, tak by člověk mohl být schopný vytvořit sám sebe, ale opravdu je představa myslí, jako hardwaru a softwaru pravdivá?

---

<sup>74</sup> NICOLA, Ubaldo. *Obrazové dějiny filozofie*. V Praze: Euromedia Group - Knižní klub, 2006. Universum (Knižní klub). ISBN 80-242-1578-0, s. 468.

<sup>75</sup> FINK, Eugen. *Filosofie Friedricha Nietzscheho*. Praha: OIKOYMENH, 2011. Oikúmené (OIKOYMENH). ISBN 978-80-7298-266-0, s. 155.

<sup>76</sup> DENNETT, Daniel Clement. *Druhy myslí: k pochopení vědomí*. Praha: Academia, 2004. Mířtři vědy. ISBN 80-200-1177-3, s. 10.

*„Mysl má víc než jen syntax, má i sémantiku. A žádný počítačový program nikdy nebude mít mysl prostě proto, že je pouze syntaktický, zatímco mysl není jen syntaktická. Mysl je sémantická – sémantická v tom smyslu, že má víc než formální strukturu: má také obsah.“<sup>77</sup>*

## 5.1. Lidská mysl a počítač

Lze popsat tak, že mentální stavy mají své vlastní intencionální vnitřní obsahy a je třeba dbát na **naši vnitřní zkušenost mentálních stavů**. Neurální procesy lze přímo pozorovat, avšak myšlenky, pocity a představy nějaké osoby mohou ostatní poznat pouze prostřednictvím jejího chování a z toho co říká. Probíhá také nepřímo, a to z tělesných projevů.<sup>78</sup>

Kdyby klasický Google překladač, měl teoreticky AI, která by byla schopna sama vyhodnocovat, překládat a následně rozpoznat tu abstraktnější hodnotu textu, čistě na základě sesbíraných dat. Jaký by byl rozdíl od textu přeložený překladatelem? Překladatel sám není schopen vše přeložit přesně tak, jak bylo myšleno, také může některé významy zásadně změnit. Ale jeho šanci zvyšuje znalost daného tématu, to si můžeme ukázat i na klasickém románu. Překladatel jej překládá s vlastní zkušeností, snaží se vcítit do daných postav, reflektuje chování postav v minulosti či v budoucnosti. Jak hlubokou znalost by musela mít taková AI, jaká data by jí musela být poskytnuta, lze něco takového vůbec?<sup>79</sup> A pokud bychom používali překladač, tak, jak ho známe, neztratil by se tento otisk překladatele celkově i v budoucnosti? Zároveň není síla počítačů právě v tom, že jsou oproštěni od těchto mentálních rysů? To, že nemají intuici, nemusí data organizovat, sjednocovat, syntetizovat nebo posuzovat. Jejich cílem je pouze generovat výsledky na základě programování od širokých týmů programátorů a není tak na nich posuzovat jejich správnost.<sup>80</sup>

Je duchovní svět schopen udržet krok s tím technologickým? Coreth ve své knize Co je člověk píše, že pokrok je pohyblivější tím, čím více se blíží k vlastním lidským hodnotám. Poukazuje

---

<sup>77</sup> SEARLE, John. *Mysl, mozek a věda*. 1. vyd. Praha: Mladá fronta, 1994. 129 s. Váhy; sv. 13. ISBN 80-2040-509-7.

<sup>78</sup> PEROUTKA, David. *Tomistická filosofická antropologie*. Praha: Krystal OP, 2012. ISBN 978-80-87183-42-7, s. 13

<sup>79</sup> FOKUS VÁCLAVA MORAVCE. Část: Umělá Inteligence. [epizoda televizního pořadu]. ČT1 Ivysílání. URL: <https://www.ceskatelevize.cz/porady/11054978064-fokus-vaclava-moravce/219411030530003/>

<sup>80</sup> BENTLEY HART, David. Reality Minus On the depressing fantasy of minds in simulated worlds. *TheNewAtlantis.com* [online]. 2022, 4.2. 2022, 2022 [cit. 2022-03-13]. Dostupné z: <https://www.thenewatlantis.com/publications/reality-minus>

na vzestup, ale i na pád, a to především pokud mluvíme o osobních hodnotách jednotlivce a jeho vlastního svobodného osobního rozhodování. Člověk už nezdědí znalosti a zkušenosti, postoje, cnosti svého otce a dále je tak moc rozvíjet věcné jmění, možná jedině skrze prostředí tvořené předky, to je určitě usnadnění (kulturní statky, mravy výchovný vliv jednoho člověka na druhého). Ale Coreth říká, že v tomto každý do určité míry začíná od začátku a nemůže prostě přijmout to, čeho se již dosáhlo a dále na tom stavět protože záleží na svobodném rozhodnutí jednotlivce.<sup>81</sup>

