

VYSOKÁ ŠKOLA EKONOMICKÁ V PRAZE

Fakulta informatiky a statistiky

Katedra informačních technologií

Studijní program: Aplikovaná informatika

Obor: Kognitivní informatika



DIPLOMOVÁ PRÁCE

Význam poznávacích procesů pro tvorbu umělé inteligence

Meaning of cognitive processes for creating artificial intelligence

Diplomant: Bc. Zdeněk Smutný

Vedoucí diplomové práce: doc. PhDr. Karel Pstružina, CSc.

Oponent diplomové práce: prof. RNDr. Jozef Kelemen, DrSc.

Školní rok 2009/2010

Motto

„Víme, že ve vědách i ve všech ostatních oblastech nepřinášejí pokrok ti, kdo křečovitě lpějí na ustáleném stavu věcí, nýbrž ti, kdo usilují o lepší, kdo se odvažují stále měnit vše, co není v pořádku.“

Isokratés

Prohlášení

Prohlašuji, že jsem diplomovou práci zpracoval samostatně a že jsem uvedl všechny použité prameny a literaturu, ze kterých jsem čerpal.

V Praze, dne 4. května 2010

Zdeněk Smutný

Poděkování

Rád bych tímto poděkoval především doc. PhDr. Karlu Pstružinovi, CSc. za osobní vstřícný přístup, nedocenitelné rady a odborné vedení mé práce.

Dále chci poděkovat MUDr. Robertu Rusinovi¹ a Ing. Janu Burianovi² za zajímavé poznámky k částem této práce a poskytnutou literaturu. Na závěr bych rád poděkoval také Ing. Tomášovi Stárkovi³, Ph.D. a Ing. Tomášovi Tvrzskému⁴ za poskytnutí výchozích materiálů k třetí kapitole této práce.

¹Neurologická klinika, Fakultní Thomayerova nemocnice v Praze, 3. lékařská fakulta UK v Praze.

²Katedra informačního a znalostního inženýrství, Fakulta informatiky a statistiky VŠE v Praze.

³Ústav řídicí techniky a telematiky, Fakulta dopravní ČVUT v Praze.

⁴Předseda představenstva společnosti Telematix Services a.s. zabývající se výzkumem a vývojem v oblasti inteligentních dopravních systémů.

Abstrakt

Práce předkládá ucelený pohled na poznávací procesy v rámci systémů umělé inteligence a nabízí jejich srovnání s podobnými procesy viděnými v přírodě včetně člověka. Historické pozadí směřující k vývoji současných koncepcí umělé inteligence umožňuje náhled na celou problematiku z určité perspektivy. Historický exkurz přechází postupně na hlavní osu zájmu: *prostředí – podněty – zpracování – odraz v kognitivním systému – reakce na podněty*, kdy srovnávám přístup a limitní možnosti člověka a stroje, respektive umělé inteligence. V závěrečné části představím dva realizované projekty inteligentních dopravních systémů a poukážu na potenciál, který se nabízí pro jejich další rozvoj v rámci hlavní osy zájmu. U každé části této práce je důraz kladen na spojitost s kognitivními procesy, jejich provázanost a závislost.

Klíčová slova: kognitivní procesy, umělá inteligence, umělý život, kognitivní věda, dějiny vědy, inteligentní dopravní systémy

Abstract

This diploma thesis brings an integral view at cognitive processes connected with artificial intelligence systems, and makes a comparison with the processes observed in nature, including human being. A historical background helps us to look at the whole issue from a certain point of view. The main axis of interest comes after the historical overview and includes the following: *environment – stimulations – processing – reflection in the cognitive system – reaction to stimulation*; I balance the approach and the limited potential of the human being against the machine (or artificial intelligence). In the last part, there are introduced two projects that have been already implemented in the intelligent transport systems, and their potential for the further expansion and development is shown here. The main emphasis is placed on the coherence between each part of this thesis and cognitive processes, and on the relation and the mutual dependence of these processes.

Keywords: cognitive processes, artificial intelligence, artificial life, cognitive science, history of science, intelligent transport systems

Obsah

Úvod	8
1. Směřování k umělé inteligenci	10
1.1. Pohled do minulosti strojů a abstraktního myšlení	14
1.2. Pojem umělá inteligence	21
1.3. Limity poznávání dané historickou zkušeností	24
1.4. Možné budoucí směřování umělé inteligence, respektive umělého života	27
1.5. Závěrečná implikace	31
2. Kognitivní procesy člověka a umělé inteligence	35
2.1. Prostředí a podněty	36
2.1.1. Prostředí člověka a umělé inteligence	40
2.1.2. Podněty a jejich zpracování u člověka a umělé inteligence	46
2.1.3. Shrnutí problematiky prostředí a podnětů	50
2.2. Myšlení a odraz v kognitivním systému	51
2.2.1. Centrum kognitivních funkcí	51
2.2.2. Učení a paměť	55
2.2.3. Problematika zapomínání	60
2.2.4. Myšlení	61
2.2.5. Člověk jako stroj	65
2.3. Zpětná reakce	70
2.4. Transdisciplinární přístup	73
3. Pilotní projekty inteligentních dopravních systémů v ČR	78
3.1. Monitorování a řízení pohybu pohyblivých objektů po pohybové ploše letišť pomocí GNSS	79
3.1.1. Základní informace	79
3.1.2. Zhodnocení projektu z hlediska kognitivních procesů	80

3.2. Informační systém pro přepravu nebezpečných věcí využívající systém GNSS	82
3.2.1. Základní informace	82
3.2.2. Zhodnocení projektu z hlediska kognitivních procesů	82
3.3. Směrování k inteligentnímu vozidlu	84
4. Závěr	86
Terminologický slovník	90
Seznam obrázků	92
Seznam příloh	93
Literatura	95

Úvod

Poznávací procesy jsou nepostradatelné pro život, nalezneme je už u nejprimitivnějších organismů a v průběhu evoluce se stále zdokonalují. Vzhledem k tomu, že se v oblasti umělé inteligence necháváme v mnohém inspirovat především živou přírodou, tak se jedná o fundamentální problematiku, která zasahuje i do umělých systémů vytvořených člověkem. Hledáme v přírodě různé analogie, podnětné myšlenky a přístupy, které můžeme aplikovat v umělé inteligenci. To je také cílem kognitivní informatiky dle [ICCI 2010], jež využívá transdisciplinárního přesahu více vědních oborů.

Jako výchozí bod své práce jsem si stanovil tzv. hlavní osu zájmu kognitivních procesů: prostředí – podněty – zpracování – odraz v kognitivním systému – reakce na podněty. Samozřejmě s přihlédnutím k faktu, že některé přístupy k umělé inteligenci se zabývají znaky, jež se objevily na závěr evoluce života, a jiné naopak na jeho počátku.⁵

Téma je zpracováno tak, aby mohlo být využito vědeckými pracovníky napříč různými obory, kterých se problematika umělé inteligence a kognitivních procesů týká, neboť snahou bylo srozumitelně vysvětlit oblast zájmu v širokém kontextu. To je také důležitým úkolem multidisciplinárně vzdělaných odborníků – sdružovat odborníky různých vědních disciplín za účelem dosažení určeného cíle.

Cíle diplomové práce

Předložená práce si klade hned několik dílčích cílů. Prvním cílem je představit východiska dnešního přístupu ke stroji, potažmo k umělé inteligenci, která jsou dána uspořádáním světa člověka napříč dějinami lidstva. Přitom je důležité vycházet ze širokých základů daných historií, filosofií, matematikou a konče dnešním pohledem (paradigmatem) na problematiku umělé inteligence. Druhým a stěžejním cílem je představit kognitivní procesy, přičemž není mým záměrem zabývat se restriktivně jen některými přístupy k umělé inteligenci, nýbrž prezentovat tyto procesy v širších souvislostech napříč různými pohledy⁶ na umělou inteligenci a jejich srovnání s člověkem. V tomto ohledu je třeba uplatnit transdisciplinární přístup, který přináší kognitivní informatika. Třetím cílem je představit současné projekty inteligentních sys-

⁵Lidské (inteligentní) chování vs. reaktivní chování inspirované jednoduchými organizmy.

⁶V obecné rovině s přihlédnutím k základnímu dělení umělé inteligence na tzv. tradiční a novou, viz část Pojem umělé inteligence.

témů a poukázat u nich na kognitivní procesy v rámci osy zájmu a potenciál, který se nabízí pro jejich další rozvoj. Těmto cílům odpovídají i tři stěžejní kapitoly této práce. Cílem, který sjednocuje zmíněné dílčí úkoly této práce, je poskytnout ucelený pohled na problematiku poznávacích procesů umělé inteligence, a to nejenom v současném kontextu (pohledu), ale také v historických a filosofických souvislostech s výhledem do možného budoucího směřování umělé inteligence. Skrze celou práci pak prochází vztah člověka a stroje, kteří existují vedle sebe a vzájemně se ovlivňují.⁷

Postup zpracování diplomové práce

Diplomovou práci jsem rozdělil do čtyř základních kapitol. První kapitola představuje východiska a teoretický úvod do problematiky umělé inteligence s historickým a filosofickým kontextem. Čtenář se v ní rovněž dozví, co stojí v pozadí našeho současného pohledu na umělou inteligenci, a je probírán kritický pohled na přírodní vědy, které utváří současný pohled člověka na svět.

Druhá kapitola prezentuje možnosti poznávacích procesů člověka, který je velkou inspirací pro tvorbu inteligentních systémů, a umělé inteligence se zřetelem na hlavní osu zájmu. Kromě toho je v jejím závěru představen transdisciplinární přístup, který se stále více uplatňuje v posledních letech a je původcem vědeckého pokroku nejen v oblasti umělé inteligence.

Třetí kapitola navazuje volně na předchozí dvě a její úloha je pouze doplňková; její snahou je totiž ukázat, jak mohou vypadat současné inteligentní systémy, které mají interagovat s člověkem (být součástí jeho prostředí). Důraz je opět kladen na kognitivní procesy a jejich případné vylepšení. Za tím účelem představuji dva projekty inteligentních dopravních systémů řešených na Fakultě dopravní ČVUT z grantu Ministerstva dopravy ČR s názvem „Účast České republiky v projektu GALILEO“. Na uvedených výzkumných projektech jsem měl možnost se podílet jako člen řešitelského týmu projektu NEBNAK.

Poslední závěrečná kapitola obsahuje zamyšlení nad možným postmoderním vývojem a perspektivami, které nové technologie nabízejí.

⁷Člověk svým myšlením ovlivňuje konstrukci a účel strojů, přičemž stroje zase ovlivňují (ulehčují) každodenní život člověka.

1. Směřování k umělé inteligenci

Než se dostaneme k vymezení současných pohledů a možností, které umělá inteligence nabízí, tak by bylo vhodné uskutečnit stručný exkurz do historie, abychom lépe pochopili, co přispělo k současnému stavu a našemu pohledu na tuto problematiku. Svůj výklad započnu již ve starověku. Nicméně, významný mezník spatřuji až v období renesance, tedy cirká před 500 lety. Tehdy v relativně rychlém sledu přicházejí nové technické vynálezy a zároveň nastává obrat v myšlení, vidění světa a přístupu člověka k přírodě (včetně vztahu vůči sobě). [Pstružina, 1998, str. 8-9] [Kelemen, 1998, str. 85][Marková, 2007, str. 71]

Náš pohled na svět je omezen našimi poznávacími schopnostmi, což reguluje (omezuje) i naši činnost (zpětnou vazbu). [Linhart, 1976, str. 9] Někdo by to mohl shrnout slovy, že jsme vězni svého vlastního těla, ovšem to není zcela výstižné.¹ Za posledních několik století se nám podařilo posunout hranice našich nedokonalých poznávacích smyslů díky různým nástrojům za tím účelem sestrojených. Takovým příkladem může být mikroskop. Lidské oko dokáže rozlišit (vidět) dva objekty pouze, když je vzdálenost jejich obrazů na sítnici oka alespoň 150 μ m. Buňky mají v tomto ohledu velikost asi 10 μ m, člověk je tedy nedokáže vidět – slévají se. [Barbieri, 2006, str. 22] Existuje mnoho jevů a objektů, které nevnímáme nebo vnímáme jinak.² Člověk by rád šel dále a věnoval složitějším nástrojům, respektive strojům³ novou kvalitu (určitou samostatnost). Zatím stále onu úlohu převážně substituujeme hlavně my sami tím, že dané stroje přes určité rozhraní ovládáme, což je ale omezující.⁴

¹I kdybychom získali nějaký nový nativní smysl s možností přímého poznání, bylo by zapotřebí nových neuronů, jež by tyto gnose vyhodnocovaly. Bylo by tedy třeba buď současný mozek zefektivnit, nebo zvětšit. Jednoduše řečeno – náš mozek je uzpůsoben možnostem našeho těla a jeho poznávání. Pokud je naše mysl zapříčiněna mozkovou činností, tak by jakékoli změny v mozku znamenaly i změnu myšlení. Možnosti jedince však budou vždy finitní.

²Dalším příkladem může být prostorové vnímání při čítí. Pokud dva předměty vytvářející tlak na prst jsou dostatečně blízko, tak je vnímáme jako jeden předmět. Experiment se dá provést jednoduše s pinzetou.

³Srovnej v terminologickém slovníku nástroj a stroj.

⁴Kupříkladu, rychlost reakce na určitou událost by mohla být příliš pomalá, což je nevyhovující. Při současné práci člověka s počítačem počítač většinu času jenom čeká (zpracovává tzv. Idle proces) na reakci člověka, což může být zmáčknutí klávesy. Mezitím procesor vykoná miliony instrukcí bez užitku. Tento příklad využívá určitého zjednodušení, neboť moderní operační systémy se snaží dané „volno“ využít pro správu systému.

Dlouhou dobu jsme tápali, co by mohlo být onou novou kvalitou (blížící se našim kognitivním procesům), jak ji uchopit a vložit do „umělého těla“. K tomu nás přiblížila až matematika 19. století a především vznik Booleovy algebry. Kromě hmotných nástrojů jsme vytvořili i nehmotné abstraktní konstrukty, které nám též pomáhají lépe poznávat svět kolem nás. Dokážeme tak pomocí matematiky vyvodit mnohé netriviální závěry, které odrážejí skutečnost. Již několik desítek let se snažíme oba směry⁵ spojit v jeden fungující celek. Takový celek vzniká i u přístupů k umělé inteligenci.⁶

Zde bych chtěl zdůraznit, že se jedná o chování podobné, nikoli stejné jako u člověka, neboť lidé si umělou inteligenci představují právě jako umělého člověka, který je nerozeznatelný od přirozeného. „Umělá intelligence by neměla být kopií té přirozené – ani jí být nemůže, má jinou podstatu, jinak je vázána na okolí, jiná je její existence, jiný je její účel – měla by ale být schopna s ní spolupracovat na úrovni, která odpovídá konkrétnímu účelu.“ [Mareš, 2006, str. 233, dále také Zeman, 1978, str. 233, či Mařík a kol., 2007, str. 109] V současné praxi umělá intelligence směřuje spíše k fragmentům lidské intelligence (intelligence v rámci určité činnosti), primárním cílem není, aby umělý systém měl komplexní inteligentní chování podobné člověku.

Zajímavý pohled, který dále rozšíříme, na společné dějiny a vztah strojů a lidí představuje JOZEF KELEMEN ve své knize *Strojovia a agenty*. „Pomocou najjednoduchších strojov, páky, kladky, naklonenej roviny. . . postavil človek pyramídy. Dejiny strojov môžeme chápať ako dejiny jeho úsilia zdokonaliť jednoduché stroje.“ [Kelemen, 1994, str. 18] Společenství lidí, kteří vytvořili kupříkladu zmíněné pyramidy, označuje LEWIS MUMFORD pojmem megastroj. Megastroj je obrovská societa lidí, kteří se snaží naplnit cíle, které jim dala vládnoucí třída. Lidé se však snažili vymanit z velkých společenských celků a postupně začali směřovat k individualismu a liberalizmu. A tak „...sa zotročený človek minulosti, ktorého živila príroda, stal dnešným, slobodnejším človekom, ktorého živí stroj.“ [Kelemen, 1994, str. 18] Díky našemu poznávání přírody jsme vytvořili „silné stroje“, které nám ulehčují práci. Na druhou stranu poznáváme i sami sebe, což otevírá dveře k „šikovným strojům“. Obloukem se tak dostaneme k důležitosti lidského poznání, tvořivosti, otevřenosti společnosti a jejich významu nejen pro jednotlivce, ale i pro lidstvo. „Technické výtvořy človeka vhodne

⁵Výpočetní stroj a program naprogramovaný např. ve Scheme, Lispu, Haskellu, C#, Javě nebo jiném vysokoúrovňovém jazyku využívajícím framework, které současné přístupy k umělé inteligenci využívají, patří nerozlučitelně k sobě.

⁶Jak bylo řečeno výše, cílem je osamostatnit stroje. Hezky to vyjadřuje Marvin Minsky z MIT (1967), podle něhož je „... umělá intelligence věda, jejímž úkolem je naučit stroje, aby dělaly věci, které vyžadují inteligenci, jsou-li prováděny člověkem.“ [Berka, 2003, Habiballa, 2004, str. 7]

začleňované do struktury jeho činnosti mali obrovský emergentný vplyv na rozvoj civilizácie i v minulosti.“ [Kelemen, 1994, str. 26]

Na jedné straně vytváříme stále dokonalejší „umělé tělo“ ve formě počítačového stroje, na druhé straně vymýšlíme programy, které jej využívají. Náš dualistický pohled na naše tělo a uvnitř schovanou novou kvalitu (řekněme mysl) se odráží i v přístupech k umělé inteligenci. Například křesťanský pohled,⁷ který je pro naši kulturu typický, nabízí tělo a duši. Tento pohled nás dost možná i svazuje – bereme ho v mnohém jako samozřejmost, dává totiž zapravdu naší přímé zkušenosti vidění člověka (tělo a mysl/duše uvnitř). Obdobný pohled v jiném kontextu přinesl v roce 1909 i WILHELM JOHANNSEN v biologii se svou koncepcí živé bytosti, která je duálním systémem tvořeným fenotypem (hardware, tělo) a genotypem (software, DNA). [Barbieri, 2006, str. 32] Dnešní pohled na počítač je v tomto ohledu dost podobný. Existuje zde hardware a software, neboli stroj a jeho program. Pokud se vrátíme k myšlence těla a uvnitř schované mysli, tak ve filozofii nalezneme dva přístupy – dualismus a monismus. V kontrastu s dualismem stojí monismus – veškeré mentální funkce jsou důsledkem fyzického těla – tělo a mysl jedno jest.⁸ [Peregrin, 2008, str. 91] Oba filosofické pohledy, jak starší dualismus, tak monismus, spadají pod filosofii mysli a zabývají se člověkem.⁹ Dle mého názoru však dualistický pohled na dnešní stroj je více vyhovující (oblast mysli ještě tak prozkoumanou nemáme), nicméně monistický pohled je zajímavý až při pohledu do budoucnosti (viz podkapitoly 1.2 a 1.5 s odkazem na umělý život).

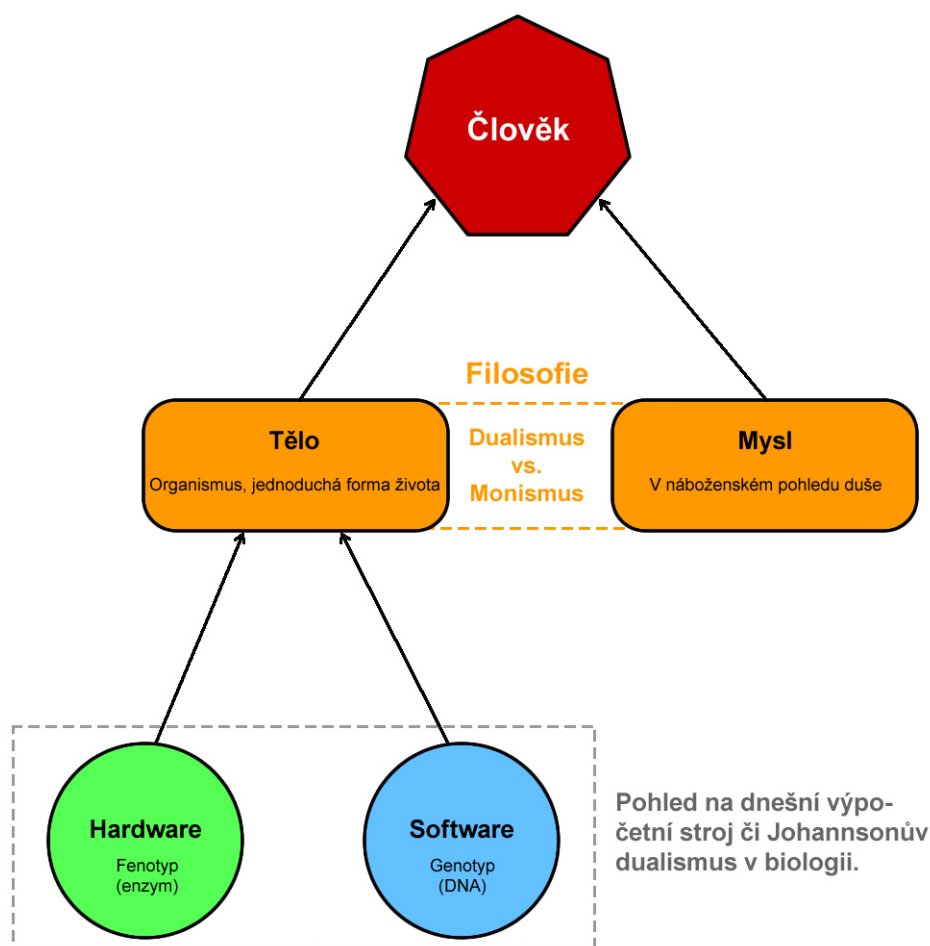
Výše vyličené pohledy na mysl a tělo, ke kterým se v krátkosti vrátím v podkapitole 2.2., najdou i své kritiky, především z hlediska „omezenosti alternativ“ na monismus či dualismus, jež jsou dány na výběr a brání nám tak nepřímou v jiném uchopení problematiky. [Searle, 1994, str. 14] Mezi takové patří i J.R. SEARLE, který řeší problém vztahu těla a mysli v jednoduchosti tak, že mozek způsobuje mysl a mysl je vlastnost mozku obdobně jako tuhost je vlastnost stolu. Přesto odmítá striktně dualistické i monistické koncepce a dodává: „Z mého pohledu se sice mysl a tělo vzájemně ovlivňují, nejsou to však dvě odlišné věci, neboť mentální fenomény jsou prostě vlastností mozku. Toto pojetí lze vyložit i tak, že v něm zastávám jak fyzikalismus, tak mentalismus.“ [Searle, 1994, str. 28] Přesto dle mého názoru při zkoumání člověka musíme

⁷Typičtějším křesťanským dualismem je Dobro a Zlo, neboli Bůh a Ďábel.

⁸Více informací například o materialistickém monismu naleznete v [Zeman, 1978].

⁹V souvislosti s člověkem a filosofií mysli uvádí KAREL PSTRUŽINA spíše pojem identismus než monismus. Významově však oba myslí to samé. Srovnej: [Peregrin, 2008, str. 91] a [Pstružina, 1998, str. 62]

mít k tělu a mysli kvalitativně jiný přístup, ale zároveň být si vědomi provázání – v tomto ohledu a z dnešního pohledu mi jeho koncepce připadá spíše dualistická i přesto, že se snaží najít nové místo v pohledech na tuto problematiku. Spíše si tedy myslím, že jeho pohled byl časem přiřazen k těm dualistickým.



Obrázek 1.1.: Přehledné představení zmíněných dualistických koncepcí.

Výše zmíněný dualistický pohled¹⁰ ponechám i ve svém dalším výkladu, kdy na jedné straně budu sledovat historický vývoj lidských nástrojů a později (výpočetních) strojů, na druhé straně mě zajímá i vývoj matematiky, idejí, které mají odraz v realitě a které vedou k dnešnímu pojmu software. Je zajímavé udělat si širší rozhled, neboť většina odborných článků se orientuje pouze na relativně nedávnou historii 20. století a nezabývá se širším kontextem. To vše spoluutváří naše současné poznávací schopnosti, jež využíváme a snažíme se je „zrcadlit“ i v přístupech k umělé inteligenci.

¹⁰Zajímavý výchozí dualistický pohled (karteziánský dualismus) pro úvod do umělé inteligence naleznete kupříkladu v [Carter, 2007].

Následující podkapitola tedy ilustruje jednotlivá období, která jsou oddělena přelomy, kdy se lidstvu otevřely nové technické možnosti nebo nastal významný obrat v myšlení a chápání světa. Nesnažím se o kompletní výčet, nýbrž jen o rychlý přehled, východiska, která nám pomůžou s pochopením současného stavu přístupů k umělé inteligenci, respektive k jejím poznávacím procesům.

Pro jednoduchý přehled jsem poznámkami na okraji označil „Směřování k výpočetnímu stroji“ jako [A] a „Směřování k současné matematice přírodních a technických věd“ jako [B].

1.1. Pohled do minulosti strojů a abstraktního myšlení

Na začátku lidské invence zpravidla stojí určitý praktický problém, alespoň tomu tak bylo dříve. Dnes už to tak docela neplatí, neboť mnohé problémy, které dnešní věda řeší, se nemusí nutně dotýkat našich každodenních potřeb či obecněji našeho okolí (přímého kontextu).

Takovým reálným problémem dozajista bylo, jak zjednodušit počítání především pro obchodníky či správce. Když se podíváme až do starověku, tak zde nalezneme první počítací nástroj, který byl rozšířený od Japonska až po Evropu – jmenoval se abakus, nebo také suan-pchan (Čína) či sarob-jan (Japonsko). Principiálně byl podobný dětskému počítadlu a aktivně se používá asi 3500 let.¹¹ Tato jednoduchá pomůcka umožňovala složitější výpočty, kdy počtáři stačilo umět počítat jen do deseti (v desítkové soustavě). [Mareš, 2006, str. 242] [A]

Problematictější je ovšem hovořit o matematice starověku, neboť vycházíme pouze z dochovaných útržků. Antické Řecko (po roce 800 př. n. l.) jako kolébka evropské kultury vycházela z mnohem starších poznatků v oblasti matematiky – Mezopotámie (Babylon 2000 př. n. l.) a Egypta (především 6. století př. n. l.). Za nejplodnější období řecké matematiky, respektive geometrie¹² můžeme považovat 6. - 3. století př. n. l., které začíná dobou THALETA Z MÍLÉTU (624 - 548) a končí ARCHIMÉDEM (287 - 212), či vlivem Říma a později obsazením Řecka Římem ve 2. století př. n. l. V době, kdy vládne Evropě Řím, se matematika nerozvíjí tak rychle jako dříve - dá se říci, že stagnuje. Ačkoli Řekové měli velké zásluhy v oblasti geometrie, už to tak neplatilo v oblasti algebry, myšleno z hlediska počtů. Co se týká určité formalizace a uchopení [B]

¹¹První zmínky pocházejí již z Akkadské říše 2500 let př. n. l.

¹²Geometrie (zeměměřičství) se původně zabývala definicemi geometrických objektů, jejich vlastnostmi a vztahy. [Černý, Kočandrlová, 1998, str. 7] O matematice v dnešním širokém pojetí ještě nelze mluvit.

světa, tak nesmíme opomenout ARISTOTELŮV (384 - 322) odkaz, který později vyústí v dnešní axiomatický systém. Naopak velkou tradici v algebře měly národy v oblasti Mezopotámie (především zásluhou Babyloňanů a Chaldejců).

Až počátkem středověku se řecké, mezopotamské a vzdálené indické učení spojí díky arabským učencům a otevírá se nový pohled, který je podobný dnešnímu přístupu k matematice. Začíná nový vzestup a rozvoj, nejdříve v arabském světě a následně i v křesťanské Evropě, jenž vyvrcholí v novověku. Evropská matematika středověku musela mnohé teprve objevit. Využívala římské číslice, což velmi znesnadňovalo veškeré počty, takže se počítalo hlavně do deseti. Pak se buď využíval abakus, nebo vágní pojmy.¹³ [Petrů, 2007, str. 26] Obrat přichází až s překlady arabských děl od 12. století. Jedná se především o díla AL-CHOREZMÍHO (780-850), jež tento arabský učenec napsal v 9. století (Číslo nula, Algebra, Algoritmus, Neznámá X) . Evropa také částečně upouští od římských číslic a nahrazuje je arabským, respektive indickým symbolickým zápisem číslic vhodnějším pro výpočty. K tomu však zcela dojde až v 16. století. [Petrů, 2007, str. 27]

Matematika nebyla středem zájmu středověkých učenců, a až na některé výjimky (PIERRE ABÉLARD, GERBERT Z AURILLACU neboli papež SILVESTR II., MIKULÁŠ ORESME) se jí zabývali spíše okrajově. Počátkem 13. století italský učenec LEONARDO FIBONACCI¹⁴ (1180 - 1250) vydává „Knihu počtů“ (Liber Abaci), kde představuje arabská čísla a arabskou algebru. Opravdový obrat přichází až s příchodem humanismu a renesance do Itálie ve 2. polovině 14. století. Evropa objevuje ztracený odkaz starověku a stává se jeho pokračovatelem. Přicházejí postupně ústřední díla tzv. „vědecké revoluce“ počínající dílem MIKULÁŠE KOPERNÍKA (1473 - 1543) „Šest knih o obězích sfér nebeských“ (De Revolutionibus orbium coelestium libri VI), sahajícím až k práci ISAACA NEWTONA (1643 - 1727) a knize Matematické principy přírodní filozofie (Philosophiae Naturalis Principia Mathematica).

Další důležitý zvrat přichází v novověku s logaritmickou stupnicí, načež vzniklo logaritmické pravítko (1623, WILLIAM OUGHTRED a EDMUND GUNTER). [Mareš, 2006, str. 247] Nebyl to ovšem jediný stěžejní objev své doby. Začala se objevovat různá mechanická počítadla, která usnadňovala a zrychlovala početní operace. Náčrtky takových počítadel nalezneme již u LEONARDA DA VINCIHO (1452 - 1519), na jejichž základě byly později zhotoveny funkční vzorky. Mezi konstruktéry těchto strojů nalezneme jména jako PASCAL (1623 - 1662) či LEIBNIZ (1646 - 1716), kteří své počítací

[A]

¹³Například: „o něco více“, „součástka průměrné délky“

¹⁴Syn italského obchodníka, který hodně cestoval po Středozezemním moři a po arabském světě.

stroje vytvořili v polovině 17. století. Byl to věk plný zvratů a tomuto rozvoji pomohl i vývoj matematiky. V této době ještě můžeme mluvit o relativně jednoduchých výpočetních operacích, jež zahrnovaly sčítání, odčítání a násobení. Rozvoj těchto počítadel a počítání závisel na rozvoji přírodních věd, obecné vzdělanosti a především poptávce po takových vynálezech. Většině lidí tak až do začátku 20. století zpravidla stačilo umět základní počty, případně ovládat jednoduchá počítadla či v lepším případě použít logaritmické pravítko. Přesto se uvádí ještě jeden milník, a to tzv. mechanický počítací stroj CHARLESE BABBAGE (1791 - 1871), přesněji jeho analytický stroj. Pracoval na něm s přestávkami několik desítek let, ovšem nikdy jej během svého života (zemřel r. 1871) funkčně nedokončil. Jednalo se na tu dobu o převratný počín, neboť tento stroj byl již programovatelný a využíval i děrných štítků. Fungující Babbageův počítací stroj byl zhotoven až v roce 1991, přičemž vážil několik tun. [Mareš, 2006, str. 250]

Na začátku 20. století přichází mechanická a později elektromechanická kalkulačka. Vystává nutnost přesnějších a náročnějších výpočtů. Lidé začínají konstruovat dříve fantaskní stroje, které dokážou létat nebo plout pod hladinou moře. Je to opět doba rychlých změn. Čísla začínají být čím dál více abstraktní a vzdálená každodennímu životu.¹⁵

Pokud bych měl obsáhnout vývoj matematiky od 16. století do počátku 20. století, asi by tato práce musela být na jiné téma. Proto udělám menší úkrok stranou a budu se zabývat matematikou, která se snažila uchopit svět a která by byla dále využitelná i pro poznávací procesy umělých systémů. Již jsem se zmiňoval o významné roli ARISTOTELE. V oblasti kategorizace jsoucna na něj později navázali mnozí myslitelé (např. JOHN WILKINS, 1614-1672¹⁶). Kromě toho se ARISTOTELES snažil formalizovat jazyk, aby lépe popisoval svět. V tomto ohledu na něj navazuje právě novověk, který se ubírá směrem k automatizovanému uvažování. Zavádí se symbolizace a formalismus práce s nimi. S touto myšlenkou přichází i LEIBNIZ, kdy předpovídá, že budoucí spory už nebudou řešit dva filozofové rozmluvou, nýbrž matematici. [Petrů, 2007, str. 35] Pravdy se budeme moci dobrat „počítáním“, což byl všeobecný optimismus té doby. V průběhu 19. století se díky problematice nekonečna vytváří teorie množin, jejíž zá-

[B]

¹⁵Dříve si člověk dokázal pod určitým pojmem představit cosi hmotného – 5 krav, $1+3=4$ litry mléka. Matematika se však stále vzdalovala a ztrácela běžným lidem ve víru neznámých veličin, vzorečků a vět. Složitě výpočty se tak v praxi snaží rozdělit na řadu jednodušších, a ty pak lidé automatizovaně zpracovávají, aniž by přesněji věděli, o co jde. Prostě jenom násobí čísla, podobně jako dnešní osobní počítače.

¹⁶Viz [Borges, 1999]

klady položil CANTOR (1845 - 1918). Z jeho myšlenek dále čerpal počátkem 20. století FREGE (1848 - 1925), jenž vytváří axiomatický systém, který se ukáže jako sporný (nekonzistentní). Vylepšený axiomatický systém přináší až ZERMELO (1871 - 1953) a upřesňuje jej FRAENKEL (1891 - 1965),¹⁷ vzniká tzv. Zermelo-Fraenkelova teorie množin. [Mareš, 2006, str. 154] Díky tomu můžeme začít vytvářet formální teorie a zároveň prokazovat správnost tvrzení určité teorie. Základy axiomatického přístupu nalezneme v antickém Řecku, kde matematik EUKLEIDÉS (325 - 260) ve svém díle „Základy“ zavedl pět geometrických axiomů, pomocí nichž byl schopen logicky odvozovat všechny v té době známé geometrické pravdy. Ovšem do 20. století se v matematické praxi spíše užívala intuice, a tak celá konstrukce teorie vypadá spíše jako názorný „logický“ postup myšlenek k nějakému cíli. Axiomatický systém nyní přináší pevné základy, na kterých se dá stavět a díky kterým lze ověřit správnost jakéhokoli tvrzení. Všeobecný optimismus kalily některé paradoxy, které se začaly objevovat především u složitějších teorií. V roce 1920 přichází tzv. „Hilbertův program“ (podle DAVIDA HILBERTA, 1862 - 1943), který se snaží o vypracování dokonalé axiomatiky, jež by zamezila i těmto paradoxům. Naděje se rozplynou o jedenáct let později, kdy GÖDEL (1906 - 1978) předloží důkaz vět o neúplnosti axiomů, jež vlastně stanovují hranice axiomatických metod. [Mareš, 2006, str. 152] Přesto matematika ve službách fyziky slaví velké úspěchy a v průběhu 20. století přináší závratné objevy.

Od 30. let 20. století můžeme mluvit o úsvitu počítačů, kdy se objevují první reléové výpočetní stroje (tzv. nultá generace). Počítače zpracovávají algoritmy, které se využívají také v matematické dedukci, kdy by se měla z axiomů formální cestou za pomoci určitého kalkulu zjistit pravdivost či nepravdivost libovolné formule. ALAN TURING (1912 - 1954) nezávisle s ALONZEM CHURCHEM (1903 - 1995) ovšem dokázali v letech 1936 - 37, že „existuje třída matematických problémů, které nemohou být vyřešeny žádným určitým jednoznačným procesem ani žádnou heuristickou procedurou.“ [Petrů, 2007, str. 49] K tomu vymyslel abstraktní tzv. Turingův stroj (1937), na kterém poukázal na problém zastavení se Turingova stroje, a tedy nerozhodnutelnost některých formulí. Bortí se tak Hilbertovy plány na dokonalou axiomatiku. Zatímco GÖDEL dokazuje, že formální axiomatika nemůže být úplná ani bezesporná,¹⁸ tak TURING předesílá, že formální axiomatika nemůže být vždy rozhodnutelná.

V roce 1948 vytváří WIENER (1894-1964) základy kybernetiky a vydává knihu „Kybernetika aneb řízení a sdělování u organismů a strojů“, kde poukazuje na obecné

¹⁷ZERMELO jej přináší v roce 1908, nicméně FRAENKEL jej vylepšuje o jedenáct let později.

¹⁸Srovnej: Gödel, K.: Filosofické eseje. Praha: Oikymenh, 1999

principy zpracování informací v umělých systémech a živých organizmech.

V souvislosti s tím se objevují i otázky, zda je možné, aby výpočetní stroje myslely. Vychází se z předpokladu, že naše mysl je abstraktní funkce (existuje zde tedy algoritmus), kterou lze provádět nezávisle na fyzické podstatě (je jedno, jestli jde o lidské tělo nebo výpočetní stroj). V roce 1950 přichází TURING se svými úvahami¹⁹ a testem, jenž představuje nový pohled na pojem inteligence a inteligentní chování.²⁰ Nicméně až konference na Dartmouth College v roce 1956 se datuje jako vznik „umělé inteligence“ a formulování oblasti jejího zájmu.²¹ [Mařík a kol., 1993, str. 19]

Když se nyní budeme soustředit již výlučně na tuto novou vědní disciplínu, tak je třeba podotknout, že problematika, kterou se zabývá, je interdisciplinární. Účelem umělé inteligence bylo řešit problém podobně, jako ho dokáže řešit člověk. Nejdříve v 50. a 60. letech přišly heuristické metody prohledávání stavového prostoru (1956) a tzv. rezoluční princip (1965) [Jiroušek, 1995, str. 9], přičemž za účelem snadnějšího programování byly vyvinuty i speciální programovací jazyky LISP či PROLOG. Po mírné krizi koncem 60. let, kdy vědci pochybovali o některých tehdejších přístupech²² či poukazovali na nové problémy, přichází opět čas boomu díky expertním systémům založeným na znalostech. Přichází i praktické využití expertních systémů (MYCIN, PROSPECTOR), což zvýšilo obecně úsilí v této vědecké oblasti. [Mařík a kol., 1993, str. 20-22] V současné době mluvíme již spíše o intelligentních systémech, např. telematických dopravních systémech. V souvislosti s expertními systémy se také mluví o zpracování neurčitosti, neboť i v našem každodenním životě běžně pracujeme s neurčitostmi. V současné době se touto oblastí zájmu zabývá tzv. soft computing.²³ Jedná se stále o symbolickou perspektivu vyjádření neurčitosti.²⁴

Trošku jiný přístup přináší ROSENBLATT (1928 - 1971), který se inspiruje živým systémem a představuje model neuronu (1958). Postupem času vznikají další přístupy k umělé inteligenci, odrážející chování přírody – evoluční algoritmy (60. léta) či multi-agentní systémy (80. léta). Tyto přístupy využívají oproti výše zmíněným (spíše symbolicky orientovaným) komplexitu systému.

¹⁹Viz Turingův článek v časopise *Mind* s názvem *Computing Machinery and Intelligence*

²⁰Neboli co by měl umělý systém splnit, abychom o něm mohli tvrdit, že je inteligentní.

²¹Srovnej: [Kelemen, 1994, str. 32]

²²Kritická kniha MINSKÉHO a PAPERTA o perceptronu a jeho limitním použití. Viz též [Petrů, 2007, str. 70]

²³Viz [Mařík a kol., 1997]

²⁴To vede k zamyšlení, že při lidské práci s neurčitostí nevycházíme z čísel či pravděpodobností, nýbrž z našich intencionálních nepopsatelných zkušeností. Reprezentace takových zkušeností je samozřejmě problematická.

Kromě nejrůznějších metod, které se uplatňují v umělé inteligenci, zde vyvstávají opět otázky položené již v 50. letech 20. století: „Mohou počítače myslet, podobně jako člověk?“ Tímto spíše filozofickým problémem se zabývá také SEARLE, především v souvislosti s komputacionismem,²⁵ který tělo představuje jako hardware a mysl jako softwarový program. Lépe řečeno, mysl je v tomto podání vlastně složitou funkcí, kterou realizuje náš mozek.

Umělou inteligenci dále rozděluje na tzv. silnou a slabou. Slabá UI by měla projít Turingovým testem, což v původním znění Turingova testu znamená, že dokáže velmi dobře imitovat chování člověka, takže běžný člověk nepozná, že komunikuje se strojem (Turingovo porozumění). Silnou UI „chápeme jako porozumění takové, že systém bude disponovat pocitem chápáním takovým, jakým disponuje lidská mysl.“ [Pěchouček, 2005]

Toto rozdělení přináší i některé problémy, především z hlediska subjektivního pohledu a vyložení daných definic. Někteří slabou UI vidí v systému GPS,²⁶ jiní již v inteligentní automatické pračce,²⁷ přičemž původním Turingovým testem by samozřejmě neprošly, ale jeho variantami ano. Naopak silnou UI někdo vidí v expertním systému PROSPECTOR nebo MYCIN,²⁸ ovšem jiní by namítali, že zde chybí onen „pocit chápání“ jako u nás lidí. [Habiballa, 2004, str. 6] [Pěchouček, 2005]

Searlova argumentace směřuje k tomu, že mysl není pouhý program v dnešním slova smyslu, a podpírá ji příkladem „Čínský pokoj“ (1980).²⁹ Kromě této „námitky vědomí“ [Petrů, 2007, str. 53] přicházejí i další problémy³⁰ ve směru striktního komputacionismu. Dalším myšlenkovým proudem je konekcionistická teorie založená ve vši jednoduchosti na komplexní síti propojených jednotek, které se vzájemně ovlivňují. Mohou tak vznikat emergentní (nové) vlastnosti, které u jednotlivých jednotek

²⁵Také se nazývá komputacionistický funkcionalismus. Toto označení se používá především ve filosoficky orientovaných textech (filosofie mysli), přičemž modelování mysli se děje na logicko-symbolické úrovni, a spadá tedy pod tradiční umělou inteligenci, viz část 1.2. Pojem umělá inteligence.

²⁶Řekne člověku přesnou polohu a člověk ani nemusí při určení této polohy rozeznat, zda mu údaje poskytl člověk nebo počítač.

²⁷V rámci svého omezeného prostředí zjistí stupeň zašpinění prádla a zvolí vhodný program k vyprání – stejně by to udělal člověk. Kdo by pak rozeznal, jestli to vyprala automatická pračka dle voleb člověka, nebo se o vše postaral stroj sám?

²⁸V rámci oblasti svého působení dokážou vyvodit z určitých dat relevantní závěr stejně jako člověk. Dokonce může vzít v úvahu více předpokladů podobně jako člověk, nicméně zde chybí ony neuchopitelné a nepředatelné zkušenosti (intencionální), které bere na druhé straně v potaz člověk. Viz [Křemen, 2007, str. 9-13] [Mařík a kol., 1997, str. 244]

²⁹Ve vši jednoduchosti jej lze shrnout do následujícího: Počítač dokáže pracovat se symboly podle programu a tím vytváří zdání inteligentního chování, což člověk dokáže také. Člověk však navíc dokáže intencionálně porozumět danému problému. Podrobněji vysvětlený příklad „Čínského pokoje“ naleznete v [Searle, 1994, str. 33].

³⁰Často se uvádí „námitka matematická“, která byla formulována ROGEREM PENROSEM, přičemž vychází kromě jiného i z Gödelovy věty o neúplnosti. [Petrů, 2007, str. 48]

nepozorujeme. Komputacionistická teorie má základy v dnešní matematice, ovšem konekcionistická teorie spíše v přírodních komplexních systémech. Jako příklad konekcionistického přístupu se uvádí umělá neuronová síť, kde jednotlivé uzly (tělo neuronu) jsou spojeny různě silnými vazbami s dalšími uzly.

Lze tedy na umělou inteligenci pohlížet z mnoha úhlů, záleží vždy jen na tom, čeho chceme dosáhnout. K výše zmíněným přístupům se budu v jednotlivých kapitolách vracet a prohlubovat je dle potřeby. V pozadí tohoto vývoje stojí matematika, na jejíchž principech dnes stavíme veškerou softwarovou výbavu počítače. Matematika se díky automatizaci výpočtů (za pomoci výpočetních strojů) opět vrátila zpět do každodenního života člověka, i když v pozměněné formě. Nová generace lidí dokonce intuitivně přemýšlí jako počítače, respektive softwarové programy v nich.³¹ Vytváří se tak virtuální prostředí, které se stává naším každodenním údělem a do značné míry nás ovlivňuje.

Výpočetní stroje se od konce 2. světové války rychle rozvíjely a postupně se rozšířily až do domácností, tak jak to známe dnes. Jediná věc, která všechny běžné počítače spojuje, je jejich konstrukce dle tzv. von Neumannova schématu.³² Nechtěl bych se zde zabývat samotnou historií výpočetní techniky posledních šedesáti let, to lze nalézt v mnoha populárních publikacích. Vývoj nástrojů a strojů určených k výpočtům lze v zásadě rozdělit do čtyř etap, kdy každý pokrok byl spíše skokový a souvisel zpravidla i s dobou, technologiemi a rozvojem abstraktního myšlení. [A]

- Starověk a středověk – nejrůznější pomůcky pro složitější výpočty (abakus)
- Novověk – složitější nástroje (logaritmické pravítko) a mechanické kalkulátory
- Počátek 20. století – elektromechanické stroje, včetně reléových počítačů
- Od 40. let 20. století – elektronické výpočetní stroje

Troufám si říct, že myšlenky předběhly možnosti techniky. Mnohé moderní rysy dnešních počítačů tak nalezneme již v 19. století, jen způsob provedení byl v té době limitující.

³¹Je zajímavé pozorovat při práci s počítačem dítě, které si velmi rychle osvojuje ovládání a dokáže lépe interagovat s počítačem než starší generace, což je dáno i možnostmi „mladého mozku“. Tuší, kde je v systému chyba, a dokáže intuitivněji reagovat, doslova se vcití do chování počítače. Mozek se přizpůsobuje i těmto virtuálním prostředím, které spojuje s reálným v jeden celek.

³²Viz <http://www.fi.muni.cz/usr/pelikan/ARCHIT/TEXTY/VNEUM.HTML>

1.2. Pojem umělá inteligence

Po historickém exkurzu by bylo vhodné alespoň zmínit proudy, architektury či dělení umělé inteligence včetně toho, co pro nás v tomto pojednání umělá inteligence bude představovat.

Na úvod bych začal klasickým problémem: Co si představit pod pojmem umělá inteligence – jaká je její definice? Pokud projdeme odbornou literaturu, tak naleznete mnoho definic, které odrážejí svým způsobem i dobu svého vzniku a pohled na umělou inteligenci. Za posledních šedesát let je vidět i posun v pohledu (veřejnosti).

Umělá inteligence byla ve svých začátcích v očích veřejnosti spojována s logickým usuzováním, s možností rychlých matematických operací. Souvisí to tak trochu i s viděním tehdejších vědců v očích společnosti, kteří konstruovali divotvorné stroje a museli být velmi zběhlí v matematice a obecně v přírodních vědách. Jako by v tom byl ukryt jejich intelekt a tvůrčí schopnosti. Časem nás v této doméně (logického usuzování) počítačí stroje předběhly – znamená to tedy, že jsou z našeho pohledu intelligentní? Nikoli, stále něco chybí. Postupem času jsou definice opatrnější a obecnější s posunem k vědomí.³³ Jako bychom se zastavili a zamysleli se: Máme stroje, které jsou v automatickém logickém usuzování lepší než sám člověk, ale pořád jim něco chybí, něco, co má i pětileté dítě, kterému předcítáme pohádku – vědomé porozumění příběhu.[Hayward, Varela, 2009, str. 162] Podívejme se na některé zajímavé definice:

Minského definice (1967):

Umělá inteligence je věda, jejímž úkolem je naučit stroje, aby dělaly věci, které vyžadují inteligenci, jsou-li prováděny člověkem. [Berka, 2003]

Kotkova definice (1983):

Umělá inteligence je vlastnost člověkem uměle vytvořených systémů vyznačujících se schopností rozpoznávat předměty, jevy a situace, analyzovat vztahy mezi nimi, a tak vytvářet vnitřní modely světa, ve kterých tyto systémy existují, a na tomto základě pak přijímat účelná rozhodnutí za pomoci schopností předvídat důsledky těchto rozhodnutí a objevovat nové zákonitosti mezi různými modely nebo jejich skupinami. [Mařík a kol., 1993, původně Kotek Z., Mařík V., Zdráhal Z.: Metody rozpoznávání a umělá inteligence. In: Kybernetika ve výzkumu a výuce, ČSVTS FE VŠSE, Plzeň, 1983, s. 16-30.]

Definice Richové (1991):

³³Viz. výše uvedený Searlův „Čínský pokoj“.

Umělá inteligence se zabývá tím, jak vytvořit počítače, které by dělaly věci, ve kterých je v tuto chvíli lepší člověk. [Akerkar, 2005, str. 2] Původně byla tato definice publikována ELAINE RICHOUVOU v roce 1991.

Definice Bourbakise (1992):

Umělá inteligence je věda, která studuje chování inteligentních systémů. Intelligentními systémy se v tomto případě myslí takové systémy, které jsou navrženy a naprogramovány chovat se analogicky (nikoli nutně stejně) jako některé lidské akce jako cítění, řešení problémů, učení se pravidel atd. [Bourbakis, 1992]

Jednotlivé definice více či méně trefně vymezují to, jak autor nahlíží na umělou inteligenci, a korespondují s cílem publikace či článku, ve kterém se nacházejí. Jako nejtrefnější a zároveň dostatečně obecná se často uvádí právě definice MARVINA MINSKÉHO principiálně vycházející z Turingova testu. Pro účely tohoto pojednání nám bude postačovat obecnější heslo: „*Umělá inteligence by měla především pomáhat člověku a rozvíjet možnosti v celé šíři jeho činnosti.*“ Stále je to svým způsobem jen nový druh (virtuálního) stroje a účelem stroje, stejně jako dříve, je kromě ekonomických aspektů i osvobodit člověka od těžké práce (fyzické i psychické). Na jedné straně to zní naivně – „pomáhat“, když stroje v našich rukách mohou i zabíjet, na druhou stranu zavádění „umělé inteligence“, ať už si pod tím představíme cokoli, trochu naznačuje, jako by nás takový stroj neosvobozoval od těžké práce, ale spíše od (svobodného či tvůrčího) myšlení. Přesto všechno toto heslo zdůrazňuje onen cíl, který vychází i z historického směřování k umělé inteligenci, a je důležité ho mít stále na mysli. Právě tato účelnost vše spojuje.

V souvislosti s výše řečeným heslem by bylo vhodné zmínit se i o jednom z řady (filozofických) pohledů uplynulého 20. století. Jedná se o problematiku tzv. nehumanismu.³⁴ [Sim, 2003, str. 19] V jednoduchosti jde o to, že kvůli rychlému tempu vývoje³⁵ dochází k „dehumanizaci“. ³⁶ Na místo člověka nastupují sofistikované počítačové systémy, které ho vytlačují. Dokonce sám člověk směřuje díky pokročilému lékařství pomalu k „novému člověku“ – kyborgovi. Končí tímto lidství tak, jak jej známe? Ačkoli byl člověk osvobozen jako jedinec v období humanismu a začal rozvíjet přírodní vědy, jež mu otevřely nové dveře poznání a svobody, přesto jsou to paradoxně právě výtvořky člověka (přírodních věd), které ho nyní opět svazují a ničí tím jeho lidskost. A právě

³⁴Představitelem myšlenky je postmoderní filosof JEAN-FRANÇOIS LYOTARD (1924 - 1998).

³⁵Myslí se tím například kapitalismus, který tlačí na efektivitu, rychlost, přesnost, což naráží na limity člověka, a nabízí řešení ve formě akurátního stroje.

³⁶Porovnej: [Kelemen, 1994]

v souvislosti s tím se mluví o umělé či nelidské budoucnosti, jejíž počátky nalezneme již dnes v umělém životě a umělé inteligenci. Přichází tedy otázka, zda výše zmíněné heslo „pomáhat člověku“ již nenabírá nového významu v tomto kontextu budoucnosti, kdy lidské a nelidské splývají do jednoho. Omezím se na konstatování, že umělá inteligence pomáhá dravému a dost možná i „nelidskému“ vývoji.

S tím souvisí má poslední poznámka, že v současné době chápeme umělou inteligenci jako oddělenou oblast od zbytku, zpravidla hardwarového, stroje. Je to oddělená kvalita, což podtrhuje i dualistický pohled vytvořený historií.³⁷ Umělý život, tak jak jej známe dnes, působí velmi primitivně a rozhodně se nepodobá vyšším formám života. „Jedním z hlavních cílů umělého života je totiž studium zákonitostí života, které jsou nezávislé na mediu, v němž se odehrává.“ [Wiedermann, 2004] Náznaky budoucího směřování ve vzdáleném horizontu ovšem předpokládají spojení obou oborů, tedy umělé inteligence a umělého života, do jednoho – řekněme inteligentního – umělého života.³⁸ S tím potom filozofové jako LYOTARD také pracují při svých úvahách. V dnešní době se „umělá inteligence zabývá znaky, které se v reálném životě objevily v závěru evoluce, zatímco umělý život simuluje vše co se objevilo na počátku.“ [Barbieri, 2006, str. 32] I když náznaky propojení tu jsou i dnes (umělý život a nová umělá inteligence).