## 5.2. Umělý smysluplný život, umělé smysluplné já

V současném světě je těžké rozeznat, co je realitou a co nikoli. A v jistém smyslu nás do budoucna nečeká vývoj jiným směrem. Virtuální realita je úzce spojena s tou nevirtuální. Mozek, který je postupně vylepšován různými technologiemi by určitě vytvořil jakéhosi „nadčlověka“. Byl by ale, takový člověk stále člověkem? Filozof David J. Chalmers ve své knize *Reality plus* předkládá teorii, že pokud bychom mozek nahrazovali po částech obvody, bude nadále schopen stejného chování jako jeho původní neurologie. Je přesvědčen, že by se virtuální realita neměla odsouvat na druhou kolej. Domnívá se, že už možná i v jedné žijeme, a to dokonce i smysluplným životem.

Ve svém článku *Realita mínus*, filozof a teolog David Bentley Hart oponuje Chalmersovi tím, že by v simulovaném světě mohla existovat pouze simulovaná já. Hart s Chalmersovou představou myslí jakožto počítače nesouhlasí. „Jak by mohla myšlenka vzniknout z čistě mechanického substrátu, přičemž výpočty jsou modelem toho, jak systém fyzických „přepínačů“ nebo operací může generovat syntaxi funkcí, které zase generuje sémantiku myšlenky, ta zase vytváří realitu nebo iluzi vědomí. Nic z toho počítač nedělá a rozhodně nic z toho neudělá mozek, kdyby byl počítač.“<sup>82</sup>

V dokumentárním snímku „Já, člověk“ připomínají fakt, že toho o lidském mozku víme velmi málo.<sup>83</sup> A lze tedy říct, že představa neurofyzologie jakožto obecnější výraz neurotechnologie,

---

<sup>81</sup> CORETH, Emerich. *Co je člověk?*. 2. vyd. Praha: Zvon, 1996. Studium (Zvon). ISBN 80-7113-170-9, s. 179

<sup>82</sup> BENTLEY HART, David. *Reality Minus On the depressing fantasy of minds in simulated worlds*. *TheNewAtlantis.com* [online]. 2022, 4.2. 2022, 2022 [cit. 2022-03-13]. Dostupné z: <https://www.thenewatlantis.com/publications/reality-minus>

<sup>83</sup> *JÁ, ČLOVĚK* [I AM HUMAN] [film]. Režie: SOUTHERN Taryn, GABY Elena. USA, 2019

kteřá by byla vyjádřena v nějakém jiném fyzickém médiu, mysl jako vzor činnosti, kterou lze přesunout do kybernetické sítě, virtuálního rámce a myšlení, zkušenost nebo všechny ostatní rysy mentality by byly pouze záležitostí spojení a vodivosti, bez jakékoli souvislosti s fyzickou složkou, procesy a individuální historií mozku a těla, spíše opravdu připomíná iluzi o tom, co mysl, mozek a počítač opravdu jsou. „Postavte si model mozku a ten vykouzlí sjednocené vědomí, jakmile zapnete elektřinu“<sup>84</sup>

Organické aspekty se prolínají s anorganickými a jsou pro koncepci umělých světů zásadní. Prolínání technologií a organismu, tělesného světa či strojovosti vede k novému chápání reality, tělesného světa, ale i toho umělého, které se v prožívání významně liší. Vznik virtuální duše jako digitalizované „camera obscura“, která v sobě značí umělejší svět i s jeho softwarovými procesy. Prolínání technologií a organismu je začátkem pro vývoj konceptů kyborga a umělé inteligence.<sup>85</sup>

---

<sup>84</sup> BENTLEY HART, David. Reality Minus On the depressing fantasy of minds in simulated worlds. *TheNewAtlantis.com* [online]. 2022, 4.2. 2022, 2022 [cit. 2022-03-13]. Dostupné z: <https://www.thenewatlantis.com/publications/reality-minus>

<sup>85</sup> CHARVÁT, Martin. *O nových médiích, modularitě a simulaci*. Praha: Togga, 2017. ISBN 978-80-7476-121-8, s. 76.