Mé pojednání však bude dodržovat současný pohled na problematiku s cílem naznačit směřování v blízkém horizontu.³⁹ Přesto nelze umělý život opomenout, neboť stejně jako umělá inteligence se snaží vytvořit kvalitu, kterou nacházíme v tom, čemu říkáme život (živá část přírody). A život umí být inteligentní. Ale vraťme se k současné umělé inteligenci.

Umělou inteligenci nejčastěji rozdělujeme do dvou hlavních skupin:

- **Tradiční či klasická umělá inteligence** je založena na vnitřní logicko-symbolické reprezentaci světa a schopnosti s touto reprezentací pracovat – konat rozhodnutí vzhledem ke svým cílům (deliberativnost). Modeluje se shora dolů. Více viz [Kelemen, 1994, str. 63][Pstružina, 1998, str. 136]
- **Nová umělá inteligence** je založena na myšlence jednoduché reaktivity, kdy i složité chování lze získat jednoduchými reakcemi – není třeba racionální usuzování. Základními myšlenkami jsou: emergentní funkcionalita (nová kvalita,

³⁷ Oddělené principy a základy pro vývoj hardwaru a softwaru.

³⁸ Profesor HUGO DE GARIS mluví o „Artilektech“ – mohutné systémy umělé inteligence, které nás předčí a nahradí jako součást přirozeného vývoje. [Sim, 2003, str. 65] Více v kontroverzní, leč zajímavé knize *The Artilect War: Cosmists vs. Terrans*.

³⁹ Lyotard by zřejmě mluvil o období posthumanismu.

kterou jednotlivé komponenty nemají, vzniká jejich interakcí), dekompozice na úrovni úloh a reaktivita. Modeluje se zdola nahoru. Více viz [Mařík a kol., 2001, str. 161][Kelemen, 1994, str. 64][Pstružina, 1998, str. 148]

Kromě tohoto rozdělení nalezneme i rozdělení na silnou a slabou umělou inteligenci. Toto rozdělení bylo zmíněno již výše a ve vší stručnosti lze říci, že silná UI myslí jako člověk, naopak slabá UI pouze jedná jako člověk. K tomu se vážou i pojmy (myšlenkové proudy) komputacionismus a konekcionismus, které jsou vysvětleny také v předchozí podkapitole či terminologickém slovníku.

Na závěr zmíním ještě tzv. distribuovanou umělou inteligenci. Nová umělá inteligence, jež přinesla jednoduchého reaktivního agenta, tak ovlivnila i tradiční umělou inteligenci, jež jako odpověď představila deliberativního (uvažujícího) agenta. Další myšlenkou bylo využít distribuovaného či paralelního zpracování informací v rámci větší skupinky agentů (multiagentní systémy). Distribuovaná umělá inteligence tak využívá různé druhy (reaktivní, uvažující, sociální) reaktivně jednoduchých agentů za účelem dosažení specifického chování celého systému či vyřešení určité úlohy.

1.3. Limity poznávání dané historickou zkušeností

Veškeré naše jednání je ovlivněno naší předchozí zkušeností (ať už přímou či nepřímou), a to platí i o přístupu k poznávání a veškeré lidské činnosti z ní plynoucí. Historické zkušenosti tak ovlivňují i naše budoucí poznávání a přístup k němu. Celé generace vytvářely cestu lemovanou úspěchy i neúspěchy, na kterou nyní navazujeme i my. Vytvořili jsme vlastní konstrukty, jež více či méně úspěšně odrážejí fungování světa kolem nás. V tomto kontextu mluvím především o matematice⁴⁰ a vědách, které ji do značné míry využívají (fyzika, chemie, informatika. . .). Ačkoli to na první pohled vypadá, že vše funguje dle přírodních zákonů, které povětšinou známe a dokážeme do nich pomocí matematického aparátu proniknout až k podstatě věci, tak tento náš optimizmus není zcela na místě. I když nám moderní matematika dává určitý vhled do skutečnosti, dokonce i možnost předvídání budoucnosti,⁴¹ přesto tento vhled není celistvý.

KURT GÖDEL v roce 1931 dokázal, že „...kromě velmi primitivních teorií (s konečným počtem možných tvrzení) žádný konečný systém axiomů nemůže být úplný.“ [Mareš, 2006, str. 183] Jinak řečeno, ať vybudujeme sebelepší axiomatický systém,

⁴⁰ . . . která má výsadní postavení i v přístupech k umělé inteligenci. . .

⁴¹ Dokážeme vypočítat, jak se bude systém chovat v realitě – známe výsledek dříve, než nastane.

přesto nebudeme schopni vždy odvodit správnost či nesprávnost každého výroku v rámci dané teorie. Ještě odvážněji můžeme říci, že matematika nemůže obsáhnout (redukovat) obrovský⁴² svět kolem nás, což se odráží i v možnostech poznávání, kdy mnohé naše poznání stavíme právě na pojmech, které se vytvářejí z pojmů jednodušších za pomoci definic. V lepším případě nebudou dosažené výsledky plně odpovídat odrazu v realitě, což rychle odhalíme, v horším případě se vydáme slepou cestou apriorního poznání (nezávislého na zkušenosti). V současné době nám v tradiční umělé inteligenci postačuje formální matematický přístup, protože její složitost není tak velká a oblast, kterou řeší, je finitní.⁴³ Ovšem co se stane, když budeme chtít vytvořit složitou umělou inteligenci, která bude na podobné úrovni jako naše mysl, totiž bez jasného ohraničení (ať už v rámci prostředí či činnosti), či lépe řečeno - ohraničená světem kolem nás, tak jak jej vnímáme i my? Či spíše nejen, jak svět vnímá člověk s omezenými přímými poznávacími možnostmi, ale ještě podrobněji, tak, jak nám to nejnovější technologie nabízejí. Narazí formální systémy na své limity, kdy nebudou schopny dostatečně popsat svět kolem sebe? Pokud by tomu tak bylo, tak můžeme říci, že nikdy nebudeme schopni – ani s pomocí současných strojů – razantně prohloubit komplexnost poznání našeho světa. Otázkou také je, kde budou ony hranice poznání současných přístupů k umělé inteligenci založených na formálních systémech? Může to tedy znamenat, že formální matematický přístup je pro možnosti poznávání umělých systémů (potažmo i nás) do budoucna omezující. Smíříme se s tím, nebo budeme hledat novou cestu? Vypadá to, že určitý nový směr přináší principiálně tzv. nová umělá inteligence. Myšlenkově je to nová cesta, kdy budujeme umělou inteligenci stylem zdola nahoru na základě jednoduchých principů reaktivity. [Kelemen, 1994, str. 63] Přesto ovšem i u tohoto přístupu využíváme možnosti informatiky (softwarová řešení) či robotiky (hardwarová řešení), které stavějí své základy na formálních systémech.

Navíc, matematika je pevně svázána s naší myslí (naše mysl ji vytvořila), s naším chápáním světa kolem nás, a je třeba podotknout, že i naše možnosti, tedy možnosti naší mysli, jsou konečné.⁴⁴ Je samozřejmě obtížné si představit, že by lidstvo přestalo využívat matematiku, když ve většině každodenních činností člověka nám

⁴²Dost možná nekonečný svět kolem nás. Nicméně, tomuto označení jsem se chtěl raději vyhnout.

⁴³Můžeme také říci ohraničená. Zahrnuje např. stroje specializované na určitou činnost (inteligentní pračka) nebo pracující v ohraničeném prostředí (autopilot v letadle).

⁴⁴Například možnostmi našeho těla a jeho uchopením světa, daným lidskou fylogenezí. Ona už jenom pouhá myšlenka na to, že naše vágní mysl vytváří exaktní disciplínu, jakou je matematika, načež se pomocí matematiky snaží přiblížit lidské mysli, stojí určitě za zamyšlením. Více v knize [Křemen, 2007], která se zabývá současnými přístupy k umělé inteligenci a zaváděním vágnosti.

slouží dobře. Matematika má výjimečnou pozici, avšak nabízí jednostranný pohled na svět. Popisujeme pomocí ní fungování světa, ovšem je třeba mít na zřeteli limity a problémy, které s sebou nese. Novověká věda, která na ní staví, je dnes přijímána až v dogmatickém smyslu obdobně jako církevní středověký výklad světa. A podobný je i její vztah vůči společnosti, kdy se na vědu obracíme stejně, jako se dříve lidé obraceli k církvi. [Kelemen, 1998, str. 58]

Nejlépe je to vidět ve fyzice, kde „fyzikové užívají matematiku jako svůj jazyk, tedy jako prostředek popisu.“ [Bais, 2009, str. 6] Postupem času se poznatky fyziky stále vzdalovaly naší přímé zkušenosti či přímé praktické ověřitelnosti různých teorií (např. teorie strun). Fyzika vytvořila složitý model světa, využívá k tomu matematické rovnice a různé konstanty, s jejichž změnou se mění i tento model světa. „V současné době ani nevíme, proč hodnoty univerzálních konstant jsou právě takové, jaké jsou.“ [Bais, 2009, str. 9][Smolin, 2009, str. 37] A doufáme, že hlubší teorie nám jednou přinese patřičné odpovědi. Stavěla se pyramida vědění, ovšem nikdo si nebyl jist, zda se staví na pevných základech, anebo zda stavíme na chybných teoriích. Rozhodčím zde byla matematika, jež nabízela aparát, který zaručoval určitou správnost pro další postup. Právě díky tomuto rozhodčímu se akcelerovala i rychlost rozvoje exaktních věd, leč v blízké době může mít opačný efekt.

Fyzika je typickou ukázkou konstruktů lidské mysli, kdy vytvořila model, který odráží, jak člověk vnímá svět (přírodní jevy a zákonitosti) kolem sebe. Fyzika se však postupně dostává do slepých uliček a její rozvoj se od 70. let 20. století rapidně zpomalil. [Smolin, 2009, str. 13] „Oba tyto objevy, relativity i kvantové teorie, znamenaly definitivní rozchod s newtonovskou fyzikou. Přes stoleté úsilí jsou však dosud neúplné. Oba ukazují na potřebu ještě hlubší teorie. Hlavním problémem obou teorií je totiž existence té druhé.“ [Smolin, 2009, str. 30] Vznikl tak mýtus o sjednocující obecné teorii, kdy se vědci předhánějí, kdo „objeví“ tu nejmenší částici, ze které je vše postaveno. [Lukacs, 2009, str. 78] Možná zde je třeba radikální změna přístupu (myšlení), o kterém sami fyzici začínají uvažovat.⁴⁵ Na současném pohledu se podílela významně i matematika, o kterou se fyzika opírá. Ta funguje jako zrcadlo našeho pohledu na svět a teoriemi, které v ní nemají odraz (oporu), se dále nezabýváme. „Již se nemusíme spoléhat na experiment, který by teorie ověřoval. Tak to dělal GALILEO. Teď stačí pouhá matematika, aby zkoumala zákony přírody. Vstoupili jsme do éry postmoderní fyziky.“ [Smolin, 2009, str. 129]

⁴⁵Viz [Smolin, 2009, str. 34]

V této ukázce na příkladu fyziky nebylo mým cílem říci, že za problémy současné fyziky stojí matematika – to by bylo příliš opovázlivé. Prvním cílem bylo poukázat na problémy, které pronásledují veškeré přírodní vědy, jež staví na tzv. „tvrdých teoriích“. „Dokud se nástroje poskytované paradigmatickým osvědčují při řešení problému tímto paradigmatickým vymezením, pohybuje se věda kupředu [...] Význam krizí spočívá v tom, že poskytují náznaky toho, že příležitost pro změnu nástrojů právě přišla.“ [Iser, 2009, str. 20, původní citace Thomase Kuhna] Druhým cílem bylo i zde poukázat na důležitou, nikoli nutně zápornou roli matematiky. Záleží jen na našem pohledu a tom, co chceme dokázat.

Problém nekritického přístupu (pohledu) vidím v dnešním výkladu prováděném především na primárních vzdělávacích zařízeních,⁴⁶ případně ve způsobu přijímání a vnitřní interpretace výkladu žákem či studentem. Zde se mnohdy axiomatický systém výstavby vysvětluje nebo přijímá dogmaticky jako skutečnost, fakt, pravda (jediné správné vyjádření reality, se kterým můžeme pracovat). Ještě jinak řečeno: „Nejsme zatěžováni komplikovanější pravdou.“ Je to problém interpretace všech přírodních věd, které se vyučují na školách, kdy zaměňujeme naší myslí vytvořený model se skutečností.⁴⁷ To se odráží pak v našem nekritickém pohledu, ohraničené tvořivosti a omezené perspektivě vidění světa. Lidská inteligence vzniká až díky společnosti, je s ní úzce spjata včetně sociálního a kulturního zázemí. „Intelekt člověka sa utvára ako emergentný efekt jeho bytia a konania v spoločenstve iných ľudí.“ [Kelemen, 1998, str. 63-64] Než společnost vytvoří vědce, vystaví ho svému vlivu, své „pravdě“ o světě, a potom i vědec pouze navazuje na vyšlapanou cestu vědění.

1.4. Možné budoucí směřování umělé inteligence, respektive umělého života

Obecně tedy není vždy prospěšné zůstávat svázán současnými přístupy či postupy, které mají své kořeny v hlubší historii.⁴⁸ Je třeba se i k problémům poznávání v umělé inteligenci postavit otevřeně, nechat se inspirovat světem kolem nás, a to v transdis-

⁴⁶Právě zde se vytváří pohled jedince na svět. Jedinec se utváří dle potřeb společnosti, aby všichni měli pokud možno stejné vidění světa a byli k sobě intencionálně blízko.

⁴⁷Kolik učitelů předkládá např. problematiku orbitalů v chemii jako určitý model a kolik ji interpretuje jako skutečnost – takto to prostě funguje. Více nemá cenu studenty zatěžovat. Realitu však modelem pouze naznačím, nikoli plně vyjádřím.

⁴⁸V tomto smyslu je matematika nejzákladnějším pilířem. V problematice umělé inteligence nalezneme určitě i další části, na kterých „stavíme“ a mohou být omezující, ovšem ty jsou samy o sobě postaveny na matematice. Kupříkladu, problém umělého těla – klasický počítač vs. kvantový počítač.

ciplinárním smyslu. Je důležité nespolehat pouze na možnosti, které přináší současná informatika, respektive přírodní vědy, nýbrž využít celé širší vědních oborů (např. biologie, psychologie, lékařské vědy, filosofie), jejichž procesy jsou někdy těžko formalizovatelné pomocí matematiky.⁴⁹ A nejenom to, další výzvou je neřídit se přímo zažitými a někdy sdílenými schématy, jež jsou postupem času redukční. Takovými schématy může být v současnosti například komputacionistický funkcionalismus nebo neurologický pohled na mozek, který je připodobňován stroji. Profesor neurologie OLIVER SACKS upozorňuje na tento problém, kdy klasická neurologie těžko vysvětluje některé duševní pochody, a dodává k tomu: „Klasická neurologie byla vždycky (stejně jako fyzika) mechanická: vychází z Jacksonovy (myšlen John Hughlings Jackson – pozn. Z.S.) analogie mozek-stroj, kterou dnes změnila na analogii mozek-počítač. [...] Neskládají (duševní pochody – pozn. Z.S.) se pouze z abstrakce, klasifikace a kategorizace, jsou v nich i city, pocity a soudy. Pokud by nám chyběly, změnili bychom se v počítače...“ [Sacks, 2008, str. 32] V mozku se nepřenáší jenom vzruchy, nýbrž jsou zde i mnohé chemické reakce, které počítačí stroje, tak jak je známe dnes, nemohou nahradit a těžko je napodobují. V tomto ohledu neznáme všechny konsekvence biochemických procesů v mozku a jejich efekty na vědomí.⁵⁰ Samozřejmě, že v tomto smyslu hovořím o směrování k silné umělé inteligenci.

Dle mého názoru může mít transdisciplinární přístup pozitivní vliv na tvorbu budoucích inteligentních systémů i na přístupy k jejich poznávacím mechanismům. Nakročeno již máme (multiagentní systémy, genetické algoritmy...). Jak by ale měl tento směr vypadat do budoucna? Můžeme navázat na myšlenku „nové umělé inteligence“ a prozkoumávat příčiny a zákonitosti v živé přírodě od zdola nahoru. Poznáním na pomezí toho, čemu říkáme život, můžeme vetknout do současných přístupů k umělé inteligenci nové směry. Tímto se dostáváme na půdu umělého života a jeho zkoumání. Takové výzkumy by se měly zaměřit především na viry a jednoduché bakterie (či spíše jejich organely). A nejenom z teoretické stránky jako doposud, výhodné by bylo s touto hmotou reálně pracovat, neboť tyto systémy nejsou ještě tak komplikované a jejich biologické procesy nejsou příliš časově náročné.⁵¹ Kupříkladu v 1 ml kultivačního media může žít až 10^9 bakterií, přičemž „za optimálních podmínek by za 48 h při generační době 20 min z jediné buňky vzniklo 2^{144} buněk, což odpovídá 4000

⁴⁹Ztratíme tak částečně kontrolu nad umělým systémem, spíše budeme pracovat s „black boxy“, kde víme, jaký máme dát vstup, aby se komplexní systém s určitou pravděpodobností nějakým způsobem zachoval. Viz dále.

⁵⁰Základní postřehy o funkci biochemických procesů v mozku lze najít v [Pstružina, 1998, str. 165]

⁵¹Než se vytvoří složitější organizmus, tak to trvá delší dobu.

násobku hmoty zeměkoule.“ [Hodek, Šulc, 2009] Výhodou je i ono pomezí mezi tím, co označujeme jako živé a co je z našeho pohledu neživé, protože to má blízko k naší schopnosti vytvářet z neživé přírody (výpočetní) stroje. Nevýhodou je finančně velmi nákladný laboratorní výzkum v oblasti biochemie, i když by se zpočátku mohly využít již hotové výzkumy a možnosti počítačového modelování. Neměla by se striktně razit cesta směrem k práci buď s živou, nebo neživou hmotou, ale najít kompromis, využít předností obou a získat tak vhodnou efektivitu celku.

V souvislosti s touto myšlenkou mi byla položena otázka, která je nasnadě: „Není to pak spíše práce s black boxy, kde nevíme, jak se přesně zachovají? I tato relativně primitivní hmota (ve srovnání kupříkladu s námi) je velmi komplikovaná a její chování v současné době do značné míry nepředpokladatelné. Jak potom nasadit nějaký takový inteligentní systém do rizikových oblastí, kde je třeba přesné a jasné rozhodnutí, na které se dá spolehnout?“ Je třeba si ale položit otázku, zda i dnešní systémy, ať už softwarové nebo hardwarové, nejsou do jisté míry takovými „black boxy“. Představím zde dva příklady:

1. Pokud máme velmi složitý program, řekněme operační systém nebo navigační systém, který má statisíce až milióny řádků kódu kupříkladu v jazyku C či C++, a snažíme jsme se jej sebelépe napsat, tak po jeho spuštění čas od času za neopakovatelných podmínek dojde k chybě.⁵² Ta vzniká kvůli velké složitosti programu a je velmi těžko odstranitelná, není to tedy klasický druh chyby. I když provedeme stejný postup, nemusí se chyba objevit. Nejsou to tedy běžné chyby, jejichž odstranění je triviální. Cílem vývojového týmu je tyto projevy minimalizovat.
2. Lidé pracující pro letecký či kosmický průmysl si velmi brzy začali uvědomovat, že na první pohled tvrdé systémy se mohou chovat jako měkké. Především v souvislosti s kritickými hardwarovými či softwarovými systémy se proto hovoří o tzv. dependabilitě. Cílem dependability je snížit možnost nepředpokládaného (chybového) chování systému na minimum a poskytnout tak službu, na kterou se dá spolehnout. Nejjednodušším způsobem je například redundance⁵³ (nadbytečnost), kdy se současně několikrát provede nějaký postup, ovšem vždy s jiným druhem algoritmu nebo diverzifikací použitého programovacího jazyka.

⁵²Respektive k nežádoucímu chování, které jsme nepředpokládali.

⁵³Původně se redundance zaváděla v oblasti hardwaru od 50. let 20. století kvůli konstrukčním problémům hardwaru a nepřiliš vyspělým technologiím výroby. Navíc, vyrobit několik jednodušších a poruchovějších zařízení bylo levnější než vyrobit jedno složitější, které by zajišťovalo stejnou spolehlivost.

Dle mého názoru každý složitější systém, který člověk vytvořil pomocí přesných exaktních postupů, je svým chováním do určité míry nepředvídatelný podobně jako „black box“. Přesto tyto systémy používáme, neboť víme, že v naprosté většině případů se zachovávají očekávaně. U kritických systémů s tím bojujeme zvyšováním spolehlivosti systémů (HW i SW), nebo například už při vývoji speciálními softwarovými nástroji pro správu komplexního kódu (C++, C# či JAVA), kdy takový nástroj upozorňuje na různé vazby mezi třídami a funkcemi.⁵⁴

Je to obdobné, jako když se snažíme prohlédnout emergentní jevy v přírodě díky systematickému vhlížení do pozorovaného systému. „Emergentní efekt je tedy zejména projevem dynamiky nikdy nekončícího výkladového procesu nahlížení světa...“ [Mařík a kol., 2007, str. 110] Záleží ale jen na nás, co chceme pozorovat a jaké máme možnosti pozorování, a to ať už vědecké, kulturní či sociální. Knihu přírody tak čteme po částech, což závisí na naší (technické) schopnosti nahlížet pouze na části světa, nikoli celek (ten nejsme schopni obsáhnout). V poznávání emergentních jevů přírody (světa kolem nás) můžeme nalézt určitou paralelu se systémy, které vytvořil člověk a u kterých tedy předpokládáme, že přesně víme, jak se bude systém chovat. Když totiž vytváříme velmi složitý systém, jež nejsme příliš schopni prohlédnout celý sami (mozek jedince),⁵⁵ tak nám některé spojitosti uniknou. Tyto spojitosti pak vytvářejí neočekávané chování. Člověk totiž z vývojového hlediska nebyl až do nedávna zvyklý vytvářet tak složité stroje, respektive využívat složité konstrukty své mysli k dosažení svých cílů. Daný „black box“ tedy vzniká především lidským omezením, které se snažíme eliminovat různými způsoby.

Z těchto důvodů si myslím, že pokud rozšíříme naše poznání v oblasti primitivního života (pro začátek), tak se dostaneme v předpověditelnosti chování (a tedy i možnostmi práce s primitivním životem) na úroveň člověkem vytvořených umělých (počítačových) systémů. To vše dává předpoklady k budoucí hybridní symbióze „živé“ a „neživé“ složky umělé inteligence, či lépe řečeno, umělého života.

Nesmíme zapomenout ani na pohled shora dolů, typický pro tzv. tradiční umělou inteligenci. Lidé chtějí aplikovatelné a použitelné výsledky hned, což jim touto formou můžeme nabídnout. Kromě klasického redukcionismu, ve smyslu rozložení problému na části, by mělo být cílem také napodobit určité lidské chování ve smyslu komuni-

⁵⁴Příkladem může být např. NDepend, www.ndepend.com Využívá se při analýze kódu, která se dále může využít pro redaktorování různých segmentů kódu.

⁵⁵Nedokážeme prohlédnout jeho části tak, abychom byli schopni předpovědět veškeré chování celku.

kace.⁵⁶ Kromě řešení specifických úloh by měla umělá inteligence zajišťovat interakce na vyšší úrovni se samotným člověkem. Umělá inteligence má člověku především pomáhat a rozšířit jeho schopnosti i možnosti. Proto by se i nadále měly rozvíjet přístupy k umělé inteligenci, které se snaží napodobit chování a rozhodování člověka v určité situaci. Možný je i hybrid tradiční a nové umělé inteligence, kdy reprezentace světa bude uchována v komplexnosti systému a interakce mezi prvky budou generovat (emergentně) rozhodnutí celku.

1.5. Závěrečná implikace

Celkový pohled na umělou inteligenci bych zakončil určitou analogií s výše zmíněnými příklady vztahující se k fyzice jako vědnímu oboru. Náš prvotní pohled na umělou inteligenci byl shora dolů – snažíme se rovnou napodobit inteligentní chování člověka (tzv. tradiční umělá inteligence) stejně, jako když NEWTON na základě svého přímého pozorování vytvářel obecné zákony přírody. Postupem času s lepšícimi se schopnostmi člověka v poznávání světa kolem sebe se fyzika a chemie zaměřila na svět i mimo naši přímou perцепci. Objevil se zcela nový mikrosvět, u něhož pozorujeme jiné chování než u větších ekvivalentů přímé zkušenosti našeho světa. Stejně i tzv. nová umělá inteligence se zaměřila na postup zdola nahoru a přišla s myšlenkami, že i jednoduché reakce mohou vytvářet složité chování. Taktéž s trochou nadsázky může kvantová fyzika vysvětlit náš svět a jeho komplexnost pomocí šesti druhů kvarků, šesti druhů leptonů a čtyř interakcí. [Smolin, 2009, str. 37]

Problémem obou pohledů je, jak je spojit dohromady. Fyzika nedokáže přesně popsat, jak vzniká náš makrosvět z prvků objevených v mikrosvětě, podobně nová umělá inteligence neví, jak z jednoduchých částí emerguje mysl. Pokud by se nám ono spojení povedlo ve fyzice, tak dokážeme o dost lépe vytvářet závěry z pozorování světa kolem nás.⁵⁷ Jestliže se nám podaří sjednotit přístupy k umělé inteligenci, tak dokážeme vytvářet umělou inteligenci různé úrovně přesně podle toho, jak budeme jako lidstvo potřebovat. Takové inteligentní systémy se budou moci lépe zařadit do našeho kaž-

⁵⁶Zaměřit se na chování celku tak, aby připadal člověku inteligentní. Nestačí jen, aby umělá inteligence řešila nějaký problém, cílem by měla být i komfortní interakce „inteligentního celku“ s člověkem.

⁵⁷V poslední době se stále více mluví o teorii dekoherence, která poukazuje na to, že naším skutečným problémem je samo pozorování, respektive nechtěné interakce pozorovaného objektu s prostředím, jež zkresluje pozorovaný výsledek. Viz [Pavlík, 2004, str. 166].

dodenního života. A díky spojení živé a neživé hmoty můžeme vytvořit umělý život, který bude z našeho pohledu téměř nesmrtelný.⁵⁸

Konečným cílem by pak mohlo být vytvoření inteligentního umělého života, což by mohlo znamenat vytvoření oné silné umělé inteligence. V oblasti slabé umělé inteligence jsou současné tendence ve vývoji dostatečné, neboť tyto systémy budou mít omezenou autonomii a budou přímo či nepřímo ovládány (omezovány) člověkem. V tomto ohledu jsou pro nás zajímavé stále tradiční přístupy k umělé inteligenci i kvůli tomu, že přinášejí v současné době největší míru uplatnění v praxi.

Netvrdím, že výše zmíněný „nový“ směr výstavby zdola nahoru sdílený umělou inteligencí a umělým životem a pracující s živým materiálem je jediný možný, spíše nám může otevřít nové kvality, které jsme doposud nebyli schopni vidět. Je to další z mnoha cest, která může otevřít další cesty, včetně té k sobě samému. Při těchto myšlenkách zastávám názor, že účelná je zde pluralita názorů a především pohledů na problematiku (se zdravou dávkou kritičnosti). Můžeme tedy říci, že na „rozdíl od moderního universalismu tu máme postmoderní pluralismus [...] rozumnost se nenachází v homologii expertů, nýbrž v paralogii vynalézajících.“ [Peregrin, 2008, str. 47 a 49] Na to jsem se příklady snažil poukázat i v částech věnované abstraktním konstruktům,⁵⁹ které vytváří iluzi jediného rozhodce, kterou cestou se máme vydat. Výsledkem by měla být zdravá rozmanitost pohledů, což je předpoklad k neustálému vývoji v dané oblasti. Bohužel, jak bylo řečeno, tato rozmanitost pohledů má svá omezení a naráží na výchovu a jednotné vzdělání jako výchozí podmínky.⁶⁰

⁵⁸Tělo takto vytvořené by nemuselo tak jako u živočichů fungovat pouze desítky let, ale díky našim vlastním zásahům i několikanásobně déle, než se nyní dožíváme.

⁵⁹Například zmíněná matematika.

⁶⁰Zajímavou postavou fyziky v tomto světle je RICHARD FEYNMAN (1918 - 1988), kromě jiného i držitel Nobelovy ceny za fyziku. Byl známý svým osobitým přístupem k vysvětlované látce. Legendární se stala kniha „Feynmanovy přednášky z fyziky“ (dále jen „Přednášky“), jež osobitým způsobem rozebírá obtížné kapitoly ve fyzice, a kniha „Šest snadných kapitol“, která je sestavena z vybraných úvodních kapitol k jeho přednáškám. V těchto kapitolách se snažil využívat matematiku minimálně a spíše dát prostor experimentům a jednoduché představě. Jeho pohled na fyziku byl neklasický, a proto mnohdy „viděl“ ve fyzice věci, které druzí neviděli.

Ve Feynmanově úvodu, který je zahrnut jak do „Přednášek“, tak do „Šesti snadných kapitol“, říká např. k výuce kvantové mechaniky: „... běžný způsob, jakým je kvantová mechanika vykládána, ji činí pro většinu z nich (studentů, pozn. autora) nedostupnou, protože absolvovat kurs kvantové mechaniky trvá příliš dlouho. A přitom při skutečných aplikacích – zvláště těch složitějších v elektronice a chemii – není složité řešení diferenciálních rovnic skutečně používáno. Takže jsem se snažil vysvětlit používání kvantové mechaniky, aniž by k tomu bylo nejprve třeba zvládnout řešení parciálních diferenciálních rovnic.“ [Feynman, 2007, str. 23]

Také uvádí jeden zajímavý postřeh k přednáškám na Caltechu, které vedl v letech 1961 a 1962 a z nichž byly poté sestaveny knižní „Přednášky“: „Každopádně nebylo mým úmyslem ztratit zájem byt i jediného studenta, jak se nejspíš stalo. Myslím, že jedním ze způsobů, jak se toho vyvarovat, by bylo investovat víc práce do vytvoření sady problémů, které by vyjasnily některé z myšlenek uváděných na přednáškách. Řešení problémů poskytuje dobrou příležitost doplnit látku uváděnou na přednáškách a umožňuje, aby se abstraktní teorie zkonkretizovala a lépe pak utkvěla v myslích posluchačů. Přes to všechno si myslím, že problém vzdělávání studentů má jen jediné řešení – uvě-

Již nyní zde existují první pokusy, kupříkladu pokusy s biologickým materiálem, které se provádějí a jsou známé jako bakteriální či DNA počítače, které řeší některé jednoduché úkoly, např. problém připálených palačinek. [Haynes a kol., 2008] V tomto případě se jedná spíše o určité „ohýbání“ biologických možností bakterií po vzoru softwarových programů založených na genetických algoritmech.⁶¹ Je však třeba se zamyslet, zda toto chování je skutečně tím, co pak běžně v přírodě pozorujeme, anebo jsme jen objevili možnost, jak mohou tyto jednoduché buňky kromě jiného fungovat. Vždy tedy záleží na tom, jaký je cíl našich pokusů. V 90. letech 20. století vědci dokázali, že soubory DNA molekul dokáží počítat obdobně jako Turingův stroj. [Kelemen, 2010, str. 213] Ovšem vypadá to spíše jako hledání analogií se současnými (neživými) systémy. Je třeba udělat krok vpřed, zbavit se tíže neživého a ponořit se plně do prostředí, kde tyto živé systémy rozvinou své kvality. Takovým může být i lidské tělo. Na druhé straně je pravdou, že člověk potřebuje nalézt něco známého při svém poznávání, paralely, aby měl odkud se odrazit a začít tvůrčí činnost. [Pstružina, 2005, str. 61]

Bezesporu to vede k většímu prozkoumání možností DNA, včetně různých organel a jejich úloh v biologických procesech. Spíše, než napodobovat chování složitých procesů (např. softwarové multiagentní systémy), se nabízí možnost pracovat přímo s živou hmotou. Právě v tomto spatřuji jednu z nových cest, kterou se může umělá inteligence vydat, směrem k umělým biologickým či semibiologickým systémům, jak bylo zmíněno. Znamená to i lepší objasnění (lidských) kognitivních procesů, které můžeme využít dále v umělých systémech. Oklikou bychom se pak mohli dostat až k tématu fungování neuronů, respektive mozku, který nefunguje jako dnešní počítačí stroje a v nich algoritmičké programy – je z jiného materiálu, jinak uspořádán, je přizpůsoben určitému tělu a reaguje v jiném prostředí. V tomto spatřuji oproštění se od pohledů, které jsou dnes upřednostňovány a ze kterých se v základu vychází.

Zajímavé pokusy v tomto ohledu, i když velmi odvážné,⁶² jsou s tzv. „biomozkem“, prováděné na University of Reading. „Biologický mozek se skládá ze souboru neuronů, vypěstovaných na multielektrodovém poli (Multi-Electrode Array, MEA). MEA je vypouklá rovina, zásobená 60 elektrodami, a ty přijímají elektrické signály generované

domit si, že nejlepších výsledků lze dosáhnout pouze přímým individuálním stykem mezi studentem a dobrým učitelem, s nímž může student všechny myšlenky rozebírat a prohovorit.“ [Feynman, 2007, str. 24]

⁶¹Genetické algoritmy se inspirovaly přírodou, a tudíž ji jen do určité míry napodobují. Jedná se o heuristický způsob řešení složitějších problémů. Genetické algoritmy spadají do tzv. evolučních výpočetních technik. Viz [Mařík a kol., 2001, str. 117]

⁶²Využívají se již celkem složité biologické systémy.

buňkami. Tyto signály pak řídí pohyby robota. "Biomozek" spolupracuje s robotem pomocí Bluetooth.“ [Ideje, 2009] Zajímavostí je, že na výzkumu se podílí profesor KEVIN WARWICK, u nás známý především díky knize *Úsvit robotů – soumrak lidstva* (Vesmír, Praha 1999). Tyto pokusy, dle mého názoru, značně přispějí ke zmapování kognitivních možností biologických organizmů, a rozšíří tak i možnosti jejich aplikací v umělých inteligentních systémech (především směrem k silné umělé inteligenci, respektive inteligentnímu umělému životu).

Co se kognitivních schopností týče, tak nové umělé inteligenci (umělému životu) by velmi pomohl výzkum na rozhraní toho, čemu říkáme život. Nejenom, že bychom zjistili nové způsoby, jak získat určité vjemy z prostředí a dále je zpracovat, ale „vedlejší“ výsledky by mohly pomoci lidem i v ostatních vědeckých oblastech. Příkladem a první vlašťovkou mohou být výše zmíněné DNA počítače. EHUD SHAPIRO se svým týmem v roce 2002 „vedli programovatelný molekulární počítač složený z enzymů a DNA molekul místo křemíkového čipu.“ [Lovgren, 2003] V našem dualistickém pohledu pak DNA představuje software a enzymy hardware. O dva roky později oznámili, že dokážou zjistit nebezpečnou aktivitu uvnitř buňky směřující k rakovinovému bujení a následně na to i účinně zasáhnout. [Benenson a kol., 2004] To již klade na DNA počítač nárok na schopnost pracovat ve specifickém „živém“ prostředí, zjišťovat a zpracovávat údaje a reagovat – pracovat tedy se vstupy a posléze i výstupy.

Příklady uvedenými na konci této kapitoly jsem chtěl podpořit své závěry ohledně budoucího vývoje umělé inteligence, umělého života či nový biologicky zaměřený přístup natural computing [Kelemen, 2010, str. 211]. S budoucím vývojem jsou spojeny i pohledy na kognitivní procesy odehrávající se uvnitř umělých systémů, přičemž jsou často srovnávány s člověkem – v této problematice jsem se omezil spíše na nastínění problémů a současných názorů na řešení s ohledem na historické souvislosti. V dalších kapitolách této práce budu pracovat především se současnými přístupy ke kognitivním možnostem umělé inteligence.

2. Kognitivní procesy člověka a umělé inteligence

Kognitivní procesy u člověka a umělé inteligence stojí v popředí zájmu kognitivní vědy, respektive kognitivní informatiky.¹ Jak již bylo naznačeno v předchozí kapitole, cílem umělé inteligence by mělo být pomáhat člověku v jeho různých činnostech, čehož dosahujeme – kromě jiného – i studiem toho, co považujeme za inteligentní chování člověka. Z tohoto důvodu musí inteligentní systémy (především tradiční umělé inteligence) také na patřičné úrovni s člověkem interagovat či vhodným způsobem komunikovat. A právě studium poznávacích procesů člověka a jejich (i částečná) simulace v oblasti umělé inteligence by nám mohla zajistit toto přiblížení. Jinak řečeno, tímto můžeme alespoň částečně spojit náš svět² se „světem strojů“ již nyní.³

V tomto kontextu si dovoluji malou metaforu z oblasti zábavního průmyslu. Ať už se natočí jakýkoli sci-fi či fantasy film, vždy je nutné, aby zde byly určité znaky podobnosti s naším světem a zákony. Vždy tam musíme najít něco známého, i když je to v jiné podobě (antropoidní mimozemšťan, fyzikální principy, zvířeti podobná stvoření). Náš mozek v tom najde příběh, do kterého se dokážeme vcítit. Pokud bychom vytvořili příběh na principech nám vzdáleného světa, těžko bychom v něm hledali správné paralely s naším světem, které máme ve svém modelu světa – výsledkem by bylo odcizení a nepochopení. Člověk a film tedy musí být vytvořeni jeden pro druhého. To samé platí například i o virtuálních online světech, kde je třeba zachovat určitý odraz našeho reálného a známého světa. Mozek jako hlavní centrum kognitivní činnosti nejen, že v sobě zrcadlí naše tělo, jež je mozkiem ovládáno, nýbrž také díky ontogenezi jedince zrcadlí i model prostředí, ve kterém existuje. Výsledkem je

¹Ačkoli není chybou, když zahrneme kognitivní procesy do prostoru zájmu kognitivní vědy, přesto nalezneme ještě užší okruh, do kterého by tato problematika měla spadat. Jedná se o kognitivní informatiku, kterou se snaží ve světě propagovat a také jako první definoval okruh zájmu prof. Yingxu Wang z University of Calgary. „Cognitive Informatics is a transdisciplinary enquiry on the internal information processing mechanisms and processes of the natural intelligence – human brains and minds – and their engineering applications in cognitive computing, computational intelligence, as well as the information/communication technology and software industries.“ [ICCI 2010]

²„Svět vždy představuje vnitřní řád (kosmos), který jej tvoří určitým ve smyslu samostatnosti a nezávislosti.“ [Pstružina, 1995] Svět vyjadřuje ohraničené prostředí, ze kterého jsme schopni přijímat podněty a dále je zpracovávat. Jedná se tedy i o prostředí, které je mimo naši přímou percepci. To dále v textu označuji jako svět, zatímco prostředím myslím část světa v dosahu našeho přímého či nepřímého vnímání. Viz „Prostředí člověka a umělé inteligence“

³I když jde zatím pouze o „slabou“ umělou inteligenci, která se pouze chová inteligentně.

naš svět a jeho určité chápání. A právě v tomto kontextu (na těchto principech) by měly pracovat i inteligentní systémy vycházející z tradiční umělé inteligence.

Z těchto důvodů by inteligentní systémy měly umět reagovat v rámci našeho světa předvídatelně – známě, racionálně, nikoli však nutně tvořivě. Jsme to právě my, lidé, kdo máme s inteligentními stroji, které vytváříme, i žít. Samotné vytváření stroje člověkem dává předpoklady k tomu, že při vývoji v nich otiskneme svoje vidění světa a možnosti, které v něm máme. Vždyť naše okolí i my sami jsme výbornou inspirací při této tvorbě. Proto je důležité studovat lidské (obecně biologické) poznávací procesy a z nich vycházet při tvorbě poznávacích procesů umělé inteligence, které jim jsou vzdáleně podobné. Obdobně jako model reality, kterou představuje.

V následující kapitole proto vycházím z našeho známého světa, přičemž budu sledovat jeho průnik se světem strojů. Je nutné podotknout, že se budu zabývat vždy jen částí našeho světa, neboť ani my nežijeme v celém světě naráz, nýbrž v jeho části, která nás obklopuje. Navíc, trendem současné umělé inteligence je úzká specializace na konkrétní činnost, respektive prostředí.

Sbližování našeho světa a světa strojů probíhá oboustranně: Člověk proniká již od mala do tajů techniky, zařazuje ji do svého každodenního života, a naopak – technický pokrok, nutnost vývoje přibližuje stroje lidem. To vše podtrhuje důležitost kognitivních procesů ve spojitosti s umělou inteligencí na jakékoli úrovni.

Co se týká obsahu této kapitoly, určitě by bylo možné pojmut problematiku rozsáhleji z hlediska hlavní osy zájmu *prostředí – podněty – zpracování – odraz v kognitivním systému – reakce na podněty* v umělé inteligenci. Avšak cílem této práce není podat zcela vyčerpávající souhrn v oblasti kognitivních procesů v umělé inteligenci,⁴ nýbrž naznačit určité horizonty, obecný pohled na problematiku v naznačeném kontextu.⁵

2.1. Prostředí a podněty

Člověk, stejně jako reálný nebo virtuální stroj,⁶ existuje v určitém prostředí, které více či méně vnímá, a musí být schopen v něm interagovat. Celá biologická evoluce stejně jako ontologický vývoj jedince je zásadně determinována prostředím, respektive světem. U člověka s tím souvisí i jeho přežití, přičemž svět, ve kterém žije, je

⁴To by bylo nad rámec rozsahu této práce.

⁵Především vzhledem k člověku.

⁶Například soft robot, viz [Kelemen, 1994, str. 38]

velmi složitý například oproti světu housenky. Značný rozdíl nalezneme právě v podnětech, které dokáže člověk a housenka vnímat a následně zpracovat (v rámci kognitivního systému). S tím souvisí i různorodost a množství receptorů, které má housenka a jež má k dispozici člověk. Jednoduše lze říci, že housenka vnímá oproti člověku jiný, lépe řečeno menší část světa, který pro její existenci postačuje.⁷

Stejně jako na housenku lze pohlížet i na člověka, jehož vidění světa je nejsložitější, ale určitě lze vytvořit i komplexnější pohled na svět než ten, který nám zprostředkovávají současné smysly. V tomto kontextu bych rád představil praktický příklad, kdy se jednotlivci rozšířilo vnímání světa. Student medicíny STEPHEN D. na sobě experimentoval s drogami (kokain a amfetaminy). Jednoho dne se probudil z živého snu, kde byl psem obklopeným velmi zajímavými a bohatými pachy. Po probuzení nejen, že svět plný pachů byl stále kolem něj, ale dokonce se mu zlepšila interpretace barev a paměť vybavovat si předměty. Otevřely se před ním nové možnosti, jaké nebyl schopen dříve vnímat, a navíc tento smysl zastřel všechny ostatní. „Zjistil, že podle pachy rozeznává přátele i pacienty na klinice [...] každý z nich měl vlastní pachovou fyziognomii, pachovou osobnost, která byla daleko výraznější než tvář.“ [Sacks, 2008, str. 166] Časem se orientoval především za pomoci čichu, přičemž měl tendence si vše očichat a osahat. Dokonce i jeho vnímání světa se změnilo: „Svět se rozpadl do nespočetných konkrétních jednotlivostí, jejichž bezprostřední a naléhavé prožitky mě přemáhaly.“ Stephen byl svým založením spíš intelektuál se sklonem k reflexi a k abstrakci. Po své proměně ale zjišťoval, že uprostřed nových zkušeností mu myšlení a kategorizace připadaly obtížné a nereálné.“ [Sacks, 2008, str. 166] Nový svět ho trochu mátl. Po třech týdnech jeho schopnosti zmizely. Tyto odchylky byly poprvé popsány již před více než sto lety. V tomto případě se jedná o záchvaty (epilepsie) v zahnuté části mozkového závitu gyrus hippocampi zvaného uncus, který je fylogeneticky součástí rhinencephalonu – tzv. starého čichového mozku. Dochází pak k celkovému posílení čichu (hyperosmii). [Sacks, 2008, str. 165] Cílem této krátké odbočky bylo poukázat na to, jak je člověk omezen ve svém přímém vnímání prostředí kolem sebe.

Člověk si brzy uvědomil svá omezení v „ohraničeném světě“ a začal uměle rozšiřovat svoje možnosti v rámci prostředí. Jeho um a myšlení mu pomohly vytvořit nástroje a později stroje (viz první kapitola), které významně zvyšují možnosti interakce s prostředím. Kupříkladu dokážeme létat mimo naši planetu v nehostinném vesmíru. Po-

⁷Samozřejmě, že housenka používá jiné receptory i způsob jejich vyhodnocení, než člověk. Nelze tedy jednoduše „rozšířit housenku“ a tím získat pohled na svět takový, jako má člověk. Obecně – čím vzdálenější živočišný druh vůči člověku, tím odlišnější vidění světa. Savci budou mít pohled na svět bližší člověku než trepka velká. Každý druh má svůj vlastní ohraničený svět.

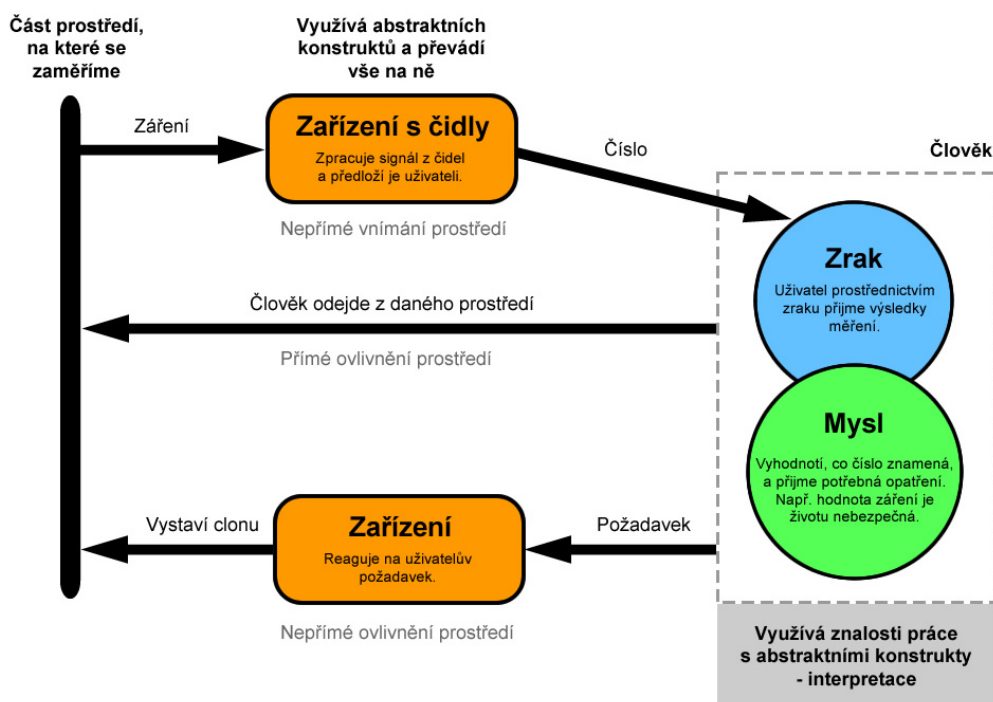
kud by hrozilo nějaké nebezpečí pro naši planetu, tak jej můžeme odhalit a zavčas na ně reagovat.⁸

Lidé nečekají na evoluci, co jim nabídne, ale sami ji vytvářejí svou činností. Takovou činností je i tvorba inteligentních systémů. Dosáhli jsme mnoha věcí, o kterých jsme dříve jako lidstvo pouze snili, a právě sny se staly důležitým motorem. Abychom však mohli stále lépe a lépe vnímat a pracovat s tímto světem, musíme co nejlépe pracovat v našem prostředí,⁹ ve kterém existujeme, a vytvořit nové nepřímé interakce s okolím. Možnosti přímého poznání, jež nabízí naše tělo, jsme téměř vyčerpali. Člověk se stále více obklopuje nejrůznějšími vynálezy, které převádějí pro člověka nevnímatelné signály z prostředí na zpracovatelné.¹⁰ Lidstvo se při (vědeckém) poznávání uzavírá do „skafandru“, přičemž většina nových údajů o světě jest zprostředkována. Vytváříme také nejrůznější zařízení a sama vytvořená zařízení převádějí údaje ze senzorů do řeči čísel a nejrůznějších abstraktních konstruktů, kterým rozumíme. Na straně příjemce je nutné tyto údaje dekodovat a zpracovat, je tedy nutná znalost například matematiky. Kupříkladu elektrokardiografie umožňuje neinvazivně snímat srdeční aktivitu. Získaná data specialista dále interpretuje nebo zpracuje pomocí fraktální analýzy s cílem zjistit, zda je člověk schopen (na základě odlišení fyziologických stavů) vykonávat určitou činnost. [Smrčka, 2002]

⁸To lidstvo bez svých vynálezů „in natura“ nedokáže.

⁹Nebo také pracovat s prostředím.

¹⁰Mezi „vynálezy“ bych zahrnul i abstraktní konstrukty, které nám umožňují vytvářet modely světa a díky nim objevovat nové spojitosti, které bychom jinak jen těžko odhalili.



Obrázek 2.1.: Nepřímá interakce člověka s prostředím.

Kromě zprostředkovaného poznávání především ve vědecké činnosti můžeme mluvit o větším kontextu dnešního člověka. Média dovolují přenášet informace nehledě na čas a místo (tzv. prostorové a časové odloučení). Jedince tak může ovlivnit i to, co se stalo na druhé straně planety, aniž by tam reálně byl. Pak lze tvrdit, že prostředí, ve kterém člověk žije a dokáže jej ať už přímo nebo nepřímo vnímat, je obrovské. Rapidní zvětšování kontextu člověka (skupiny lidí) lze sledovat již od vzniku masové komunikace (polovina 15. století) a v posledních letech ještě akceleruje díky informačním technologiím. Podněty, které tedy vyhodnocujeme každý den, pro nás nemusejí být přímo dostupné, ale zprostředkované, ať už se jedná o jakoukoliv úroveň. Oproti klasickému pohledu na život člověka to vypadá téměř jako virtuální či zprostředkovaný život.¹¹

Jednou z takových nepřímých (zprostředkovaných) interakcí pro člověka by mohla být i vyspělejší umělá inteligence, neboť i naše myšlení a schopnost zpracování množství informací je omezené, a tak by se umělá inteligence mohla stát dalším „prodlou-

¹¹Jsmo tak vystaveni obrovské možnosti manipulace. V dnešní době informačního věku nám masmédiá vytvářejí zprostředkovaný svět, jehož pravdivost nejsme schopni ověřit. Ve své krajnosti takový svět nemusí ani reálně existovat – protíná se tak s virtuálním. Navíc cílené manipulativní techniky dokonce ovlivňují naše reálné rozhodování – koho volit, komu dát peníze, jak žít, co je správné. Symbolická moc dnes vládne světu.

žením naší ruky“. Na jedné straně nám slabá umělá inteligence již dnes pomáhá předzpracovávat informace, abychom mohli abstrahovat na vyšší úroveň (přímá interakce s člověkem), na druhé straně by nám mohla silná umělá inteligence představit jiný pohled na svět, neboť její svět (prostředí a jeho vnímání) bude vypadat jinak než náš.¹²

Pokud se podíváme na horizont blízké budoucnosti umělé inteligence, je třeba brát v úvahu její pohled na svět, jež následně bude měnit i náš pohled a chápání zákonitostí světa, kde žijeme. Stále to však bude pouze zprostředkované a převedené na naše vlastní možnosti chápání. V dlouhodobějším horizontu směrem k silné umělé inteligenci však vyvstává problém, který dnes v opačné podobě nalezneme ve vztahu expertních systémů a lidského experta. V současnosti nedokážeme zcela zrcadlit zkušenosti experta do uměle vytvořeného systému, proto mají současné expertní systémy spíše poradní hlas. Je to způsobeno intencionalitou,¹³ kterou nejsme schopni přenést a kterou jazyk nedokáže vyjádřit. U umělých systémů budoucnosti narazíme na problém komplexity, kdy tyto systémy nám budou moci vysvětlit jednotlivosti či fungování, ovšem za tím vším bude stát komplexnější vidění, které nám nebudou schopné uspokojivě zprostředkovat.

Intelligentní systémy (využívající určitou aplikaci umělé inteligence) vytvořené člověkem nemusí – avšak mohou – být svázány s časem a místem. Variabilita v tomto ohledu je velká – dokonce čistě softwarová aplikace může být díky svým mnoha instancím kdykoli a kdekoli přítomna téměř okamžitě. Rychlost přenosu či reprodukce v kyberprostoru je oproti reálnému materiálnímu světu nesrovnatelná.¹⁴

V následujícím pojednání se budu zabývat srovnáním prostředí u člověka a stroje, jež může využívat některou z aplikací umělé inteligence. Dále mě zajímá, jaké vstupy využíváme a jaký je jejich význam z hlediska množství a složitosti na zpracování. Z toho vyplývá i to, jakým způsobem prostředí ovlivňuje daný umělý či přírodní systém.