### 5.3. Transhumanismus

V transhumanismu je vize člověka, který bude své části těla postupně nahrazovat nejrůznějšími variacemi součástek. Biblická koncepce člověka je založena na stvoření „Imago dei“, tedy člověk stvořen jako obraz Boží. Transhumanismus znamená něco jako přesahující, mimo jdoucí, ještě za to typicky lidské. Dokument *Já, člověk* představuje, nově vyvinuté technologie, které umí člověka „vyléčit“, dát mu naději na normální život a to i třeba jen na pár minut, stanou se však i tyto stejné technologie do budoucna naší hrozbou?

V dokumentu byl dokonce experiment, kde jeden respondent dostane otázku a druhý respondent mu poskytne odpověď, ale pouze telepaticky, přičemž se oba respondenti celou dobu nachází úplně v jiných zemích. Z toho vyvěrá mnoho otázek. Je člověk s takovou technikou v těle stále člověkem? Jaký by byl svět, ve kterém by si člověk nemohl být jistý ani vlastní hlavou, vlastními myšlenkami.<sup>86</sup>

Každý, kdo vlastní mobilní telefon či počítač, by měl mít na paměti, že je důležité pro jeho bezpečnost, updatovat tato zařízení. U některých zařízení to však nelze nebo je to příliš nákladné. Představme si někoho, kdo má například kardiostimulátor nebo defibrilátor s technologií bluetooth. Takovou technologii lze ovládat bezdrátově. Scénář, kdy by nad ní někdo převzal kontrolu, by pro zkušeného hackera, nepředstavoval zas tak těžký úkol.<sup>87</sup>

Většina tzv. kyborgů vykazuje atributy, které mohou připomínat Nietzscheho „nadčlověka“ – transčlověka. Technologické součástky, které si implantují na lidské já, které ve svém základu zůstává do značné míry nedotčeno. Technologická změna je však považována za umělou a kyborg se staví do opozice proti „lidské přirozenosti“. Řada optimistů vidí v transhumanismu potenciál pro rozvoj až překonání mentálních a fyzických vlastností člověka až po vizi nesmrtnosti. Francis Fukuyama považuje transhumanismus za hrozbu současné civilizace.

---

<sup>86</sup> *JÁ, ČLOVĚK* [I AM HUMAN] [film]. Režie: SOUTHERN Taryn, GABY Elena. USA, 2019

<sup>87</sup> SPITZER, Manfred. *Kybernemoc!: jak nám digitalizovaný život ničí zdraví*. Přeložil Iva KRATOCHVÍLOVÁ. Brno: Host - vydavatelství, 2016. ISBN 978-80-7491-792-9, s. 122

Upozorňuje na mnoho etických otázek, ale hlavně varuje před transhumanistickým ideálem nesmrtelnosti. Říká, že právě smrtelnost nám dovoluje jako celku přežít a přizpůsobovat se.<sup>88</sup>

David J. Chalmers je přesvědčen o tom, že **mozek nahrazený po částech obvodu** bude vykonávat stejnou funkci jako původní neuronový systém mozku. David Bentley Hart mu ale ve svém článku oponuje „zdá se mnohem pravděpodobnější, že proces nahrazování některých neuronů počítačovými čipy by byl jen o málo víc než velmi pomalý proces sebevraždy, který by nezpůsobil stejné chování jako živá mysl, ale pouze postupné vyšínutí a omámení“<sup>89</sup>.

### 5.3.1. Člověk, stroj a tělesnost

Martin Charvát popisuje zcela jiné a nové **chápání těla ve virtuální realitě**. Toto „kybertělo“ ve virtuální realitě je charakterizováno jako tělo, které přijímá jiné podněty, ty jsou technologicky determinované a mechanicky se uvádí v pohyb. Objekty ve virtuálním prostředí jsou sice ve svém základu stejné, ale digitálně diferenciované. Lze tedy ve virtuálním prostředí poznávat animální přírodu, nicméně **artificiálně reprezentovanou** v jejímž základu je digitální kód.