2.1.1. Prostředí člověka a umělé inteligence

Prostředí chápu jako samostatný řád, který nás bezprostředně obklopuje. Je to určitý celek, jež má odraz i v našem kognitivním centru. V souvislosti s pojmem „prostředí“

¹²Tento rozdíl by šlo připodobnit setkání dvou vývojově rozdílných civilizací. Jinak chápe svět a pracuje se světem pravěký člověk a jinak dnešní Evropan vybavený mnoha elektronickými pomůckami.

¹³V souvislosti s tím se hovoří o tušení nebo šestém smyslu, který nelze racionálně zdůvodnit. Toto tušení sahá za hranice přímého vědomí.

¹⁴Lze si to představit na příkladu, kdy softwarovou aplikaci distribuji přes internet nebo pomocí pevného nosiče, řekněme poštou. Za normálních okolností bude první varianta rychlejší.

jsem zmínil i pojem „svět“, v tomto kontextu jej chápu jako prostředí, které není nutně bezprostředně kolem nás. Nicméně, i když jednáme pouze v rámci prostředí, kde se nacházíme, vždy při jednání využíváme svých zkušeností ze světa včetně zprostředkovaných zkušeností. Při myšlení tedy může svět otisknutý v nás vytvářet širší kontext a spojitosti.

Člověk i stroj jsou součástí světa, ve kterém fungují. Těch světů máme hned několik a každý svět determinuje i specifické prostředí, ve kterém se daná entita aktuálně nachází.

U člověka můžeme hovořit o následujících světech:

- fyzický, též nazýván reálným světem
- virtuální, do kterého spadají představy, myšlenky balancující i na hranici světa snů; dále zahrnuje komunikaci bez nutnosti fyzické přítomnosti
- svět snů, jenž je málo prozkoumán, a tedy ho ani nezrcadlíme u strojů, ani v této stati se o něm nebudu dále zmiňovat; více o světě snů [Pstružina, 1995], [Marková, 2007, str. 87]

U stroje se v současné době vyskytují pouze dva světy:

- fyzický
- virtuální,¹⁵ těch může být teoreticky neomezené množství

Zatímco člověk žije střídavě ve všech třech výše zmíněných světech, tak (softwarový) stroj v praxi pracuje pouze v jednom z nich (kupříkladu softwarový autonomní systém). Myšleno ve smyslu, že jejich cílem (náplní existence) je pracovat buď ve fyzickém, nebo virtuálním světě a také ho přímo svými reakcemi zpětně ovlivňovat.

S pojmem virtuality je spjato i její dělení z pohledu člověka a dnešních technologií, podobně jako u umělé inteligence na tzv. silnou a slabou virtualitu. Silná virtualita je dána technologií (pro pohyb v 3D virtuálním světě – helma, speciální rukavice aj.): „Pocit vnoření přichází ze zařízení, která izolují naše smysly natolik, že se člověk cítí přenesen na jiné místo.“ [Horrocks, 2002, str. 36] Slabá virtualita zasahuje i náš fyzický svět, typickými příklady jsou komunikace přes mobilní telefon, domácí počítač či bankovní terminál.¹⁶ Toto dělení uvádím spíše pro úplnost. V následujícím pojednání se budeme pohybovat z tohoto pohledu především na úrovni slabé virtuality.

¹⁵Virtuální ve smyslu interakce člověka a počítače bývá označováno výrazem virtuální realita, který nahradil dříve používaný termín virtuální prostředí či v armádě syntetické prostředí. [Heim, 1998, str. 5]

¹⁶Toto dělení a výchozí pohled do oblasti virtuality nalezneme v první kapitole [Heim, 1998].

Nyní by bylo vhodné srovnat fyzický a virtuální svět ve vztahu k člověku a stroji. V následující stati se zaměřím pouze na prostředí, kde se aktuálně daný stroj či člověk nachází, především kvůli fyzickému světu, jenž je velmi rozsáhlý.

2.1.1.1. Fyzické prostředí a receptory

Abychom mohli mluvit o fyzickém prostředí a bytí v něm, tak je třeba uvést určité poznatky o receptorech, které reagují na některé fenomény prostředí (vnějšího či vnitřního). Člověk nejčastěji mluví o pěti¹⁷ základních smyslech, což je už souhrn určitých receptorů, často uspořádaných a soustředěných na jedno místo – orgán (oko, ucho. . .). Kromě těchto máme i receptory na vnímání bolesti (nociceptory), teploty (termoceptory) a další, které jsou rozprostřené po celém těle, ovšem zaměřené pouze na určitý typ podráždění. Pak jsou ovšem „receptory“, které nejsou umístěné na jednom místě a jsou složeny z různých dalších receptorů. Představitelem takového vnímání může být „hlad“, který se uvádí jako neohraňčený pocit, na jehož vzniku se podílí více jevů. Obecně bychom mohli rozdělit lidské receptory na vnitřní, vnější a kombinované podle toho, na jaké prostředí reagují. Veškeré tyto stimuly pak vyhodnocuje mozek. V dalším povídání se zaměříme na receptory, jež nám zprostředkovávají vnější prostředí. „Hovoříme-li o lidském myšlení a jeho prvotním zaměření na jemu vnější jsoucna, pak receptory se chovají tak, že každý z receptorů je neustále rozprostřen mezi jsoucný a vytrhává z nich vnější stimuly, které pak přeměňuje do formy dále zpracovatelné lidským myšlením.“ [Pstružina, 2005, str. 49] Receptory tak představují první bariéru, která selektuje, co bude jednotlivec vědomě vnímat, také je to prvotní zatížení – omezení jen na určité vjemy – které se dále projevuje při dalším zpracování. Je to složitý mechanismus výběru stimulů a převodu vnějších jsoucen na vnitřní reprezentaci, se kterou lidské myšlení dále pracuje. Kromě stimulů z vnějšku mysl pracuje i s mentály, což si můžeme představit jako vnitřní myšlenkové stimuly, které vznikají uvnitř mysli (představy, pojmy, pocity, jazyk). [Pstružina, 2005, str. 50] Na okraj ještě zmíním, že i mnohé buňky v těle člověka mají vlastní specializované receptory, které vyhodnocuje sama buňka se zřetelem na svůj účel a kód, jenž určuje, zda a jak na percepci bude reagovat. U inteligentních systémů nás také budou zajímat jen vnější receptory, které využívá člověk při svém vědomí pro život ve fyzickém prostředí. Jedná se tedy především o pět základních lidských smyslů.

¹⁷Zrak, čich, chuť, hmat, sluch.

V souvislosti s umělými systémy se spíše než o receptorech mluví o čidlech či senzorech, v tomto směru pokládám veškeré tyto výrazy za ekvivalentní. V umělých systémech nalezneme mnoho různých druhů senzorů, které zaznamenávají podněty z prostředí a převádí je na dále zpracovatelný signál. Jsou mnohdy vyrobeny ke speciálnímu účelu – úzce zaměřené a vhodné pro jednoduché reaktivní chování, bez nutnosti jejich hlubšího vyhodnocení (přijímám nebo nepřijímám).¹⁸ Problematictějšími jsou až složitější zařízení, která samotné signály musí vyhodnotit. Takovým příkladem může být rozpoznávání hlasu či obrazu, kdy samotné zachycení obrazu nebo zvuku nám žádné informace neposkytuje. [Mařík a kol., 1997, str. 178] Primární úlohu zde hraje právě ona interpretace, o kterou se stará umělá inteligence zpravidla na softwarové úrovni. Výsledkem pak může být rozpoznání určitého slova, které teprve poté spouští příslušnou reakci. Obdobné je to i u člověka, kdy význam slov (zvuků či viděného textu) „vzniká až při střetu slyšeného nebo viděného s vnitřním obsahem mysli.“ [Pstružina, 2005, str. 49] V tomto ohledu můžeme říci, že hledáme správný podnět. Téma zpracování podnětů, respektive lidského myšlení, ovšem rozvedu až dále.

Vnímání a existence subjektu v prostředí je ohraničena jeho schopností přijímat podněty, které svým bytím vytvořilo prostředí, dále je správně zpracovat a interpretovat.

2.1.1.2. Virtuální prostředí

Když otevřeme otázku receptorů ve virtuálním prostředí, tak narazíme na problém: Dané prostředí nemá fyzickou podobu. Proto bych spíše než o receptorech a podnětech, které vytváří, hovořil o pravidlech.¹⁹ Pro člověka je to prostor, který má oporu v modelu mysli nejenom ve vztahu k fyzickému,²⁰ ale i ve vztahu možných světů – fantazii. Pokud vezmeme v úvahu počítačem zprostředkovanou virtuální realitu, tak zde opět nalezneme vztahy známé z fyzického světa, ovšem naše interakce je na jiné – nefyzické úrovni.²¹ Tyto akce a reakce virtuálního prostředí mají podstatu v pravidlech. Pokud znám pravidla, tak mohu ve virtuálním světě žít. Předvídám totiž jeho chování a sleduji své cíle v něm stejně jako ve fyzickém světě. Tento svět je však

¹⁸Pro příklad uvedu magnetické, indukční či optické senzory, které se používají u běžných „reaktivních zařízení“ typu dopravníkový pás, kdy senzor upozorňuje na vybočení pásu.

¹⁹Jsou to pravidla, která jsme vytvořili my, lidé, a jednoduše je můžeme změnit. Podle svého uvážení tedy můžeme nejen přizpůsobovat program danému prostředí, ale i pravidla světa danému programu, včetně rozšíření světa o nová pravidla. Fyzikální zákony kolem nás změnit nebo vytvořit nové nedokážeme, můžeme je pouze lépe využívat.

²⁰Při práci ve virtuálním prostředí využíváme i stejný jazyk, respektive slova jako „stisknout tlačítko“, „otevřít složku“ a další.

²¹I když klikání myši či využití jiného polohovacího zařízení probíhá ve fyzickém světě, samotná cílová akce se odehrává ve virtuálním světě. Fyzický svět je prostředník.

proměnný podobně jako naše myšlenky. Snadno mohu změnit pravidla světa a tím změnit jeho podstatu. Pokud jsem tedy zvyklý v daném prostředí na určité „uživatelské rozhraní“ a někdo jej naprosto změní, tak těžko budu moci při další návštěvě ve změněném prostředí dále interagovat, dokud se nenaučím jeho pravidla. Žití v tomto prostředí se podobá spíše modelu světa, vytvořenému naší myslí – endoceptu, viz níže. Zatímco myšlenkový model světa můžeme „snadno“ upravit stejně, jako když změníme ve virtuálním světě pravidlo,²² tak ve fyzickém světě je vše ontologicky stále stejné a neměnné. Těžko se změnit viditelné světlo v radioaktivní záření a obráceně, virtuálně to je možné, ale ve fyzickém světě značně nepravděpodobné.

Počítačem zprostředkovaný virtuální svět by se neměl snažit příliš skokově měnit, nýbrž zachovávat určitá pravidla tak, aby se tento svět stal pro člověka intuitivní a známý. Několikrát jsem použil sousloví „počítačem zprostředkovaný“, právě tak by bylo možné obecněji říci „strojem zprostředkovaný“, ze kterého vyplývá, že pokud chceme v takovém světě být (vyjma našich vlastních představ), tak musíme využívat prostředníka pro komunikaci – fyzický svět. A v tomto směru je nasnadě, že zde člověk používá své receptory pro komunikaci. Nejčastějšími smysly, které zapojujeme při takové činnosti, jako je práce s počítačem, jsou zrak a sluch.²³ Tyto smysly nám umožňují získávat podněty a ty dále vyhodnocuje kognitivní systém.

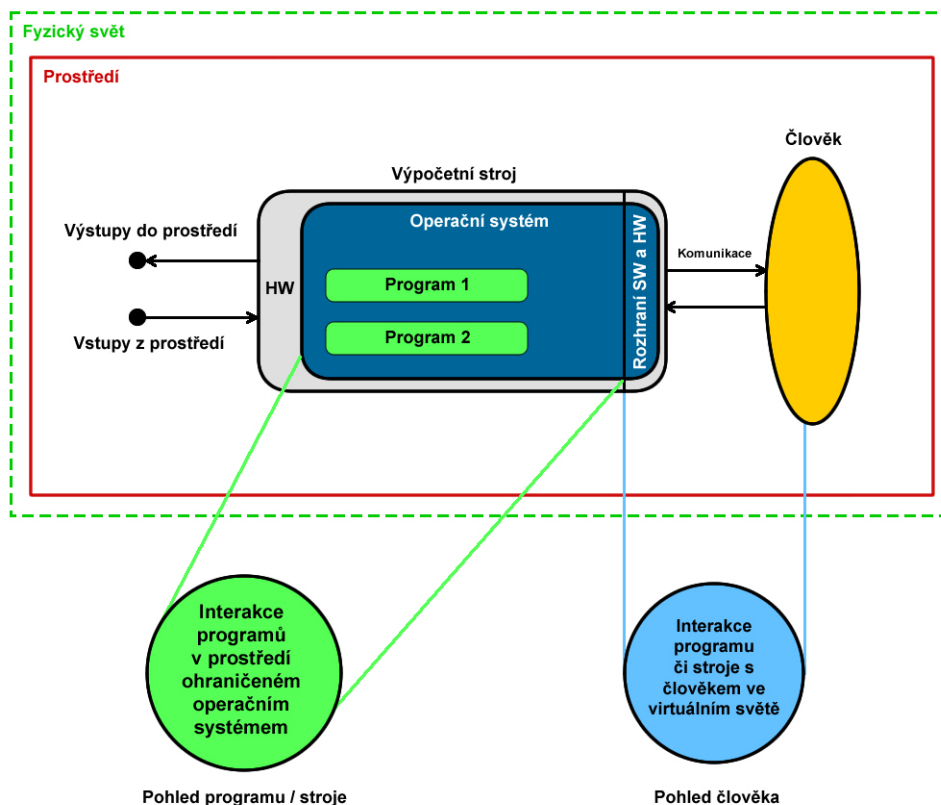
Systémy pracující pouze na softwarové úrovni opět existují ve světě (pevných) pravidel. Právě tento virtuální svět nás s nimi spojuje, a zatímco fyzická zařízení, na kterých spouštíme softwarovou aplikaci, mohou být různá, tak ovládání (komunikace) je principiálně stále stejné. Když odhlédneme od rozhraní uzpůsobeného pro interakci s člověkem, tak se nám otevře část světa, která je před našimi zraky schována. Je to vidění světa skrze program, který je ohraničen operačním systémem. Což je obdobné, jako když odhlédneme od člověka jako celku – od toho, jak se chová nebo vypadá, a zaměříme se na jeho části – na to, jak funguje uvnitř. V tomto pohledu uvidíme virtuální svět zcela jinak, přičemž zde najdeme analogie z našeho fyzického světa. Jednak tu můžeme nalézt pevné zákonitosti prostředí, které nám udává operační systém,²⁴ ale také interakce mezi programem a operačním systémem. Ačkoli tyto interakce probíhají principiálně jinak než ve fyzickém, potažmo virtuálním světě, ve kterém

²²Zatímco pod pojmem pravidlo v souvislosti s člověkem jsem uváděl příklad s uživatelským ovládáním, pro softwarového robota tím může být třeba komunikační protokol nebo API (Application Programming Interface).

²³Počet zrakových receptorů je přibližně $2 \cdot 10^8$ a sluchových $3 \cdot 10^4$. Přičemž 80% informací získáváme za pomoci zrakového analyzátoru. [Zeman, 1978, str. 162]

²⁴Dokonce v tomto pohledu, kdy je program prováděn, jsou tyto zákony neměnné, stejně jako fyzikální zákony našeho světa.

interaguje program s člověkem, tak přesto program vyčkává na určité podněty zprostředkované operačním systémem, jež dále zpracovává. Tento pohled na virtuální svět vzhledem k interakci s člověkem není podstatný, ale nabývá na důležitosti ve chvíli, kdy interakce s člověkem není vyžadována. Jedná se například o autonomní softwarové systémy, jako jsou vyhledávací boti nebo v jistém smyslu antivirové programy.



Obrázek 2.2.: Dva pohledy na virtuální svět.

Dále nalezneme zařízení (programy), která přímo či nepřímo využívají určité senzory ve fyzickém světě, nicméně s lidmi komunikují skrze rozhraní ve virtuálním světě. Takové systémy mnohdy postavené na některých metodách umělé inteligence ovlivňují skrze uživatele (člověka) i fyzický svět. Příkladem může být dnešní mobilní telefon. GPS²⁵ přijímač v telefonu je „receptor“, ke kterému má program přímý přístup přes definovaný port. Zachytává tak signál, který vysílají družice, a aby získal potřebnou informaci (o své pozici, rychlosti...), tak jej musí umět zpracovat. Také

²⁵Global Positioning System – v dnešní době nejpoužívanější. Kromě tohoto amerického systému pro určování polohy existuje i ruský GLONASS či evropský GALILEO.

může nepřímo zprostředkovaně získat informace ze vzdáleného senzoru. K tomu se dnes využívá například technologie GPRS²⁶. Přes GPRS se připojí k internetu, kde podle daného protokolu osloví specifický server, který shromažďuje informace o počasí v různých lokalitách naší planety. Po odeslání požadavku s uvedenou pozicí ve WGS²⁷ souřadnicích přijde nazpět odpověď o počasí v lokalitě, kde se nachází. Po vyhodnocení získaných údajů naplánuje daný program vhodnou trasu (s využitím heuristického prohledávání) s tím, že zohlední, že má být námraza a sníh. Využije tedy při svých plánech jen silnice vyšší třídy, které budou pravděpodobně udržované. Své závěry včetně informací, které vypracoval, předloží uživateli, který je rovněž vyhodnotí a případně provede změny v trase dle svého uvážení. Tímto umělý systém ovlivňuje naše rozhodnutí a tím nepřímo působí i na fyzický svět.

2.1.2. Podněty a jejich zpracování u člověka a umělé inteligence

Prostředí svým přirozeným bytím (existencí) vytváří podmínky pro vznik podnětů. Podněty nebo také stimuly jsou vybrané části světa, které dokážou naše receptory zachytit – zaznamenáváme tím jejich existenci. Možnosti získávání a posléze zpracování určitých podnětů v přírodě jsou dány biologickou evolucí, která probíhá na naší planetě. V oblasti umělé inteligence hraje důležitou roli člověk. Jelikož sám vytváří inteligentní stroje, zužuje (selektuje) jim množství a druh podnětů, které získávají z okolí, pouze na ty, které jsou nezbytně nutné k tomu, aby stroje plnily svůj úkol. Toho docílujeme specifickými senzory nebo programy pro zpracování získaných podnětů. Tak je tomu alespoň v případě, kdy inteligentní stroj pracuje se stimuly pocházejícími přímo z našeho reálného, respektive fyzického, světa. Příkladem takového stroje může být například agent s externím řízením reagující na světlo – když světlo sílí, tak změni směr svého pohybu. V tomto případě prostředí přímo determinuje chování agenta, které je reaktivní. „Racionalita není vlastností agenta, ale určitou charakteristikou vztahu agenta a jeho prostředí.“ [Mařík a kol., 2001, str. 166] „Základním problémem (reaktivního agenta – pozn. Z.S.) je návrh nejvhodnějšího souboru senzorů, který umožní robotům situovaným v konkrétním prostředí dosahovat autonomně určených cílů.“ [Mařík a kol., 2001, str. 169] Příkladem reaktivního agenta z živočišné části přírody můžeme uvést mravence, který je jako jednatel při svém rozhodování ovlivněn koncentracemi určitých látek, jež jeho receptory dokážou rozpoznat. Podle toho se pak pohybuje ve svém prostředí. [Mařík a kol., 2001, str. 166]

²⁶General Packet Radio Service

²⁷Myšlen WGS84 – World Geodetic System 1984.

Jednodušším příkladem z rostlinné říše nám mohou být některé rostliny, které se otáčejí za zdrojem světla (heliotropismus).

Inspirací je nám samozřejmě i člověk a jeho chování, které je determinováno i vnitřními myšlenkovými pochody. Kromě reaktivního agenta existuje též deliberativní agent,²⁸ jehož rozhodování je podmíněno také vnitřním zpracováním. Rozhodování takového agenta není přímo vázáno pouze na prostředí, nýbrž i na zpracování signálů za pomoci programu, který určuje jeho konečnou reakci (rozhodnutí). Obdobnou paralelu deliberativnosti najdeme u člověka, který některé podněty musí vyhodnotit a teprve poté vhodně reagovat,²⁹ za účelem dosažení svého cíle. Stejně tak je nám vlastní i reaktivní, či lépe řečeno, reflexivní chování. „Mícha je centrem reflexů [...], reflex je přímou odpovědí organismu na určité vnější nebo vnitřní podněty“. [Klíma, 2008, str. 130] Reakce na vnější stimul je rychlejší díky reflexivnímu oblouku, kde řídicí úlohu má šedá hmota míchy. Mozek je informován vzestupnými drahami, ale nemůže zasáhnout do toho (ovlivnit), co se děje. Co se týká reaktivního jednání, tak zde nám inspirací byly spíše jednodušší formy života.

Po základním vhledu do oblasti percepce přichází důležitá otázka: Jakým způsobem s prostředím kolem sebe interagujeme? Člověk filtruje podněty přicházející z receptorů po aferentních drahách do mozku. Do samotného vědomí se tedy dostane jen zlomek podnětů, jež byly zachyceny. Mozek dále srovnává přijaté podněty se svým modelem světa, a tím se dostáváme již k tématu dalších podkapitol, ale myslím si, že je vhodné zmínit celý kognitivní proces alespoň v hrubém náčrtu pro pochopení celé problematiky. Karel Pstružina v souvislosti s vnitřním modelem světa poukazuje na tzv. endocepty. „Když tedy dospělý člověk vnímá své okolí, pak jsem předchází jeho vlastní myšlenkový model světa. Tento model světa budeme nazývat endoceptem, neboť v protikladu k percepci vyjadřuje vnitřní aktivní vybavování světa. Nelze jej nazvat zkušeností, protože zahrnuje nejen smyslově představivostní, ale i pojmové modely světa.“ [Pstružina, 1994] Důležitou ideou v tomto kontextu je niterní model světa, něco jako jeho odraz v nás, který nám umožňuje díky myšlenkovým procesům, jež je využívají, předvídavost.³⁰

²⁸ Stejně principy uváděné u agentů ve spojitosti s fyzickým prostředím nalezneme i v čistě virtuálním prostředí. V tom je určitá záludnost, neboť přístupy k umělé inteligenci ve virtuálním a fyzickém světě se prolínají. To je dáno i tím, že virtuální svět ve spojitosti s člověkem má ve fyzickém světě do značné míry svůj předobraz.

²⁹ Vyhodnocení probíhá v mozku, přičemž reakci si uvědomujeme a mnohdy ovlivňujeme (zkušenostmi či učením), ať už se nakonec jedná více o emocionální, či racionální jednání. Důležitou roli zde hraje zmíněný vnitřní model mysli, jež předchází samotným vjemům.

³⁰ V určitém smyslu můžeme mluvit o neustálém předvídání budoucnosti, obdobně jako v jiném konstruktu naší mysli – matematice (viz podkapitola 1.3.). V tomto případě to znamená, že pokud na-

Je podstatné zmínit, že ačkoli to vypadá, že jednotlivé děje – vnímání a zpracování či myšlenkové procesy – jsou odděleny, tak tomu tak není. Jedná se o současný děj, kdy percepce probíhá souběžně s představivostí a pamětí. Vnímání „je neoddělitelné od počitků, není jejich důsledkem.“ [Lukacs, 2009, str. 74]

U člověka se ve smyslu předvídavosti mluví spíše než jen o modelu světa o „nastavbě“ ve smyslu procesu plánování. „Lze říci, že plány obsahují schémata a scénáře možných aktivit, které lidské myšlení vytváří, a pak vzhledem k aktuálním vjemům jimi prochází a vybírá z nich ty, které mu umožňují svou představu o sobě samém jako budoucím naplnit.“ [Pstružina, 2005, str. 55] Je to samozřejmě něco více než jenom jednoduchá reaktivita pozorovaná kupříkladu u reaktivních agentů či bakterií. Ačkoli nalezneme pokusy o tvorbu vnitřní reprezentace světa u uvažujícího agenta [Mařík a kol., 2001, str. 176], přesto tyto pokusy jsou zatím úzce zaměřené na jisté prostředí (jeho stavy) a postrádají univerzálnější reprezentaci (model). Navíc, takový model je pak u každého umělého agenta ve vztahu k objektu stejný, ale endocept má každý člověk jiný, tedy i vztah každého člověka k určitému objektu je jiný. Dle mého názoru má být niterní model světa směrem k silné umělé inteligenci něco více než „mapa“, která nám pomůže dostat se z bodu A do bodu B, má spíše zahrnovat systém dřívějších zkušeností (v nejobecnějším významu) včetně vyhodnocení užitku (slasti), které dřívější řešení přineslo.³¹ Každý inteligentní systém by reagoval na základě svých zkušeností z prostředí, kde je nasazen, jinak. I přesto, že na začátku, po sjetí z výrobní linky, by měly všechny systémy stejné možnosti, postupem času by se přizpůsobovaly prostředí.³² Inteligentní systém by sám sebe na základě zkušeností upravoval na míru potřebám svého nasazení.

Je v tom skrytá autonomie umělého systému, záleží tedy i na člověku, jakou auto-

příklad vidíme část objektu, spojujeme tuto část s námi známým celkem. Pokud vidím auto z boku, kde má dvě kola, předpokládám, že další kola má i na druhé straně.

³¹Zkusím tuto myšlenku demonstrovat na příkladu, který je inspirován rozhovorem s Karlem Pstružinou. Předpokládejme, že člověk i stroj by měli podobné možnosti vnímání prostředí v oblasti automobilové dopravy. Představme si dopravní situaci na křižovatce. Ve chvíli, kdy na semaforu ze zelené přejde stav na oranžovou, tak bychom už neměli vjíždět do křižovatky. Opatrný řidič by zastavil, ale v tomto případě máme řidiče se „zkušeností“, který spěchá a situaci vyhodnotí jinak – začne plánovat. Je si vědom prodlevy mezi tím, kdy mu naskočí červená, a tím, kdy je na řadě zelená. Také vidí kamion, kterému bude trvat déle, než se rozjede, což platí i o autě s vlečným zařízením na druhé straně. Navíc má dobré rychle reagující auto, ve kterém se pohybuje určitou rychlostí. Promítá své zkušenosti do aktuální situace a rozhodne se naopak zvýšit rychlost a křižovatkou projet. Celé toto vyhodnocení se nemusí udát ve vědomí, ale i mimo jeho rámec (nevědomě). Za jiné situace v rušném provozu by se rozhodl naopak zastavit a neriskovat. Inteligentní dopravní systém by se v této situaci choval stále stejně podle určitého vzoru.

Obdobný příklad by byl se stejnou křižovatkou v noci, kdy by řidiči naskočila červená. Nicméně křižovatka by byla prázdná a nikde žádné auto. Řidič by se tedy mohl rozhodnout pomalu křižovatkou projet. Přitom ve svém nitru kalkuluje s různými možnostmi, včetně té, že policie čeká za rohem.

³²Dnes se tento problém řeší zakázkovou výrobou a širokou oblastí řízení podnikové informatiky.

mii dovolí. Avšak právě díky oné variabilitě každého člověka se dokáže lidstvo jako celek velmi rychle přizpůsobovat měnícímu se prostředí, ve kterém existuje. V tomto směru je to určitě zajímavá myšlenka nejen ve spojení s multiagentními systémy, které jsou založeny na spolupráci (komunikace a kooperace) mnoha agentů. Ve svém důsledku nemusí vést výsledek (vyhodnocení podnětu) nutně k racionálnímu řešení, což ovšem stále spadá do chování člověka, k němuž bychom měli chtít směřovat.

Po vyhodnocení následuje reakce. I při reakci na podněty využíváme svůj vnitřní model světa a korigujeme své reakce. „Předpokládá se, že své akce a interakce s objekty reprezentujeme jako niterné modely. Předpovězené senzorické důsledky jsou srovnávány se skutečnými důsledky a užity k optimalizaci motorické kontroly.“ [Koukolík, 2010, str. 29]

Jiný pohled na zpracování podnětů je třeba vidět ve virtuálním světě, který je výtvorem člověka, když nebudeme brát v úvahu oblast lidských myšlenek u jedince, která je tomuto světu inspirací. Opět záleží, zda náš pohled bude z hlediska uživatele nebo programu. V prvním případě se jedná o prostředí pravidel, které vytvořil člověk, kde podněty vznikají při interakci člověka (skrže komunikační rozhraní) a stroje na obou stranách. Prostředníkem je zde fyzický svět, kdy výsledky své práce musí program člověku představit tak, aby je mohl vnímat (zpravidla jako zvuk a obraz). Zpracování podnětů na straně softwarového programu probíhá dle algoritmu, přičemž by měl zahrnovat veškeré situace, ke kterým může dojít. V opačném případě dochází k chybě a neočekávanému chování. Příkladem takového vzájemného působení může být nejen online svět her nebo offline práce v tabulkovém kalkulátoru, ale také práce s expertním systémem, se kterým v dialogovém režimu člověk komunikuje.

Trošku jiná situace nastane, když je stroj využíván jen jako mediátor v komunikaci mezi lidmi, i když pravidla virtuálního světa platí i zde. V zásadě je to však obdobné, neboť se oba dva (komunikátor i komunikant) nacházejí ve virtuálním prostředí omezeném pravidly.

Z pohledu programu veškerá interakce probíhá pouze v prostoru ohraničeném operačním systémem či sítí – kyberprostorem. Pokud půjdu do důsledků, tak i program mohu brát jako ohraničené prostředí, ve kterém mohou iteračně interagovat softwaroví agenti mezi sebou (např. v rámci programu Netlogo). Podněty jsou opět zpracovávány algoritmicky, ovšem není zde nutná interakce agentů mimo toto prostředí (např. s člověkem). V tomto případě se jedná čistě o interakci mezi programem a prostředím i v případě, že spolu interagují dva programy (procesy), neboť prostředníkem jejich

komunikace je operační systém, který vytváří samotné prostředí. Lépe řečeno, každý program má vlastní prostředí dané tzv. kontextem programu, a pokud chce cokoli udělat, tak musí volat služby nabízené operačním systémem. Interakce (např. komunikace) dvou procesů je tedy problematika spojená právě s operačním systémem. Každopádně i zde platí jednoduché pravidlo, když jeden něco komunikuje, tak druhý by měl počkat, dokud předchozí neskončí.³³

Ve fyzickém světě vše, co člověk (stroj) přijímá, je zprostředkované a vyfiltrované – pracujeme pouze s částečnými odrazy jsoucna v nás, s neúplnými informacemi. Nelze tedy vzít jsoucno, včlenit jej fyzicky do nás a mít tak o něm veškeré informace. Ve virtuálním světě při komunikaci to lze, příkladem může být předávaná struktura popisující objekt mezi dvěma procesy. Nutno dodat, že obdobné je to i u lidského myšlení jednotlivce.

Specifickým problémem při interakci nejen v prostředí, ale i mezi jednotlivými subjekty v rámci komunikačního schématu, je šum. Jedná se o zkreslení, ke kterému dochází v průběhu přenosu médiem. V souvislosti se zpracováním takto poškozených signálů jsou důležité opravné prostředky, které využívají jak stroje, tak člověk. U člověka takovým prostředkem může být kontext situace, ze kterého odvodí i to, co už nezaznamenal. V případě stroje záleží na druhu komunikace, a podle toho i na zvolené metodě – příkladem může být kontrolní součet.

2.1.3. Shrnutí problematiky prostředí a podnětů

V této podkapitole byly popsány obecné principy interakce člověka a stroje ve fyzickém a virtuálním prostředí. Dalším bodem mého zájmu se staly receptory a vytváření podnětů se zevrubnými náznaky principů zpracování, které budu dále rozvíjet. Vzhledem k šíři tohoto tématu jsem se zaměřil především na oblast, kde se střetává člověk a stroj, neboť právě zde je část, kterou se snaží naplnit tradiční umělá inteligence. Zároveň jsem představil i pohled čistě virtuální bez nutnosti lidské interakce, jenž je v současnosti typičtější především pro novou umělou inteligenci.³⁴

S prostředím a podněty souvisí i chování, které se nejčastěji omezuje buď na reaktivní nebo uvažující, kdy reaktivitu nalezneme především u nové umělé inteligence či

³³K čemuž se používají jednoduché semaforey nebo zamykání (mutex). Více o těchto principech a světě, kde softwarový program běží, lze nalézt v [Čada, 1993] a nověji zpracované v [Fišer, 2003].

³⁴Především z hlediska autonomních systémů s adekvátně navrženými senzory tak, aby plnil celý systém svůj cíl.

umělého života, kdežto uvažování je typičtější pro tradiční umělou inteligenci, ovšem prolíná i distribuovanou umělou inteligenci.

2.2. Myšlení a odraz v kognitivním systému

Tato podkapitola je specifická tím, že se zabývá evolučně posledními fenomény spojenými především s typicky inteligentním chováním člověka. Je to tedy oblast, kterou zastřešuje tradiční umělá inteligence využívající tzv. kognitivní systém.³⁵ Z těchto důvodů je třeba blíže se seznámit s vyššími kognitivními procesy u člověka. V této části je také srovnávám s jednoduššími, avšak do značné míry analogickými principy u výpočetních strojů. Vybírám zde pouze některé fundamentální procesy, které člověka a stroj spojují a diskutuji i otázku pohledu na člověka jako na stroj.

2.2.1. Centrum kognitivních funkcí

Mozek člověka jako centrum lidských kognitivních funkcí je nám stále záhadou. Zatímco naše mechanické popisy fungování ostatních orgánů v lidském těle celkem vystihují jejich podstatu, tak u mozku nedokážeme vysvětlit všechny jeho funkce, respektive fenomény, které s jeho činností pozorujeme. Klasický mechanický (fyzický či neurologický) pohled nám dává mnoho informací, od průměrné hmotnosti 1,35 kg, počtu buněk a jejich propojení³⁶ až k samotné činnosti na úrovni neuronů a přenosu vzruchů synapsemi. Nezodpovězenou otázkou zůstává naše vyšší mozková činnost – mysl, kam spadá i lidská inteligence. Jediné, co zatím z fylogenetického hlediska můžeme říci, je, že mozková kůra, kde vzniká naše racionální uvažování, se objevuje už u plazů. Ti ovšem svůj mozek budovali na bazálních gangliích³⁷ (nuclei basales), které jsou součástí koncového mozku (telencephalon). Bazální ganglia, kde je šedá hmota soustředěna ve shlucích, mají mnohé „konstrukční“ problémy, především při jejich zvětšování a zásobování. Kvůli tomu také po 240 milionech let, kdy vládli naší planetě, nezačali být inteligentní obdobně jako člověk.³⁸ Savci začali svůj mozek nao-

³⁵ Jež plní následující funkce: paměti, autonomního řešení, učení, plánování, vyhledávání a rozhodnutí. [Pstružina, 1998, str. 137]

³⁶ „Počet buněk v jednom krychlovém milimetru šedé kůry se pohybuje mezi 100 000 až 150 000. Každá buňka může být spojena vlákny s dalšími buňkami – zde se čísla pohybují mezi 5000 až 200 000.“ [Klíma, 2008, str. 106]

³⁷ Zajišťují koordinaci pohybů a ovlivňují i naši emoční stránku (mají blízko k limbickému systému).

³⁸ „Existují výpočty, podle nichž by náš mozek musel být asi pětikrát větší, než je, kdyby byl zkonstruován jenom na principu bazálních ganglií a měl přitom obsahovat stejný počet nervových buněk, které má.“ [Klíma, 2008, str. 104]

pak budovat na mozkové kůře³⁹ (neocortex), která se postupem vývoje začala vrásnit (rýhovat). „Uspořádání nervové tkáně v ní vede k nadměrnému zmnožení nervových buněk a jejich spojů. Tvoří se jich nadbytek. [...] Tím dochází k nové kvalitě nervové činnosti, kterou lze nazvat opravdovou inteligencí.“ [Klíma, 2008, str. 106] K zajištění života, jako má zvíře, bychom tedy velkou část neocortexu ani nepotřebovali. Proč se mozek našich předků začal zvětšovat? Nejznámější je teorie o příjmu masité potravy, neboť mozek spotřebuje na svou velikost hodně energie, ale také zvětšování sociální skupiny a nutná komunikace mezi jejími členy. [Krámský ed., 2009, str.189] V současné době již vědci ve svých názorech kombinují více aspektů, jež se podílely na kvantitativních změnách mozku.⁴⁰

Uvádí se, že právě mozková kůra je centrem naší vyšší mozkové činnosti, avšak zatím neznáme přesné principy, jak mysl z mozku „povstává“. Vzniká tak několik filozofických teorií, které se snaží vysvětlit vztah mysli a těla (tzv. mind-body problém). Především se jedná o dualistické koncepce, které poukazují na souvislost a propojenost mysli a těla (jejich interakci). „Řekněme, že myšlenkové procesy jsou doprovozeny procesy předávání bioelektrických a biochemických impulsů mezi neurony.“ [Pstružina, 1994] Ale stejně tak zde nalezneme striktní dualistickou separaci, kde mysl a tělo na sebe vůbec nepůsobí.⁴¹ [Hajnal, 2005] Dualismus odmítá redukci mysli a těla jeden na druhého – mysl je kvalitativně něco jiného než tělo a obráceně. Naproti tomu stojí monistické koncepce, které buď pokládají mysl za jedinou reálnou a fyzický svět je produktem mysli (mentalismus či idealismus), nebo přesně naopak prohlašují, že fyzikálně či mechanisticky lze vysvětlit i mysl (fyzikalismus či materialismus). [Krámský ed., 2009, str. 19]

V oblasti umělé inteligence, i vzhledem k tomu, že má ještě daleko k tomu, co označujeme jako mysl, se uplatňuje monistická koncepce (materialismus), tedy, že mysl můžeme redukovat na softwarovou reprezentaci, ať už se jedná o konekcionismus nebo funkcionální komputacionismus. [Krámský ed., 2009, str. 21] Kritici namítají, že právě toto omezení – redukce – nám nedovoluje přiblížit se naší mysli v silném slova smyslu.

³⁹Toto poněkud odvážnější tvrzení mohou opřít o [Klíma, 2008, str. 106]. Pro vysvětlení ještě uvedu, že součástí zmiňovaného koncového mozku (též velký mozek, telencephalon) jsou bazální ganglia (šedé shluky uvnitř bílé mozkové hmoty), limbický systém (neboli také čichový mozek, rhinencephalon) a mozková kůra (tvořená šedou hmotou, neopallium či neocortex). Podrobnější informace o vývoji neocortexu a srovnání jeho velikostí u různých savců viz [Koukolík, 2006, str. 9 - 20].

⁴⁰Podrobněji se lze s těmito aspekty seznámit v [Krámský ed., 2009, str. 190 - 194].

⁴¹Příkladem takového myšlení je teorie okasionalismu, která počítá s agensem mimo náš svět (bůh), jež ke každé fyzické akci vytvoří událost druhého typu.

Tyto myšlenky mají své filozofické kořeny a přinášejí nám i různé modely, jak by mohly probíhat procesy naší mysli.⁴² Na druhé straně, s rozvojem možností počítačového modelování vytváříme i simulace mozku.⁴³ Jako příklad bych uvedl současný projekt „Blue brain“ pod vedením prof. HENRYHO MARKRAMA ve Švýcarsku. Tento projekt vychází z 15 let trvajících experimentů se šedou kůrou mozkovou potkanů.⁴⁴ Cílem by mělo být v příštích pěti letech vytvořit model této části mozku, a to kompletně ve 3D s veškerými interakcemi, ke kterým v neocortexu potkanů dochází. [Blue Brain, 2010] Raději bych upřesnil, že se při modelování zahrnou všechny známé interakce, ke kterým dochází.

Nápad inspirovat se mozkem, respektive fyzikálními procesy v mozku, a představit umělou neuronovou síť přišel již v 50. letech 20. století. Tehdy představil FRANK ROSENBLATT jednoduchou učící se umělou neuronovou síť – tzv. perceptron. Vycházel samozřejmě ze starších prací, které se snažily jednoduše popsat fungování neuronu, ze kterých se mozek skládá. WARREN MCCULLOCH (1898 - 1969) a WALTER PITTS (1923 - 1969) vytvořili už ve 40. letech 20. století formální popis fungování neuronu. [Petrů, 2007, str. 67] Přestože všechny tyto inspirace živou přírodou nám daly nové technologie a možnosti, jak řešit některé problémy, tak k opravdovému fungování neuronu mají daleko. Problémem je zde naše velké abstrahování či zjednodušení složitých elektrochemických procesů, jež mezi neurony probíhají. Proto s postupným vývojem počítačové techniky přichází i větší možnosti podrobnějších softwarových modelů, jako ve zmíněném projektu „Blue brain“.

V rámci tvorby inteligentních systémů v dnešní době rozlišujeme mezi dvěma hlavními architekturami či proudy, jejichž cílem je modelování lidského nervového systému: konekcionistickou a logicko-symbolickou.

„Konekcionismus je založen na myšlence, že na mnohé složité jevy, myšlení a inteligenci nevyjímaje, lze pohlížet jako na emergentní vlastnosti paralelních dějů v rozsáhlé síti třeba i jednoduchých a vzájemně podobných aktivních prvků (formální neurony), mezi nimiž existují interakční vazby (links).“ [Mařík a kol., 2001, str. 38] Mezi takové systémy řadíme i zmíněné neuronové sítě, jejichž opětovný vzestup na oblíbenosti nastal v 80. letech 20. století.⁴⁵ Jedná se o komplexní nelineární dynamické sítě,

⁴²Viz [Pstružina, 1998]

⁴³A tak trochu doufáme, že se nám tam podstata mysli alespoň částečně ukáže.

⁴⁴Na rozdíl od lidského mozku není mozková kůra rýhovaná a má velikost hlavičky špendlíku.

⁴⁵Ačkoli perceptron vznikl již v 50. letech, tak bohužel kniha „Perceptron: An Introduction to Computational Geometry“ (Minsky a Papert) zabrzдила rozvoj konekcionistického přístupu (viz první kapitola). Podle jiných názorů byl útlum způsoben tím, že logicko-symbolická perspektiva byla stále ještě zajímavější. [Petrů, 2007, str. 70]

kteřé mají blízko k chaotickému chování (deterministický chaos). Z hlediska vztahu mysli a těla tedy není překvapivý názor, že lidská mysl sama emerguje ze složitého dynamického systému, jakým je náš mozek, a je jednoznačně determinována fyzickým mozkem a prostředím, ve kterém se nachází. Jinak řečeno, chování či činnost naší mysli je teoreticky předvídatelná. [Veselý, 2005]

Starší a také tzv. tradiční přístup k umělé inteligenci staví na jiných principech: logicko-symbolických, symbolicko-reprezentačních, algoritmických či komputacionistických. [Mařík a kol., 2001, str. 20] Vychází se zde z matematiky, respektive logiky, a práce se symboly za pomoci programu (algoritmu). Využívá se zde reprezentace světa, která je uložena v paměti. Příkladem takového systému může být heuristické prohledávání stavového prostoru, o kterém jsem se zmiňoval v souvislosti s vyhledáváním optimální trasy u navigačních systémů. Pokud přejdeme názor, že mysl je pouhá funkce, jenž je typický pro funkcionalistický komputacionismus, tak v tomto ohledu u systémů tradiční umělé inteligence panuje názor, že je to postačující prostředek pro zajištění inteligentního chování umělého systému. [Veselý, 2005] Což znamená, že tímto tradičním přístupem dokážeme zajistit kupříkladu chování podobné člověku, ovšem už nikoli jeho intencionální prožívání – vědomí.

Jak jsou na tom oba proudy dnes? Určitě je třeba zdůraznit, že za posledních třicet let dosáhl konekcionistický přístup v neuronových sítích velkých pokroků.⁴⁶ „Konekcionistické sítě jsou schopné řešení stále složitějších úloh a v mnoha případech začaly konkurovat logicko-symbolickým systémům.“ [Veselý, 2005] Výhodou tohoto řešení je i to, že není třeba hledat algoritmus řešící danou úlohu. Kritici však dodávají, že až na několik zajímavých aplikací řeší, oproti systémům založených na logicko-symbolické architektuře, jednoduché problémy, a spíše, než jako jeden z hlavních proudů, by měl být brán jako druhotný. Dále poukazují na neschopnost těchto sítí učit se něčemu jako je logika tak, jak ji dnes jako lidé chápeme a používáme nejen u logicko-symbolického přístupu. S tím úzce souvisí i to, jakým způsobem reprezentovat v takových sítích symboly. Nabízí se možnost reprezentace holisticky, tedy aktuální konfigurací sítě [Veselý, 2005], nebo holografickou pamětí [Petrů, 2007, str. 72]. Důležitou vlastností holografické paměti je její schopnost i částečného vyvolání (při poškození),⁴⁷ což koreponduje se subjektivní zkušeností s vybavováním vzpomínek. Také možnost uložení množství dat na malém prostoru je daleko větší. Představitel této myšlenky je K.

⁴⁶Kohonenovy sítě, Hopfieldovy sítě, modulární neuronové sítě, fuzzy neuronové sítě aj.

⁴⁷Každá část obsahuje informace o celku. To znamená, že i z vybrané části uvidíme celek, ovšem jen v dané perspektivě. Tedy obraz bude nejasný či zamlžený v poškozených částech.

H. PRIBRAM, nicméně už se tím dostáváme na hranici v současnosti aplikovatelných teorií.

K tomuto se váže i objevení tzv. neuronálního rytmu [Gray, Singer, 1989]. Jedná se o pokus s kočkami, kdy jim byly předkládány do zorného pole různé pruhované vzorce. Neurony pak na různé vzorce reagovaly jinými frekvencemi, přičemž jednotlivé neurony byly od sebe velmi vzdálené (v rámci dané oblasti). Výsledkem je tedy konstatování, že zpracování určitých vjemů závisí na správném rytmu a sladění různých neuronů. Což také nahrává konekcionistickým myšlenkám, včetně zmíněných teorií o reprezentaci a ukládání informací.

Posuňme se kousek dál od konekcionistické pevně spojené sítě k multiagentním systémům distribuované umělé inteligence. Otevře se nám nová oblast, která je založena na volné interakci. Zajímavým počinem může být zvyšování deliberativnosti (intencionality) agentů a vytváření složitých interakcí směrem ke kolektivní inteligenci. Celý systém složený z více úzce profilovaných subsystémů by pak mohl jednat inteligentně, podobně jako člověk při svém racionálním chování.

V současné době v oblasti aplikované umělé inteligence máme několik proudů, které směřují především ke specializaci na určité prostředí, tedy i úlohu. Do budoucna bude důležité rozvíjet současné přístupy k modelování lidského nervového systému, nicméně jako zajímavé se jeví spojený výzkum v oblasti biologie a biochemie a následná simulace procesů v počítačovém modelu. To nám může dát další impuls k rozvoji kognitivních procesů umělé inteligence.

2.2.2. Učení a paměť

Paměť a učení patří k základním charakteristikám lidské inteligence a obě tyto činnosti, tedy pamatování a učení, jsou úzce propojené. Tento vztah může být například na úrovni toho, co si člověk při učení dlouhodobě zapamatuje a co naopak zůstane nepovšimnuto. K těmto procesům se dále přidává třetí, který je využívá především v tvůrčím slova smyslu – myšlení.

Vhodné by bylo začít obecnými definicemi, jež vyjadřují co je to učení a paměť. Učení můžeme chápat „jako proces získávání individuální zkušenosti.“ [Hanušová a kol., 2006, str. 69] Paměť si můžeme představit jako záznam v obecném smyslu, který dále hraje důležitou úlohu jak u živého organismu, tak u výpočetního stroje. Důležitá je schopnost vybavování z paměti, především u živých organismů.

U člověka rozlišujeme několik druhů paměti. Naznačím v této části alespoň základní dělení dle [Rusina, 2004].

Základní dělení je na explicitní a implicitní paměť. Explicitní paměť je vědomě používána k zapamatování informací, a dále ji rozdělujeme na epizodickou a sémantickou. Epizodickou paměť můžeme nazvat také emoční, neboť jsou to vzpomínky na události, ke kterým máme osobní vztah (první polibek, zážitky z dětství). Naopak sémantickou můžeme označit jako faktografickou. Jsou to různé faktické informace, které přijímáme, například letopočty nebo některé vzorce.

Implicitní paměť není nutně vědomě vybavována, lze říci, že je do značné míry automatická. Skládá se z našich motorických dovedností, chování či zkušenosti. Opět se rozděljuje na procedurální paměť a priming. Procedurální paměť je vlastně motorické učení a dovoluje automatizovat určitou činnost do budoucna. Příkladem může být jízda na kole nebo nejrůznější sporty. Čím více nějaký sport provozujeme, tím více automaticky provádíme specifické motorické akce. Pro tento typ paměti je ve srovnání s explicitní pamětí (epizodická + sémantická) typické pomalé učení, ale na druhou stranu velká odolnost vůči zapomínání. Priming si lze představit jako zvyšování paměťové výkonnosti při opakovaném vystavení podobnému stimulu. Příkladem může být neúplná črta, kdy ji postupně dokreslujeme, dokud jednotlivec nepozná, co to je. Když mu ji někdy v budoucnu předložíme znovu, tak ji pozná dříve.

S pamětí souvisí i další dva důležité pojmy: krátkodobá a dlouhodobá paměť. Krátkodobá paměť⁴⁸ uchovává informaci 30 až 40 s, přičemž je zde důležité soustředění na činnost či objekt. Naopak dlouhodobá paměť dovoluje vybavování s větším časovým odstupem (minuty a více) i po zaměření pozornosti jiným směrem. Při práci s explicitní pamětí, především z hlediska vybavování, hraje důležitou úlohu hippocampus, který funguje obdobně jako knihovník v knihovně. „Všechna data, která jsou zachycena v různých korových centrech (např. zrakové či sluchové vjemy, řeč apod.) a uvědomována, se dostávají do hippokampální krajiny, kde jsou zpracována a dále převedena do příslušných asociačních oblastí mozkové kůry k následnému uložení.“ [Rusina, 2004] Problémem lidské dlouhodobé paměti je její nepřesnost, zkreslenost při pozdější vybavnosti. Což znamená, že si mnohdy vybavujeme i to, co jsme si v dané souvislosti reálně nezapamatovali – tak trošku si některé věci domýšlíme.⁴⁹ Už SIGMUND FREUD (1856 - 1939) upozorňoval, že naše vzpomínky jsou směsí pravdy a

⁴⁸ELEANOR ROSCHOVÁ spíše uvádí, že krátkodobá paměť trvá do 20 s, každopádně můžeme říci, že je to individuální, ovšem řádově v desítkách sekund. [Hayward, Varela, 2009, str. 127]

⁴⁹Ovšem vždy si můžeme vybavovat pouze to, co bylo do paměti dříve uloženo.

výmyslu. [Mollon, 2000, str. 60] Naproti tomu paměť počítače je přesná a má i kontrolní mechanismy, které ověřují, zda nebyl záznam poškozen nebo nevhodně změněn. U systémů s umělou inteligencí je třeba opět rozlišovat práci s pamětí podle toho, v jakém prostředí pracuje, zda ve fyzickém nebo virtuálním. Principiálně se při práci s fyzickou pamětí vychází z tzv. von Neumannova schématu. Technologie záznamu a uchování dat se mění, avšak principy práce s ní zůstávají na hardwarové úrovni stejné. Na softwarové úrovni, lépe řečeno na úrovni operačního systému, prodělala práce s fyzickou i virtuální pamětí za posledních 20 let mnoho změn. Souvisí to především s vývojem operačních systémů a v poslední době s přechodem na 64-bitové procesory, které umožňují práci s (virtuální) pamětí, z dnešního pohledu téměř neomezené velikosti.⁵⁰ [Fišer, 2003] Omezení daná fyzickou velikostí paměti však stále zůstávají. Reprezentace dat je fyzicky provedena binárním způsobem realizovaným v elektronických obvodech (operační paměť či EEPROM u flashdisků) nebo magnetickým záznamem (pevné disky).

Pokud umělý systém interaguje ve fyzickém prostředí, tak podobně jako člověk zaznamenává pouze určité odrazy světa kolem sebe, které převádí na signály a posléze i struktury (proměnné). Tyto struktury buď okamžitě v této své instanci využije a uloží je do operační paměti (krátkodobá, závislá na instanci spuštěného procesu, avšak rychlá), nebo je uloží na pevný disk (dlouhodobá, pomalá),⁵¹ odkud mohou být kdykoli vyvolány, respektive přesunuty do operační paměti. Co kam se uloží, je plně determinováno programem. Je to ale statický přístup předem daný kódem programu a nemění se. Selektce a ukládání do paměti u člověka je dynamická, mění se v čase v závislosti na lidském myšlení jako celku. Do procesu ukládání tak zasahují například emoce či aktuální naladění. Na druhou stranu je třeba podotknout, že umělé systémy nepotřebují k dosažení svého cíle tak složitý proces, jaký nalezneme u člověka (prostředí i činnosti, které vykonávají, jsou jednodušší).