To, značně komplikuje dělení živého na neživé. Nebo přítomností a tzv. telepřítomností ve virtuálním prostředí. Tělo a naše individualita se stává **mnohočetnou**. Oproti žité reálné zkušenosti se objekty ve virtuální realitě nemusí definovat zákony, které aplikujeme v běžném životě. To vede k jinému chápání vztahu mezi námi a prostředím.<sup>90</sup> Přičemž důležitou otázkou je, jak zajistit, aby při užívání systému AI byla dodržována základní lidská práva a zákony, pravidla a normy. Člověk je dodržuje ze základních lidských hodnot a strachu ze sankcí. Ale jak zajistit, aby je dodržovala AI a technologie, pokud tento strach necítí?

---

<sup>88</sup> CARNEGIE ENDOWMENT FOR INTERNATIONAL PEACE. The World's Most Dangerous Ideas. *MyWire* [online]. [cit. 15-02-2022]. Dostupný z: <http://www.mywire.com/pubs/ForeignPolicy/2004/09/01/564801?page=4>

<sup>89</sup> BENTLEY HART, David. Reality Minus On the depressing fantasy of minds in simulated worlds. *TheNewAtlantis.com* [online]. 2022, 4.2. 2022, 2022 [cit. 2022-03-13]. Dostupné z: <https://www.thenewatlantis.com/publications/reality-minus>

<sup>90</sup> CHARVÁT, Martin. *O nových médiích, modularitě a simulaci*. Praha: Togga, 2017. ISBN 978-80-7476-121-8, s. 72-73

## 6. Závěr

Systémy umělé inteligence jsou na vzestupu. Již dnes je můžeme najít na každém kroku. Prostupují postupně každou částí našich životů, nás samotných a našeho světa ve kterém žijeme. Utváří tak nová pojetí reality a jejího chápání. Je tématem velmi aktuálním, které se dotýká každého z nás. Přesto by se mohlo zdát, že se jí nevěnuje dostatečná pozornost a je předmětem pouhých diskuzí a plná mystifikací.

Cílem této práce bylo ve zkratce představit nejdůležitější pojmy, které jsou s AI spojeny. Nastínila její vznik, vývoj a využití. Představili jsme si její obrovský potenciál v podobě pomocníka a průvodce současným světem, jehož je zároveň spoluvůrce. Také si dala práce za cíl představit několik příkladů, jak lze tyto systémy i zneužít na úkor jejího uživatele. Představili jsme si příklady, jak AI proměňuje vývoj s pojetím nás samých a světa a dochází tak k ohrožení naší lidské přirozenosti a svobody člověka.

Je třeba na tento vývoj adekvátně reagovat. Edukovat a informovat uživatele o bezpečnostních rizicích užívání umělé inteligence. Zapojit toto téma i do diskusí s širokou veřejností. Neustále kriticky vyhodnocovat a reagovat na chování umělé inteligence a neudělovat jí právo na samostatná rozhodnutí. Člověk by měl být vždycky součástí řízení systému a být za něj odpovědný. Takovou pozornost si AI a s ní spojená dilemata zaslouží už v samotném vzdělávání, aby tak budoucí generace mohly pružněji reagovat na její vývoj a reflektovat její vliv na naši společnost.

Primárním cílem vývoje AI by neměl být stejně nebo lépe smýšlející systém, než je člověk sám. Jsem přesvědčena, že by pak ztratila svůj smysl. Tyto systémy jsou tak úspěšné právě pro jejich neschopnost lidských slabostí. AI by měla být primárně pomocníkem, průvodcem, se kterým půjdeme bok po boku světem. Navzájem se doplňovat a reagovat. A pokud se vrátíme k Alanu Turingovi, tak by T-test mohl být dostačující podmínkou pro detekci inteligence námi vytvořených strojů a systémů.



## Použité zdroje

### Bibliografie

ARO, Jessikka. *Putinovi trollové: skutečné příběhy z frontových linií ruské informační války*. Přeložil Lenka FÁROVÁ. Praha: N Media, 2020. Edice N. ISBN 978-80-907922-1-0.

BARTOŇ, Michal. *Svoboda projevu: principy, garance, meze*. Praha: Leges, 2010. Teoretik. ISBN 978-80-87212-42-4.

BRDIČKA, Bořivoj. Co dokážou stroje schopné hlubokého učení. <em>Metodický portál: Spomocník </em>[online]. 25. 04. 2016, [cit. 2022-02-01]. Dostupný z WWW: <https://spomocnik.rvp.cz/clanek/20855/CO-DOKAZI-STROJE-SCHOPNE-HLUBOKEHO-UCENI.html> ISSN 1802-4785

BUILDING TRUST IN HUMAN-CENTRIC ARTIFICIAL INTELLIGENCE: Communication from the commission to the European parliament, the council, the European economic and social committee and the committee of the regions [online]. 8.4. 2019, 11 [cit. 2022-03-24]. Dostupné z: <https://digital-strategy.ec.europa.eu/en/library/communication-building-trust-human-centric-artificial-intelligence>

BURDA, František. *Za hranice kultur: transkulturní perspektiva*. Brno: Centrum pro studium demokracie a kultury, 2016. ISBN 978-80-7325-402-5.

CARNEGIE ENDOWMENT FOR INTERNATIONAL PEACE. The World's Most Dangerous Ideas. *MyWire* [online].[cit.15-02-2022]. Dostupný z: <http://www.mywire.com/pubs/ForeignPolicy/2004/09/01/564801?page=4>

CNBC. Interview With The Lifelike Hot Robot Named Sophia (Full) | CNBC. YouTube [online]. 2017 [cit. 2022-3-17]. Dostupné z: <https://www.youtube.com/watch?v=S5t6K9iwcdw&t=289s>

CORETH, Emerich. *Co je člověk?: Základy filozofické antropologie*. Přeložil Bohuslav VIK. Praha: Zvon, 1994. ISBN 80-7113-098-2.

## ČL. 11 EVROPSKÉ ÚMLUVY O OCHRANĚ ZÁKLADNÍCH PRÁV A SVOBOD.

DENNETT, Daniel Clement. *Druhy myslí: k pochopení vědomí*. Praha: Academia, 2004. Mistři vědy. ISBN 80-200-1177-3.

FINK, Eugen. *Filosofie Friedricha Nietzscheho*. Praha: OIKOYMENH, 2011. Oikúmené (OIKOYMENH). ISBN 978-80-7298-266-0.

GERBINA, TV Science Dezinformace: K problému falešných zpráv. Sci. Tech. Inf. Proč. 48, 290298(2021). <https://doi.org/10.3103/S0147688221040092>

HAVEL, Ivan M.. Přirozené a umělé myšlení jako filozofický problém [online]. Glosy.info, 3. prosince 2004. [cit. 30.1.2022] Dostupné z: ISSN 1214-8857.

HIMMA, Kenneth E.; TAVANI, Herman T. (ed.). The handbook of information and computer ethics. John Wiley & Sons, 2008. ISBN 978-0-471-79959-7.

CHARVÁT, Martin. *O nových médiích, modularitě a simulaci*. Praha: Togga, 2017. ISBN 978-80-7476-121-8.

CHUA, IS, Ritchie, CS & Bates, DW Zlepšení komunikace s vážnými nemocemi pomocí umělé inteligence. *npj číslice. Med.* 5, 14 (2022). <https://doi.org/10.1038/s41746-022-00556-2>,

KOLOUCH, Jan. *CyberCrime*. Praha: CZ.NIC, z.s.p.o., 2016. CZ.NIC. ISBN 978-80-88168-15-7.

KOROLOV, Maria, Miroslav NEČAS a Saniye ALAYBEYI. Umělá inteligence pro podnikové nasazení. Computerworld [online]. 2021, červenec 2021, 2021(07-08), 8 [cit. 2022-03-19]. Dostupné z: <https://i.iinfo.cz/files/computerworld/6/umela-intelligence-1.pdf>

LECUN, Yann; BENGIO, Yoshua; HINTON, Geoffrey. Deep learning 2015. [28.1.2022]. Dostupný z <http://www.cs.toronto.edu/~hinton/absps/NatureDeepReview.pdf>

LEE, Kai-fu. *Supervelmoci umělé inteligence: Čína, Silicon Valley a svět v éře AI*. Přeložil Petr HOLČÁK. Praha: Argo. Crossover. ISBN 978-80-257-3050-8.

MAŘÍK, V. Umělá inteligence 1. 1. vyd. Praha: Academia, 1993. 264 s. ISBN 80-200-0496-3. Kapitola 1, Některé charakteristiky umělé inteligence.

NICOLA, Ubaldo. *Obrazové dějiny filozofie*. V Praze: Euromedia Group - Knižní klub, 2006. Universum (Knižní klub). ISBN 80-242-1578-0.

PEROUTKA, David. *Tomistická filosofická antropologie*. Praha: Krystal OP, 2012. ISBN 978-80-87183-42-7.