Pokud se ocitneme ve virtuálním prostředí na programové úrovni (procesní), tak je princip poněkud odlišný, a to v jednom základním bodě. Zatímco ve fyzickém prostředí pracuje stroj i člověk s odrazem prostředí a jeho převedením na vnitřní struktury, respektive symboly (nepracujeme tedy přímo se jsovcy/objekty), tak ve virtu-

⁵⁰Procesorem je dána velikost adresového prostoru. Teoreticky lze u 32-bitových procesorů mít adresní prostor $2^{32} = 4$ GB, ovšem u 64-bitových procesorů je to obrovských $2^{64} = 16$ EB (exabytů), ovšem současné procesory to mají omezené (oseknuté) na stále nepředstavitelných 2^{48} .

⁵¹Toto dělení je výstižné, ale velmi zevrubné a našlo by se dost námitek, že ne vždy to tak funguje (kupříkladu swapovací prostor na pevném disku, který program využívá). Po svém ukončení však ztrácí program možnost se k nim v budoucnu (při dalším spuštění) vrátit. Záleží tedy na samotném programu (jak pracuje se svými strukturami) a také na správě operačního systému.

álním naopak pracujeme s celými objekty⁵² (především u objektově orientovaných jazyků), neboť objekty jsou již v požadované formě, se kterou dokáže program pracovat.

V souvislosti s porovnáním paměti výpočetního stroje a člověka shledávám další možnosti rozvoje ve struktuře a ukládání dat na softwarové úrovni, neboť technologických či principiálně jiných možností pro ukládání dat v současnosti nemáme (poměr cena/výkon).⁵³ Softwarově tak budeme moci modelovat podobné chování, které u člověka může být spojeno s jiným způsobem ukládání a práce u krátkodobé⁵⁴ a dlouhodobé paměti, dále také z hlediska umístění a druhu paměti, kdy degenerativní poruchy paměti nevyřadí naráz veškeré schopnosti pamatovat si. Samozřejmě, že při přejímání obdobných principů fungování jsou naším cílem jen ty žádoucí a řešící určitý problém.

Člověka determinuje ještě jedna důležitá schopnost, bez níž by nevzešla dnešní civilizace (včetně vědy) a ani složité sociální chování – schopnost učení. U člověka jde opět o složitou činnost a souhru mnoha mozkových center, které využívají stimuly získané prostřednictvím receptorů a dále s nimi pracují, ukládají do paměti nebo pozměňují souvislosti, a tím také konstituují náš vnitřní odraz světa. Dále musíme být schopni interakce a především odpovídající zpětné reakce do prostředí kolem nás. Schopnost učení nalezneme už u jednoduchých živočichů⁵⁵ a na různé úrovni i skrze celou evoluční hierarchii druhů.⁵⁶ Z tohoto hlediska však v aplikacích nedokážeme využít veškeré možnosti učení pozorované u člověka. V souvislosti s umělou inteligencí nás budou zajímat jen některé, které dokážeme zjednodušeně modelovat či využít jim podobných principů.

Rozhodl jsem se zde představit v dnešní době velmi diskutovaný příklad (neurologický experiment), který může být inspirací pro umělou inteligenci. Jedná se o studii [Rizzolatti a kol., 2001] provedenou týmem kolem profesora GIACOMA RIZZOLATTIHO v roce 2001 a spojenou s objevem tzv. zrcadlových neuronů. Je to jednodu-

⁵²Samozřejmě, objekty mohou být zjednodušeným modelem z fyzického prostředí, ovšem v tomto případě se jedná pouze o virtuální prostředí, tedy tuto návaznost zanedbávám. Ovšem u strojů, které pracují ve fyzickém prostředí, je tento virtuální objekt odrazem fyzického objektu.

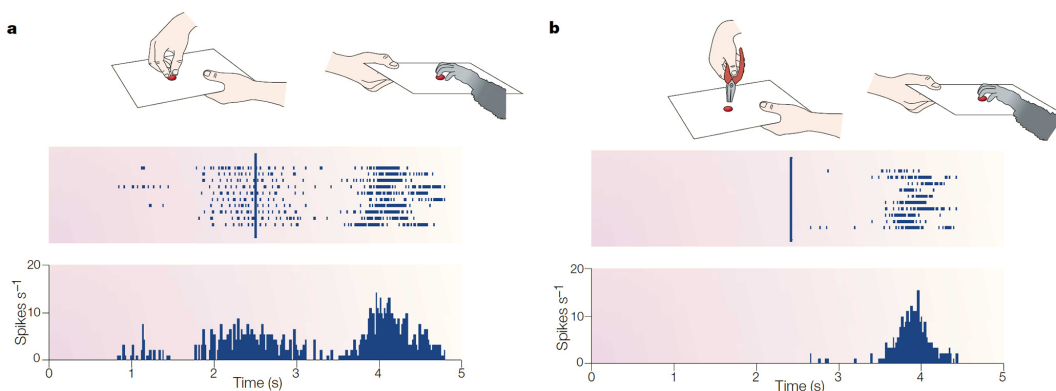
⁵³Pro lepší pochopení bych to upřesnil na představě softwarové struktury – frameworku. Jedná se o vytvoření určitého virtuálního prostředí s jasně definovaným rozhraním a operacemi, které lze volat.

⁵⁴Krátkodobá paměť má úzký vztah k pozornosti, slouží v podstatě k prezentování informací po určitou dobu tak, aby s nimi mohl mozek pracovat. Má-li se informace uložit, musí se zapojit hippokampy, a tedy dlouhodobá paměť. Stejně je tomu i u analogie v podobě počítače (operační paměť vs. pevný disk).

⁵⁵Lépe řečeno i na buněčné úrovni, už zde totiž funguje složitý mechanismus učení. [Barbieri, 2006, str. 97]

⁵⁶Viz problematika buněčné paměti v souvislosti s učením bakterií v části 2.4. Transdisciplinární přístup.

chý systém, který můžeme najít nejen u lidí, ale i u zvířat. „Zrcadlové neurony jsou součástí systému, jehož činnost je podkladem schopnosti imitovat akce i učení napodobováním.“ [Koukolík, 2006, str. 87] Při svých pokusech využívali opic, kterým byla předváděna jednoduchá činnost – vzít si rukou jídlo. V prvním případě byla ruka, která to předváděla, lidská (tedy podobná ruce opici). Když se opice dívala na danou činnost, tak mozek opice vnitřně opakoval danou činnost obdobně, jako by to dělala sama opice.⁵⁷ Pokud opice samotnou činnost prováděla, tak se činnost určitých oblastí ještě zvýšila. Když ovšem stejnou činnost – úchop jídla – předváděl člověk s kleštěmi, tak mozek opice nereagoval (kleště nezná, není to součást těla opice). Jedná se o jednoduchý způsob učení se jeden od druhého, přičemž v tomto kontextu přicházejí nové teorie o vzniku jazyka a řeči, neboť původní teorie NOAMA CHOMSKYHO se jeví jako nedostatečná. [Koukolík, 2006, str. 86] Tyto teorie včetně toho, jakým způsobem souvisí se zrcadlovými neurony, osvětlím v části nazvané Myšlení.



Obrázek 2.3.: Mozková aktivita opice při jednotlivých úkonech. Převzato z [Rizzolatti a kol., 2001].

V uvedeném příkladu studie o zrcadlových neuronech se jeví jako zajímavá myšlenka zrcadlení vlastního těla v sobě samém a využití tohoto záznamu pro vzájemné učení činností – tedy rozšíření možností interakce (například agentů). Nemusí to být tak komplexní děj jako u zmíněné opice a rozhodně se to týká systémů budoucnosti – zejména směrem k silné umělé inteligenci.

⁵⁷ K tomu se využívala metoda PET (pozitronová emisní tomografie), jejíž zjednodušený princip je založen na myšlence, že pokud je některá část mozku zatížena, tak potřebuje více kyslíku a živin. Dále se využívají radiofarmaka s krátkou dobou rozpadu. Pokud se ve správný čas provede experiment, tak nejvíce rozpadajících se radiofarmak uvidíme v zatížené části mozku. V tomto případě šlo o oblast ventrální premotorické kůry opic (oblast F5), která je stejná jako Brocova oblast BA 44 u člověka. [Koukolík, 2006, str. 87]

Schopnost učení mají umělé systémy postavené na tradiční umělé inteligenci. Jedná se o oblast strojového učení, kdy existuje několik metod, které se využívají v různých aplikacích⁵⁸ umělé inteligence. Každá z nich má svoje výhody i nevýhody, neexistuje však jedno univerzálně aplikovatelné schéma. Cílem je najít co nejlepší metodu pro určitou aplikaci – nutně se uplatňuje diverzifikace metod. Touto cestou jde i současný vývoj s přihlédnutím k prostředí, ve kterém se daný systém nachází. Pokud by se zvyšovalo množství vjemů, tak bude nutné najít vhodnou inspiraci (metodu) u živých organismů (člověka). Tím, co tedy bude do budoucna rozvíjet oblast učení, budou určitě stále složitější percepce z prostředí a nutnost jejich efektivního zpracovávání.

2.2.3. Problematika zapomínání

Se zapamatováváním se neodmyslitelně u člověka pojí i zapomínání. V současné době máme tři teorie o zapomínání, dle [Hayward, Varela, 2009, str. 128]. První teorie představuje paměť jako časem se rozkládající, obdobně jako se rozkládají organické zbytky v přírodě. Druhá předpokládá, že dříve uložené informace jsou vytlačovány novými. A konečně třetí tvrdí, že zapomínání neexistuje, jenom my, či lépe řečeno, naše vědomí k nim ztrácí přístup. Ani jednu z těchto teorií nelze uspokojivě prokázat tak, aby vyloučila zbylé dvě. V tomto ohledu jsme v začarovaném kruhu vědeckého poznávání.⁵⁹

Při tomto zamyšlení nad člověkem je zajímavé, že vzpomínky, které byly uloženy do paměti v mládí, nejsou tak náchylné k zapomínání. Naopak ty, které byly uloženy v pozdější době z hlediska ontologického vývoje jedince, jsou více náchylné. Příkladem může být postupné degenerativní onemocnění – Alzheimerova choroba. [Klíma, 2008, str. 100]

Problematiku zapomínání zatím příliš v umělých systémech neřešíme. U strojového učení sice nalezneme příklady systémů, které se učí a zároveň také zapomínají (dříve naučené příklady jsou vytlačovány novými) především v oblastech, kde se mění již naučený koncept.

Zapomínání je v sociální skupině přijímáno jako nežádoucí jev, ovšem z pohledu fungování těla organismu jde o jev velmi důležitý. Každopádně i v oblasti informatiky

⁵⁸Pro příklad uvedu: empirické učení, učení jako prohledávání (rozhodovací pravidla), učení jako aproximace (neuronové sítě).

⁵⁹Třetí zmíněná teorie je z pohledu dnešní vědy nevyvratitelná. „Pokaždé, když si vzpomeneme na něco, co jsme předtím zapomněli, počítá se to jako důkaz teorie. Když si ale na něco vzpomenout nemůžeme – dokonce i když si nevzpomeneme ani do konce svého života – můžeme pokaždé tvrdit, že jsme k dané věci prostě ztratili přístup spíše, než že by se nám z paměti zcela vytratila.“ [Hayward, Varela, 2009, str. 129]

nalezneme pokusy zavést přirozenou selekci. S těmito snahami se setkáváme především v souvislosti s problematikou vyhledávání a spravování dokumentů na internetu. V České republice je takovým projektem G-Hive⁶⁰ pro distribuované sdílení XML dokumentů – když dokument přestane být důležitý (není o něj zájem), tak začne postupně mizet ze serverů, až může zmizet zcela a zbude po něm pouze základní záznam, že existoval.

Myšlenka zapomínání je lákavá i pro budoucnost. Čím lépe budou umělé systémy jednat ve stále komplexnějším prostředí, které jim budeme odkrývat, tím více podnětů budou nuceni zpracovávat. Dnešní velikosti pevných pamětí jsou obrovské, přesto mohou u takových systémů existovat různé limitní faktory, například velikost stroje či hardwarová vyspělost. To platí i v oblasti umělé inteligence, kdy určující bude, jakým směrem se bude dále ubírat, a to především v oblastech kooperace s člověkem nebo práce ve složitém prostředí.

2.2.4. Myšlení

Myšlení – neustálý pohyb uvnitř nás, který zpracovává stimuly, komparuje je s vnitřním modelem světa, a vytváří nám tak naše vidění světa. „Lidské myšlení tak vykonává neustálý cik-cak pohyb, neustále nastavuje generované endocepty, když se projektuje mezi jsoucna, a při tomto svém pohybu mluví: To je to, co je již v endoceptivní struktuře obsaženo, nebo to je nové.“ [Pstružina, 2005, str. 63] Najdeme v tom základní operace identity a negativity.⁶¹ Díky schopnosti rozlišování (to je to a to není ono) vzniká základ pro další kognitivní procesy vnímání, poznávání, cítění, myšlení či vyjádření významu. Je též označován jako elementární znak inteligence. [Marková, 2007, str. 51] Výchozím v tomto ohledu je myšlení v opozitech, dělání rozdílů (kupříkladu teplo a chlad).

Myšlení v opozitech je pro člověka zásadní a teorie, jež na něj poukazují, nalezneme už ve starověkém Řecku. Příkladem nám může být HÉRAKLEITOS (535 - 475) a jeho teorie neustálého plynutí a změny, kde představil myšlenku protikladnosti, kterou najdeme u všeho živého (živý a mrtvý, starý a mladý, velký a malý). ARISTOTELES později pracuje ve své Metafyzice například s možností a skutečností, ovšem je zde razantní posun – nemožnost současné koexistence antinomií (protikladů). Zatímco HÉRAKLEITOS ještě uvažuje, že člověk může být zároveň mladý a starý, tak v poz-

⁶⁰Projekt na Katedře informatiky Přírodovědecké fakulty Univerzity J. E. Purkyně v Ústí nad Labem.

⁶¹Kromě těchto dvou základních operací lidského myšlení máme i další. Kupříkladu komparace, generalizace, abstrakce a další. Podrobněji jsou tyto operace rozebrány v [Pstružina, 2005].

dější době se s příchodem logiky⁶² objevuje separace. Tím se také usnadňuje kategorizace. Tento princip striktní separace dále konstituuje evropské myšlení až do 19. století. [Marková, 2007, str. 61-64] Z těchto základů logiky pak vyvstává axiomatická matematika. Tím zákony logiky ovlivňují i náš pohled na svět a na těchto základech s ním dále pracujeme.

Řecké myšlení je nám Evropanům blízké. Zajímavější proto je podívat se na stejnou problematiku antinomií skrze čínské myšlení. Typické pro toto myšlení v opozitech je dynamika, neustálý pohyb a změna oproti řeckému „statickému“ myšlení. Existují zde komplementární opozita, nikoli separovaná. Jedno doplňuje druhé a sama existovat nemohou, klasickým příkladem je symbol Jin-Jang. „V teoriích starověké Číny byly dvě komponenty páru opozit navzájem závislé, ve věčném cyklu vlnového pohybu. [...] Nikdy nebylo možno mít buď jedno nebo druhé, protože oboje koexistovalo v kontextu pohybu a neustálé změny.“ [Marková, 2007, str. 60] Zdejší matematika neměla jako výchozí bod axiomy, přesto dosáhla pozoruhodných výsledků směrem k pragmatické účelnosti. [Marková, 2007, str. 66]

V úvodu této části jsem rychle tematicky přešel od rozlišovací schopnosti, která je základem dalších poznávacích procesů, k myšlení v opozitech, jež z ní vychází, a dále k rozdílu v jeho chápání různými národy starověku. Cílem bylo ukázat výchozí operaci lidského myšlení v kontextu dalšího rozvoje jednotlivých kultur. Nejenom, že každý jednatel uchopuje svět trochu jinak, ale i jednotlivé kultury dostatečně vzdálené se mohou z hlediska myšlení (přístupu ke světu) vyvíjet jinak. I takový jednoduchý rozdíl, jako je chápání opozit, se ve výsledku odráží i v dalších lidských činnostech – lékařství, matematice, stavebnictví a jiných.

S tímto postojem k chápání myšlení souvisí také naše řeč, protože myšlení využívá na vyšší (vědomé) úrovni naší vnitřní řeč. Nalezneme zde i rozdílné praktiky při učení klasické řeckořímské a hindské gramatiky. Zatímco řeckořímská gramatika je postavena na statických pravidlech s cílem naučit snadno a dobře jazyk, tak hindská gramatika se neučí explicitně, „ale objevuje gramatický řád spolu se čtenářem, poznává jazyk spolu s ním. Nevšímá si definic a její hlavní zájem není zmrazit, co je dynamické v komunitě domácích mluvčích, ale dojít smyslu a významu prostřednictvím dynamické intersubjektivit.“ [Marková, 2007, str. 68] Tedy i různé aspekty myšlení se odráží v dalších činnostech s tím spojených, v tomto případě k výuce gramatiky.

Dostáváme se tak do problematiky řeči včetně té mentální (vnitřní). Opět se vrá-

⁶²Zákon o nonkontradikci (neprotiřečení).

tím k NOAMU CHOMSKYMU, který přichází s myšlenkou o vrozených pravidlech univerzální gramatiky, neboť pokud by se mělo malé dítě učit pouze na základě řeči svých rodičů, tak je nepravděpodobné rychlé osvojení pravidel pozorované u dětí. To naznačuje i studie z roku 2003 pod vedením dr. MARIACRISTINY MUSSO, kde zjišťovali, jak bude zatěžována Brocova oblast při učení gramatických pravidel italštiny, japonštiny a syntetického nereálného jazyka. Výsledkem bylo, že daná oblast byla aktivní pouze, když se subjekt učil reálný jazyk. Na umělá nereálná gramatická pravidla, ač využívala italská či japonská slova ze slovníku, Brocova oblast nereagovala. [Koukolík, 2006, str. 86] Odlišný úhel pohledu od CHOMSKYHO přinesl SEARLE, který říká, že za tím nemusíme hned hledat komplexní systém pravidel, ale spíše „předpokládat, že fyziologická struktura mozku určuje možné gramatiky bez meziroviny pravidel nebo teorií.“ [Searle, 1994, str. 55] V poslední době se množí námitky⁶³ k Chomskyho teorii, především ve smyslu její redukce pouze na gramatiku.⁶⁴ Z těchto důvodů se problematika řeči rozšiřuje i směrem k sémantice a fonologii, což jsou další dva subsystémy, které se paralelně se syntaktickým podílejí na vzniku řeči.

Brocova oblast je kromě řeči důležitá i z hlediska dříve zmiňovaných zrcadlových neuronů v systému učení napodobováním. Zřejmá spojitost je nejenom v tom, že ona oblast je aktivnější, pokud se učíme imitací, nýbrž i v souvislosti s (pojmovou) řečí, která zde vzniká a která v sobě otiskává vnější jsoucna. S řečí je spojena, alespoň částečně, také myšlenková práce se světem. Řeč nám zprostředkovává i věci za obzorem, které v té chvíli nevnímáme, ale v myšlenkách s nimi pracujeme. Výchozí je přiřazení pojmů⁶⁵ k jsoucnům, a právě díky tomu s nimi můžeme pracovat a nacházet nové zákonitosti či analogie i mimo prostor a čas, ve kterém se nalézáme.⁶⁶ Učení je pak důležitým pilířem lidského myšlení, neboť nabízí novou kvalitu oproti učení jiných savců právě díky spojení s řečí. V řeči odrážíme fungování části světa, který vnímáme. Stejně je tomu také u abstraktních konstruktů využívajících svou „vlastní řeč“ k popisu skutečnosti (například fyzika využívá matematický aparát).

⁶³Jde o problémy, které nezapadají do Chomskyho hypotézy nebo ji vyvracují. Více viz [Koukolík, 2006, str. 86]

⁶⁴Povšimněme si úzkého spojení i s otázkou: Mohou počítače myslet? Zatímco od 50. do 80. let 20. století s velkým rozvojem výpočetní techniky přicházel i velký optimismus, že to je jenom otázkou času, kdy budou stroje myslet, tak optimismus ochladl, když SEARLE zveřejnil svou námitku vědomí, kde upozorňoval, že mysl není jenom syntax, nýbrž má i sémantiku. Obdobně to je u Chomskyho teorie, kde vychází jen z gramatiky (morfologie a syntax), a teprve nyní také směřuje k sémantice. Ale jak jsem již naznačil, v tomto pohledu nebyl CHOMSKY rozhodně sám – viz část Člověk jako stroj.

⁶⁵Nejenom pojmů řečových, ale také mentálních, například pocit.

⁶⁶Zvířata se také dokážou chovat inteligentně, ale nemají (vnitřní) řeč stejně jako my. Zpravidla pracují pouze s věcmi, jež bezprostředně vnímají, přičemž nemají pojmy, a pracují tedy s představami (pocity) včetně těch, které nedokážeme řečí vyjádřit. Například pokud vidí myš sýr, tak jej může spojit s pocitem nasycení. [Krámek ed., 2009, str. 57]

Touto krátkou exkurzí od základních myšlenkových procesů k řeči jsem chtěl ukázat, že řeč je úzce spojena s lidskou inteligencí a tvořivostí. Do jisté míry to naznačuje, co dnešním inteligentním systémům v tomto ohledu chybí směrem k lidské inteligenci. U těchto systémů ovšem o myšlení ve stejném kontextu jako u člověka mluvit nemůžeme. Vhodnější termín je zpracování. A i zde si musíme dávat pozor na to, že zpracování informací člověkem a výpočetním systémem je zcela odlišné, avšak často zaměňované. [Searle, 1994, str. 53]

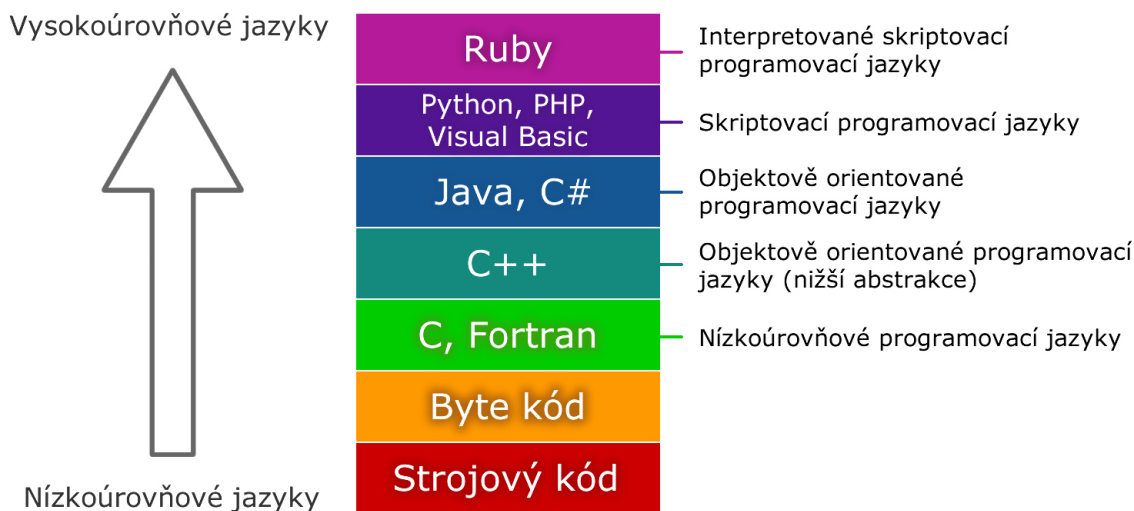
V tomto pojednání nechci představit myšlení, či lépe řečeno zpracovávání u stroje, nýbrž využít spojení „myšlení a stroj“ ve smyslu toho, jak se lidské myšlení a chápání světa odrazilo v návrzích výpočetních strojů a jejich softwarových programů, a tedy i aplikací umělé inteligence.

Na začátku jsem představil základní operace identity a negativity. Bylo by chybné představit stejnou analogii u strojů ve formě dvojkového kódu, který je pouze prostředkem. Avšak analogii najdeme ve formě programování, nejlépe v klasickém imperativním programování.⁶⁷ Jde zde o styl, jakým programátor píše kód, a tedy i o skutečnost, jak nad tím přemýšlí: jakou navrhne strukturu a členění programu, jak určitou funkci programu vyjádří v daném jazyce.⁶⁸ Imperativní jazyky se svou konstrukcí velmi podobají právě separovaným antinomiím – když nastane A, udělej toto, pokud nikoli, tak udělej tamto (nic mezi tím neexistuje). Díky jednoduchým konstrukcím můžeme vytvářet složité programy, které překvapují svou složitou funkcionalitou.

Se stylem programování se pojí i překladače, jež jsou optimalizovány pro zpracovávání kódů v určité formě, a také frameworky, nad kterými staví vysokoúrovňové programovací jazyky. Pro představu zde uvádím obrázek se zběžným přehledem programovacích jazyků. Čím vyšší úroveň jazyka, tím silnější abstrakce od detailů (návrhu a hardwarového fungování) počítače.

⁶⁷Třeba funkcionální programovací jazyky jako příklad nejsou až tak vhodné, neboť vycházejí z matematiky. Jsou tedy abstraktnější než jazyky imperativní. Navíc i tento typ programovacích jazyků (např. Scheme nebo Haskell) využívá různých frameworků vystavěných na imperativních jazycích, neboť pokročilejší funkcionální konstrukce (vycházející z matematiky) se musejí převést na imperativní – vhodnější pro zpracování strojem.

⁶⁸Je diametrálně odlišné myšlení a přístup k řešení u imperativního a funkcionálního programování (pokud jde funkci vůbec v obou programech vytvořit). Je to jiný styl myšlení o problému. Lidé, kteří umí programovat v C, JAVĚ či C++ (imperativní jazyky) velmi obtížně přecházejí na práci v jazycích funkcionálních. Dokonce již existují tzv. multiparadigmatické jazyky, které umožňují jak některé konstrukce imperativní, tak funkcionální.



Obrázek 2.4.: Přehled různých úrovní programovacích jazyků.

To, jakým způsobem myslíme, vtiskáváme do softwarové výbavy stroje, jež řídí jeho činnost (a nikoli tedy to, co stojí v pozadí za tímto myšlením, z něhož naše myšlení povstává). To by mohl být důležitý obrat v přístupu k umělé inteligenci. Zatímco současná umělá inteligence vychází z výsledného chování a přístupu člověka ke světu, lépe řečeno současné myšlení (v rámci západoevropské civilizace⁶⁹), nezajímá se, odkud toto myšlení, respektive inteligentní chování, povstává a kde jsou jeho základy. V tomto ohledu nás může obohatit právě kognitivní věda a výzkum mysli na hranici subjektivity. Je to jako objevit jiný pohled na problematiku: Přestaneme se na ni dívat zvnějšku a přesuneme se více dovnitř ve smyslu introspekce. Tato myšlenka je reálná pouze ve spojení s člověkem, u něhož budeme nacházet potřebnou inspiraci, stejně jako je tomu u toho, co pokládáme za inteligentní chování.