SEARLE, John. *Mysl, mozek a věda*. 1. vyd. Praha: Mladá fronta, 1994. 129 s. Váhy; sv. 13. ISBN 80-2040-509-7.

SNYDER, Timothy. *Tyranie: 20 lekcí z 20. století*. Přeložil Martin POKORNÝ. V Praze: Paseka, 2017. ISBN 978-80-7432-838-1.

SPITZER, Manfred, Václav BĚLOHRADSKÝ, Martin LESKOVJAN, Jeremiáš HAVRANU, Jana KARLOVÁ, Martin SOUKUP a Jan D. BLÁHA, SOKOLÍČKOVÁ, Zdenka, ed. *Kyber a eko: digitální technologie v enviromentálních souvislostech*. Brno: Host, 2019.

SPITZER, Manfred. *Kybernemoc!: jak nám digitalizovaný život ničí zdraví*. Přeložil Iva KRATOCHVÍLOVÁ. Brno: Host - vydavatelství, 2016. ISBN 978-80-7491-792-9.

STODOLA, Jiří. Principy a dilemata informační etiky: svoboda slova, právo na informace a cenzura v kontextu problému dezinformace. *Proinflow: Časopis pro informační vědy*. Katedra informačních studií a knihovnictví, 2019, **2019**(11), 9. ISSN 1804-2406. Dostupné z: doi: <https://doi.org/10.5817/ProIn2019-1-6CCBY 3.0 Cz>

TEGMARK, Max. *Život 3.0: člověk v éře umělé inteligence*. Přeložil Markéta IVÁNKOVÁ. Praha: Argo, 2020. Zip (Argo: Dokořán). ISBN 978-80-257-3173-4.

TVRDÝ, Filip. *Turingův test: filozofické aspekty umělé inteligence*. Praha: Togga, 2014. Scholia (Togga). ISBN 978-80-7476-043-3, s. 15.

## Internetové zdroje

BENTLEY HART, David. Reality Minus On the depressing fantasy of minds in simulated worlds. *TheNewAtlantis.com* [online]. 2022, 4.2. 2022, **2022** [cit. 2022-03-13]. Dostupné z: <https://www.thenewatlantis.com/publications/reality-minus>

Čínské kamery se zdokonalily na zadržených Ujgurech. Čechy snímají na ulici i úřadech. <https://www.aktualne.cz> [online]. online, 2021, 24.8. 2021 [cit. 2022-03-18]. Dostupné z: <https://zpravy.aktualne.cz/domaci/armada-skoly-i-urad-vlady-cesko-monitoruji-cinske-kamery-pod/r~59ac5ce0da6d11eb94d2ac1f6b220ee8/>

Elements of AI: Související oblasti. <https://www.elementsofai.cz> [online]. [cit. 2022-03-14]. Dostupné z: <https://course.elementsofai.com/cs/1/2>

FOKUS VÁCLAVA MORAVCE. Část: Umělá Inteligence. [epizoda televizního pořadu]. ČT1 Ivysílání. URL: <https://www.ceskatelevize.cz/porady/11054978064-fokus-vaclava-moravce/219411030530003/>

FOKUS VÁCLAVA MORAVCE. Umělá Inteligence. Část: Myslet jako člověk: Karel Oliva a Tomáš Mikolov I. [epizoda televizního pořadu]. ČT1 Ivysílání. URL: <https://www.ceskatelevize.cz/porady/11054978064-fokus-vaclava-moravce/219411030530003/cast/682788/>

FORUM Magazín Univerzity Karlovy: Dataholka: Jak sociální sítě zneužívají toho, co nás dělá lidmi. [www.ukforum.cz](http://www.ukforum.cz) [online]. Web Univerzity Kralovy, 2021, 23 listopad 2021 [cit. 2022-03-23]. Dostupné z: <https://www.ukforum.cz/rubriky/alumni/8132-jak-socialni-site-zneuzivaji-toho-co-nas-dela-lidmi>

HRONOVÁ, Zuzana. Čeští studenti vyvíjí antivir na dezinformace. Nechceme, aby pravda prohrála, říkají. [www.aktualne.cz](http://www.aktualne.cz) [online]. 2.11.2021 [cit. 2022-04-02]. Dostupné z:

<https://zpravy.aktualne.cz/domaci/vyvijime-antivir-na-dezinformace-tri-cesti-kluci-vedi-jak-na/r~68363d08367a11ec98380cc47ab5f122/>

FACEBOOK BUYS ISRAELI FACIAL RECOGNITION FIRM FACE.COM: <https://www.bbc.com/news/technology-18506255>. BBC [online]. BBC: Oxford University press, 2012, 19 červen 2012 [cit. 2022-03-23]. Dostupné z: <https://www.bbc.com/news/technology-18506255>

OXFORD ENGLISH DICTIONARY: <https://www.oed.com/viewdictionaryentry/Entry/271625>: Artificial intelligence. *Oxford English Dictionary* [online]. Oxford University press [cit. 2022-03-02]. Dostupné z: <https://www.oed.com/viewdictionaryentry/Entry/271625>

*JÁ, ČLOVĚK* [I AM HUMAN] [film]. Režie: SOUTHERN Taryn, GABY Elena. USA, 2019

MATAS, Jiří. Lídři to s předpovědí nástupu autonomních aut přehnali, říká odborník na umělou inteligenci [online]. 12.12.2021 [28.1. 2022]. Dostupné z: <https://www.e15.cz/rozhovory/lidri-to-s-predpovedi-nastupu-autonomnich-aut-prehnali-rika-odbornik-na-umelou-inteligenci-1386151>

POČET PŘÍPADŮ SPALNIČEK DOSÁHL REKORDNÍHO POČTU V EVROPSKÉM REGIONU:WHO. <http://www.euro.who.int/en/media-centre/sections/pressreleases/2018/measles-cases-hit-record-high-in-the-european-region>. Citováno 2. dubna 2022.

POŽÁR, Rudolf. Svoboda slova je nebezpečná, cenzura je ale mnohem horší, tvrdí expertka z newyorské univerzity. <https://ct24.ceskatelevize.cz/> [online]. Web, 2021, 19.11.2021 [cit. 2022-04-01]. Dostupné z: <https://ct24.ceskatelevize.cz/svet/3402340-svoboda-slova-je-nebezpecna-cenzura-je-ale-mnohem-horsi-tvrdi-expertka-z-newyorske>

RESPONSIBLE AI PRACTICES. Google [online]. [cit. 2022-03-22]. Dostupné z: <https://ai.google/responsibilities/responsible-ai-practices/>

THOUSAND OF LEADING AI RESEARCHERS SIGN PLEDGE AGAINST KILLER ROBOTS. *Support the Guardian* [online]. 2018, 18.7. 2018 [cit. 2022-04-01]. Dostupné z:

<https://www.theguardian.com/science/2018/jul/18/thousands-of-scientists-pledge-not-to-help-build-killer-ai-robots>

TOMÁŠ BOROVIČKA: Umělá inteligence přispívá k větší udržitelnosti zemědělství. *Https://prg.ai/* [online]. 2021, 18.11.2021 [cit. 2022-03-25]. Dostupné z: <https://prg.ai/tomas-borovicka-datamole-ai-umela-inteligence-zemedelstvi/> /

V AKADEMII VĚD ČR SE DISKUTUJE O PRÁVU A ETICE UMĚLÉ INTELIGENCE. *Akademie věd České republiky* [online]. 12.9.2019 [cit. 2022-04-01]. Dostupné z: <https://www.avcr.cz/cs/veda-a-vyzkum/humanitni-a-filologicke-vedy/V-Akademii-ved-CR-se-diskutuje-o-pravu-a-etice-umele-inteligence//>

VLACH, Jan a PŘÍNOSIL Jiří. *Lokalizace obličeje v obraze s komplexním pozadím*. [online]. 11.4.2007 [cit. 28.1. 2022]. Dostupné z: <http://www.elektrorevue.cz/cz/clanky/zpracovani-signalu/0/lokalizace-obliceje-v-obraze-s-komplexnim/>

WATSON A., Které zdroje jste použili k získání informací/zpráv o vypuknutí koronaviru v posledním slabém období?, 2020. <https://www.statista.com/statistics/%201111783/coronavirus-news-sources-by-age%20-skupina-uk/>. Citováno 2. dubna 2022.

WATTS, J. Máme 12 let na to, abychom omezili katastrofu způsobenou změnou klimatu, varuje OSN, TheGuardian . <https://www.theguardian.com/environment/2018/oct/08/global-warming-must-not-exceed-15c-warns-landmark-un-report>. Citováno 2. dubna 2022.