Dle mého názoru je to možnost, jak se lépe přiblížit k pokročilejším operacím lidského mozku, jako jsou generalizace nebo extrapolace na obecnější úrovni aplikovatelné v přístupech k umělé inteligenci.⁷⁰

2.2.5. Člověk jako stroj

Důležitým východiskem k úvahám o myšlení je zájem člověka o sebe sama. Oblastí, které se zabývají studiem člověka z nejrůznějších úhlů, je velmi mnoho. V tomto po-

⁶⁹Viz myšlení starověkého Řeka vs. myšlení Číňana zmíněné výše.

⁷⁰Podrobněji o operacích lidské mysli v [Pstružina, 2005, str. 46].

jednání nás však zajímá především (intelligentní) chování člověka, a proto je vhodné se zaměřit zejména na centrum kognitivních procesů – mozek.

„Sokratovský požadavek – Poznej sám sebe – již dávno překročil oblast filosofie a etiky a uskutečňuje se pomocí psychologických výzkumů.“ [Linhart, 1972] Psychologie se zabývá vnějším chováním člověka, jeho myšlením a tím, jak prožívá různé situace. Jiný pohled nabízí neurologie, která zkoumá poruchy centrální nervové soustavy zapříčiněné přímým fyzickým poškozením. Ve vši jednoduchosti zde opět nalezneme paralelu v dualistickém pohledu na oblasti, kterými se psychologie (mysl) a neurologie (fyzická podstata) zabývá. Oba pohledy dále zastřešuje mladý obor neuropsychologie a také psychiatrie, která k léčbě pacientů využívá obou přístupů a sleduje i jejich následky pro život jedince i společnost, ve které osoba žije.

Lidé rádi hledají analogie, podobné jevy či znaky, jenž jim vytváří řád. V tomto ohledu se vedle přírodních věd rozvíjela i psychologie a neurologie. Lidé pak začali objevovat jisté analogie psychologie a matematiky, respektive teorie systémů⁷¹ [Kelemen, 2010, str. 83], či neurologie a fyziky [Sacks, 2008, str. 32]. Analogie se však postupem času s vývojem počítačích strojů přesunula od matematiky a fyziky k softwaru a hardwaru.⁷² Psychologové a neurologové se uchýlili „k analogiím, metaforám a obrazům převzatých z přírodních a člověkem vytvořených jevů, které jsou lidem v dané době dostupné.“ [Hayward, Varela, 2009, str. 124] Vznikl i vědní obor, který vzešel z obdobných pozorování u živých organismů a strojů – kybernetika. WIENER a ROSENBLUTH (1900 - 1970), průkopníci v oblasti kybernetiky „boli unesení analógiami, ktoré si postupne uvedomovali napríklad pri porovnávaní elektrických obvodov a biologických organizmov.“ [Kelemen, 2010, str. 92] Vlastně lze prohlásit, že „metoda analogie je základním kybernetickým přístupem při zkoumání zákonitostí účelného řízení a sdělování v systémech.“ [Říhová, Vítek, 1989, str. 36] Není proto daleko k představám a přístupu k člověku jako ke stroji a obráceně.⁷³ Tato inspirace – hledání analo-

⁷¹Zatímco Josef Kelemen naráží na vývoj psychologie směrem k formalismu v průběhu 20. století, podobnému jako u matematiky, tak z této myšlenky dále vycházejí i další a formují ji dle své potřeby a tehdejšího vývoje vědy. V tomto případě směrem k teorii systémů: „Psychické systémy jsou dynamické a mají svou strukturu, která je určována povahou vztahů mezi jejími činnými elementy, mezi smyslovými orgány, efektory a ústředím. O chování z teoretického hlediska mluvíme tehdy, jestliže popisujeme závislost výstupů na vstupech systému.“ [Linhart, 1972, str. 14] V psychologii se v tomto ohledu rozvíjí materialisticky zaměřený a formální behaviorismus. Opět je zde ona paralela mezi člověkem a strojem.

⁷²Již v první kapitole jsem zmiňoval neurologický pohled na mozek jako na stroj, respektive později počítač. Každopádně různé analogie mozku a nějakého výdobytku techniky najdeme napříč historií. Kupříkladu pro objevitele funkce neuronu CH. S. SHERRINGTONA byl mozek telegraf, FREUD ho srovnával s hydraulickým systémem a LEIBNIZ jej přirovnával k mlýnu. [Searle, 1994, str. 47]

⁷³Pro představu uvedu, jak viděla v 60. letech 20. století kybernetika analogii člověka a stroje. Matematický stroj má pět základních složek: vstup, výstup, řadič, operační jednotku a paměť. To u člověka odpovídá aferentním drahám ze smyslových orgánů, eferentním drahám například k motorickým

gií – nás v mnohém obohatila novými myšlenkami a pokrokem v oblasti techniky, ale zároveň nás uzavírá a omezuje v širších souvislostech. Především pokud si to sami neuvědomujeme.

Podívejme se na dva zajímavé příklady z oblasti neurologie, které souvisí s kognitivními procesy u člověka. Jejich studium nám může napomoci k rozluštění, jak dané procesy v lidském těle fungují, a tak je i lépe zrcadlit v přístupech k inteligentním systémům, potažmo umělé inteligenci. Především nám ukazují, jakým způsobem se u „postíženého“ jednotlivce změnily jeho kognitivní schopnosti, a tím se mu změnilo vnímání prostředí i pohled na svět. Neboť jak bylo zmíněno již dříve, naše vnímání světa není jediné možné či zaručeně správné (plné). Stejně tak, jako se u těchto pacientů změnilo vnímání světa, bude se měnit vnímání a zpracování podnětů i u budoucí generace lidí, včetně jejich zprostředkování přes různé systémy využívající nějakou aplikaci umělé inteligence. I toto je třeba si uvědomit a mít na zřeteli při úvahách v širším kontextu.

Všechny uvedené příklady vycházejí z [Sacks, 2008], nicméně další příklady z oblasti neurologických a psychických defektů lze nalézt v [Sacks, 2009a] nebo [Sacks, 2009b].

Zpracování percepce z prostředí:

DR. P. byl muzikant a učitel na hudební škole. Postupem času ovšem přestal poznávat obličeje svých žáků, přátel i svůj vlastní. Dokonce „neviděl“ člověka, i když stál přímo před ním, dokud jej tento člověk neoslovil (nebo nevydal nějaký zvuk). Oči měl však v pořádku, v prostoru se pohyboval normálně a vyjadřoval se plynně a vybraně. Rozhodně na něm nikdo nepozoroval jakoukoli demenci. Tohoto svého problému si však nebyl vůbec vědom. Při základních testech, kdy měl rozpoznat platónská tělesa (základní geometrické útvary), si vedl bezchybně, také karty a dokonce karikatury rozpoznal přesně. Problém nastal, když měl říci, co je na fotografii, nebo pracovat s předměty kolem sebe – pokud se mu dala rukavice, dokázal ji skvěle popsat, ale chyběla integrace (úsudek) všech těchto poznatků (zrakových vjemů). To samé platilo, i pokud si zul botu – nedokázal ji pak najít na zemi vedle své nohy, dokud mu ji někdo nepodal.

svalům, řadič je obdobný mozkové kůře, která řídí nervové procesy v těle, operační jednotka odpovídá logickým funkcím probíhajícím v mozku a konečně paměť člověka ukládá informace obdobně, ovšem jiným způsobem. [Zeman, 1964, str. 73]

A přestože kybernetici jedním dechem dodávali, že stroj pouze napodobuje formální a logické operace mozku, tak se tento dovětek poněkud ztrácí v rychlém rozvoji technického pokroku. Matematický stroj by se mohl v silném slova smyslu podobat mozku až tehdy, „kdyby měl i humorální, biochemické a emocionální složky, které jsou neoddělitelně spjatý s logickými operacemi mozkové kůry člověka.“ [Zeman, 1964, str. 16]

Osoby rozpoznal, jen když byly něčím zvláštní (knírek, brada), proto také rozpoznával bezchybně karikatury. Při pozorování scén v televizi (bez zvuku) vůbec netušil, co se na obrazovce děje – nerozpoznal postavy ani nedokázal určit emoční náboj, chyběla mu empatie. Časem Sacks zjistil, že ke všemu včetně lidí přistupuje jako k předmětu, vše bylo „to“, nikoli „já“, „ty“ nebo „on“. Když vyprávěl příběh, tak byl perfektní v popisech, líčil krajinu i města. Ovšem neříkal nic o emocích jednotlivých postav nebo o jejich obličejích, a to přesto, že knihu četl ještě v době, kdy na něm nebyly pozorovány tyto odchylky. „Stejně jako počítač byl lhostejný vůči viděnému, vnějšímu světu – a co víc, konstruoval pro sebe představy předmětů úplně stejně jako počítač: podle charakteristických tvarů a schematických vztahů. Dokázal tato schémata sledovat, aniž je vnímal jako skutečné věci.“ [Sacks, 2008, str. 27] V této chvíli bychom mohli říci, že DR.P. byl na úrovni slabé umělé inteligence, kdy se pouze chová jako člověk, avšak chybí zde úsudek, jakým disponuje běžná lidská mysl. Přesto mu to nebrání být plně integrován ve společnosti (i když někdy udělá pár excesů).

Diagnóza byla vnitřní agnózie.⁷⁴ DR. P. však vůbec netušil, že přišel o pohled na svět, jaký mají ostatní lidé. Pokud je nějaká část mozku poškozena, mozek se snaží nahradit její funkci jinými částmi.⁷⁵ DR. P. začal „vidět sluchem“, identifikoval předměty podle hudby a rozeznával tak i své přátele: „To je přece Karel – znám jeho pohyby, jeho vnitřní hudbu.“ [Sacks, 2008, str. 30] Hudba se tak pro něj stala opravdu vším.⁷⁶ Měl specifické popěvky pro vlastní činnosti, jako bylo jídlo, oblékání a jiné. Dokud si zpíval, tak se dokázal oblékat, pracovat s předměty. Jakmile byl přerušen, tak všechny předměty zmizely stejně, jako když nám někdo zhasne světlo.⁷⁷

Tento příklad je zajímavý i tím, že nezapadá úplně do neurologického pohledu na svět o tom, jak fungují pochody v mozku. Neurologové vidí v abstrakci⁷⁸ výchozí bod pro kategorizaci, což je základem pro úsudek. Pointa je, že to, co ničilo DR. P. jeho úsudek, byla právě abstrakce. „Ono absurdní abstrahování, skrze něž vnímal, mu bránilo vidět vnější skutečnost a zničilo jeho schopnost úsudku.“ [Sacks, 2008, str. 31] OLIVER SACKS tak upozorňuje na zažité pohledy, především vzhledem k dílům, která napsal JOHN H. JACKSON (1835 - 1911) a KURT GOLDSTEIN (1878 - 1965). Opět je zde

⁷⁴Což je neurologická porucha poznávání.

⁷⁵A právě gyrus fusiformis (součást lobus temporales) je klíčovou součástí systému rozpoznávání tváře, zároveň František Koukolík zmiňuje tuto část mozku jako jednu z mnoha propojených styčných uzlů hudebního modulu. [Koukolík, 2010, str. 98] To nám dává tušit i ono přizpůsobení skrze hudbu.

⁷⁶O vlivu hudby na lidský mozek pojednává další kniha [Sacks, 2009b]. O neuronálních korelátech hudby se pak více dozvíte v [Koukolík, 2010, str. 89 - 108].

⁷⁷Lepším pocitovým vyjádřením, vzhledem k tomu, že šlo o činnosti, je, že někdy ztratíme myšlenku o tom, co jsme vlastně chtěli právě udělat.

⁷⁸Myšlenkový proces zajišťující rozpoznávání podstatných vlastností a vztahů, které to přináší.

možné ovlivnění těchto pánů začátkem 20. století a příklon k formalismu a kategorizaci světa. Obdobný pohled dnes máme kupříkladu i v ontologickém inženýrství, leč závěrečná aplikace probíhá ve světě námi daných pravidel (internetu), kde je ontologická kategorizace využita pro sémantické vyhledávání nebo multilinguální slovníky. Takové aplikace kategorizace však neztrácejí nic na své podstatě – řeší problém.

Vyhodnocování podnětů:

Pan MACGREGOR (93 let) byl stížen Parkinsonovou chorobou⁷⁹ a žil v Domově sv. Dunstata. Měl ovšem problém, který si neuvědomoval – nakláněl se o 20 stupňů na pravou stranu. Teprve, když viděl video své chůze, tak připustil, že s jeho pohybem není něco v pořádku, jenže sám si to neuvědomoval. Tělo má v tomto ohledu trojí jištění: zraková kontrola, vestibulární systém (nejznámější ústrojí rovnováhy, které se nachází ve vnitřním uchu) a propriocepční systém.⁸⁰ Parkinsonova choroba však narušila souhru těchto receptorů, které zajišťují správné držení těla. Pan MACGREGOR to vyřešil tak, že si na brýle přimontoval malou vodováhu – od té doby už chodil zpříma.

Ačkoli to na první pohled vypadá, že člověk při svých činnostech využívá pouze několik málo smyslů, jejichž podněty nevědomě vyhodnocuje, tak ve skutečnosti je to složitější mechanismus, který nám zajišťuje danou činnost. V tomto případě je to obdobné jako u dříve zmíněné dependability systému,⁸¹ kde máme také několik typů jištění. Řídící jednotka (dependabilního systému) pak vyhodnotí, který subsystém má pravdu, a podle toho se zachová, pokud ovšem ve výsledcích nebude nějaká shoda, tak nastává neobvyklé chování (zpravidla chybné), stejně jako u pana MACGREGORA.

Jak doufám, tak tyto příklady dostatečně ilustrují, proč je v mnohém náš pohled na člověka podobný pohledu, jaký máme na stroj, a jak jednoduché je abstrahovat až k této paralele. Hledáme neuronální (anatomické) koreláty našeho jednání. Tento čistě mechanický pohled, i když může být zavádějící, však plní v mnohém svůj účel – neurologie člověku především pomáhá. Vše má tedy dvě strany mince. Obdobný pohled na člověka nám nabízí i moderna se svým logickým myšlením, aneb ten, kdo myslí racionálně, myslí také správně, ovšem nikoli vždy jako člověk. „Vzato vážně to ovšem znamená, že myšlení už nakonec jenom kontroluje, zda vůle, nepozornost či

⁷⁹ „Přesto byl duševně cílý a smysly mu sloužily dobře; fyzicky byl zdatný a silný.“ [Sacks, 2008, str. 79]

⁸⁰ Vědomí relativní polohy trupu a končetin v prostoru, které zprostředkují nevědomě receptory ve sva-
lech, kloubech a šlachách. [Sacks, 2008, str. 80]

⁸¹ Schopnost poskytovat službu, na kterou se dá spolehnout.

netrpělivost někde nenarušuje toto vyvozování pravdy z pravdy, jinak řečeno: Rubem tohoto ideálu je myšlení, které samo nemyslí, myšlení jako dokonalý stroj.“ [Petříček, 2009, str. 28] Tím se dostáváme také k tomu, jak chápeme ideál člověka skrze jeho sociální atributy.⁸² Takový člověk je pak bohužel degradován na jednoduchého a předvídatelného agenta.

Pohled na člověka jako na stroj je dán také naším neustálým hledáním analogií a omezeným vnímáním světa skrze vnější receptory. Nejsme si vědomi, že nedokážeme uchopit celek člověka najednou. Vnímáme jen některé fenomény spojené s jeho bytím. Je to stejné jako nebýt si vědom rozdílu mezi kočárem a moderní limuzínou. V našem (zobecněném) pohledu mají oba čtyři kola a dokážou jezdit – jsou vlastně téměř totožné. Ale jsou opravdu identické? Pokud si budeme stále vědomi těchto rozdílů, tak nás to může při hledání zajímavých a užitečných analogií posunout dále.

Psychologové se již pomalu odvracejí od tohoto pohledu uplatňovaným i behaviorismem, respektive materialismem, a dívají se na člověka v nových souvislostech. Sami to nazývají „kognitivní revolucí“, kterou končí nadvláda behaviorismu a otevírají se nové možnosti směřování výzkumu. Ovšem neznamena to, že bychom měli vše, čeho jsme ve spojení s behaviorismem dokázali, odsunout zcela mimo. Naopak, jeho výsledky jsou stále důležité pro klinickou či pedagogickou psychologickou praxi. [Říčan, 2007, str. 144] Obecně lze uvést následující shrnutí: Příklon k transdisciplinárním pohledům, jež nabízí také kognitivní věda, otevírá nové možnosti nejen v oblasti zkoumání (vědy), ale i v oblasti samotného přístupu k člověku.

2.3. Zpětná reakce

Reakcí se myslí odezva na nějaký popud, v případě živých organismů se jedná o odpověď na určitý podnět. V tomto případě se zaměřím na reakce do vnějšího prostředí vzhledem k systému, o kterém budeme mluvit. Stejně, jako je zmíněné v problematice podnětů, tak i nyní je určující prostředí.

Začneme tedy reakcí ve fyzickém prostředí. Zpětná reakce do prostředí u člověka může být buď promyšlená, nebo reflexivní, záleží na druhu a intenzitě podnětu. Tomu velmi volně odpovídá deliberace a reaktivita u stroje. Zatímco reflexivní reakce nemůžeme vědomě ovlivnit a podobají se reaktivnímu chování v nové umělé inteligenci, tak vědomé reakce jsou zajímavé pro systémy tradiční umělé inteligence. Na rozdíl

⁸²Záleží samozřejmě na pohledu každého z nás, ale téměř vždy je spjat s konzumním životem. Konzumní pohled je pak násilně propagován za účelem zisku.

od nich však člověk není exaktní a neustále se přizpůsobuje, respektive kontroluje, zda v prostředí vše probíhá tak, jak předpokládá, přičemž zde existuje zpětnovazební systém.

Podívejme se na interakci člověka ve fyzickém prostředí. Když získá podněty z vnějšku a následně reaguje určitým způsobem na stimul (např. zvedne ruku), tak vždy tyto skutky předchází odraz v podobě „modelu mysli“. Když započne se zvedáním ruky, tak receptory zachycují, co se v prostředí (vnějším i vnitřním) děje, srovnává se to s kýženým výsledkem v mysli a dochází případně ke korekcím pohybu (zpravidla automaticky a nevědomě). Takovým ověřováním může být sledování činnosti očima a vyhodnocování mozkiem, jak se daří plnit cíl. Ono ověřování je pro člověka důležité především proto, že postrádá přesnost. Zatímco stroj je naprogramován – posuň rameno o 100 mm – člověk je naopak ve svém důsledku vágní, svou přesnost získává až zpětným ověřováním. [Koukolík, 2010, str. 30]

Tuto techniku neustálého ověřování nalezneme v kybernetice, konkrétně v teorii regulace, a příkladem může být termostat. [Říhová, Vítek, 1989, str. 50] Složitějším příkladem typickým pro dnešní dobu může být automobil. Máme zde různé řídicí jednotky pracující s CAN sběrnici,⁸³ které přijímají signály z desítek senzorů, přičemž jednotlivé subsystémy neustále kontrolují, zda není třeba nějaká korekce, například za účelem bezpečnějšího chování vozidla na zledovatělé ploše. Přesto se v mnoha jiných umělých systémech princip neustálého ověřování nepoužívá u většiny činností, neboť takové stroje jsou úzce profilované na nepříliš složité a měnící se prostředí – kupříkladu robot na montážní lince.⁸⁴

Ačkoli stroje mohou přímo interagovat s fyzickým prostředím, u mnohých nalezneme jinou „zprostředkovnou“ reakci. Tím prostředníkem je člověk. Cílem takového (inteligentního) systému je ovlivnit rozhodnutí člověka, který provede ve fyzickém prostředí příslušnou akci.⁸⁵ Jako ukázkou takového systému můžeme použít již zmíněný příklad s mobilním telefonem vybaveným GPS přijímačem v části nazvané Virtuální prostředí.

Přesuňme se nyní plně do virtuálního prostředí. Ve vsí jednoduchosti v tomto světě pravidel akce stíhá reakci – deterministicky.⁸⁶ V případě, že na určitou akci není

⁸³Controller Area Network neboli Lokální síť řídicích jednotek.

⁸⁴Při posunu ramena nekontroluje neustále svůj pohyb. Až ve chvíli, kdy se stane něco neočekávaného, tak zkusí program podle předem daného scénáře určitou korekci (např. opakuje akci, zvýší tlak), a pokud je neúspěšný, povolá k sobě člověka.

⁸⁵Příkladem takových systémů je automobilová navigace, expertní systém a další.

⁸⁶Ve smyslu předem daného.

možnost odpovídajícím způsobem reagovat, nastává neočekávané chování končící výjimkou. Tento obecný princip platí pro všechny programy, tedy nejen pro ty, které využívají nějakou aplikaci umělé inteligence. To je z pohledu programu na té nejnížší úrovni. Na vyšší úrovni ale tyto jednoduché, avšak složité uspořádané (akce a) reakce tvoří složitější chování, do kterého může vstupovat i člověk, jak již bylo naznačeno v části Prostředí a podněty. I na této vyšší úrovni pohledu se pohybujeme v relativně jasně omezeném světě, což znamená, že i reakce jsou snáze předvídatelné (okruh reakcí na určitý podnět ve srovnání s fyzickým prostředím je malý).

Specifikou člověka a jeho promyšlených reakcí je také (krátkodobé a dlouhodobé) plánování. Ovšem umělé deliberativní systémy tuto oblast lidské přirozenosti zvládají pouze na logické úrovni v úzké oblasti (úloze) a z pohledu člověka v krátkodobém horizontu.⁸⁷ Příkladem plánování takového systému může být šachová partie, kdy na každý tah má umělá inteligence určitou odpověď, přičemž jeho cílem je dát co nejdříve mat. Z pohledu člověka se stále jedná o krátkodobé plánování, dlouhodobé plánování je určeno komplexností prostředí a kupříkladu u sociálních vazeb může být provázáno mnoha druhy různých reakcí.⁸⁸ I v rámci krátkodobého plánování nezávisí lidské plánování jen na volbě nejlepší možnosti, ale i na dřívější zkušenosti. Zásadním rozdílem mezi plánováním umělé inteligence a člověka je tedy zaměření na logickou úroveň, ve které jsou stroje rychlejší, s menším rozsahem úlohy (kontextu) a jejího cíle. Stroj je naprogramován jít po nejlepší možné cestě za svým cílem, zatímco u člověka se v této oblasti uvažování objevují i typicky lidské atributy jako zkušenost,⁸⁹ emoce či altruismus.

Ve sféře uplatnění těchto umělých systémů nám však tento typ plánování vyhovuje vzhledem k účelu, ke kterému danou aplikaci umělé inteligence využíváme. Rozšiřuje naše možnosti směrem k rychlým exaktním závěrům, a tak nám pomáhá. Pokud bychom chtěli konstruovat autonomnější jednotky, je oblast plánování v komplexním prostředí zajímavou výzvou pro budoucnost.

⁸⁷Za krátkodobý horizont považují z hlediska člověka hodiny.

⁸⁸Včetně toho, že naše akce sledují dílčí reakce, které vedou k určitému cíli (velké reakci).

⁸⁹Včetně určité zkušenosti, která byla v souvislosti s krátkodobým plánováním představena jako ukázka chování řidiče na křižovatce (část Podněty a jejich zpracování u člověka a umělé inteligence), kdy člověk nehledí jen na daná pravidla silničního provozu, nýbrž i na své dřívější zkušenosti. V tomto smyslu se míní řídit se pravidly jako člověk, nikoli formálními postupy (též nazývané pravidly) jako stroj, viz [Searle, 1994, str. 51].

2.4. Transdisciplinární přístup

V druhé kapitole jsem důraz kladl na člověka a seznámení s jeho poznávací schopností ve srovnání se strojem. Ovšem na konci první kapitoly jsem představil svou vizi o budoucím směřování vývoje a výzkumu v oblasti umělé inteligence, kdy jako výborný zdroj inspirace může být pomezí toho, co považujeme za živé a neživé (jednoduché organizmy). Domnívám se, že by bylo vhodné alespoň nastínit jeden příklad takové inspirace s transdisciplinárním přesahem. Vhodná doba přišla až nyní, kdy jsme se seznámili s různými pohledy na kognitivní procesy v závislosti na prostředí, včetně důležité souhry umělé inteligence a člověka. Cílem této části je ukázat výhodnost transdisciplinárního přístupu kognitivní vědy (kognitivní informatiky) v oblasti umělé inteligence z hlediska inspirace pro její aplikace.

V současné době, zjednodušeně řečeno, rozdělujeme chování agentů (systému) na reaktivní nebo deliberativní.⁹⁰ Pro reaktivní chování nám byly předlohou jednoduché formy života a u deliberativního chování zase uvažující člověk či přístup tradiční umělé inteligence, která se snaží dosáhnout podobného chování, jako má při určité činnosti člověk.

Posuňme se směrem k inspiraci jednoduchými formami života, jež využívá také oblast umělého života. V souvislosti s umělým životem se hovoří také o problematice minimálního života, který v současné době reprezentují kupříkladu některé bakteroidy. [Wiedermann, 2004] Systémy umělého života se zaměřují především na výpočetní schopnosti živých organizmů, které přejímají. Bakteroid však využívá jak výpočetní, tak nevýpočetní mechanismy a důležitější je zde biologický pohled.⁹¹ Splňuje základní definici minimálního života: autonomnost, reprodukci a schopnost vývoje. Pokročilejší je „motilní bakteroid vykazující známky inteligentního chování – pohybuje se v směru gradientu stravy.“ [Wiedermann, 2004] Čímž jsme se vrátili zpět k reaktivním agentům v umělé inteligenci, což potvrzuje oborovou blízkost umělého života a nové umělé inteligence.

Pokud budeme dále zkoumat a modelovat chování buňky⁹² (v tomto případě z hlediska sémantické biologie), získáme zajímavou inspiraci pro systémy umělé inteligence, respektive umělého života. Jedná se o to, jak buňky reagují na signály. Dlouho

⁹⁰Oba druhy agentů sdružuje v multiagentních systémech distribuovaná umělá inteligence.

⁹¹„Příznačné pro teoretické studium všech zmíněných modelů je odklon od původní biologické motivace v okamžiku, kdy se ukazuje jejich praktický význam.“ [Wiedermann, 2004] Příkladem mohou být P-systémy.

⁹²Záleží vlastně jen, jak se na danou buňku budeme dívat a jak složitý model buňky či jejího chování připravíme. Viz zmiňovaný perceptron versus projekt „Blue brain“.

se myslelo, že jde o jednoduché reaktivní chování – určitý signál může vyvolat pouze určitou reakci.⁹³ Vycházelo se z toho, že buňka sama o sobě je relativně primitivním stavebním kamenem celého organismu.

Představím zde příklad života buňky. Začněme u kmenové buňky, která není diferencována a specializuje se na určitý buněčný typ dle potřeby. Řekněme, že její specializací (přeměnou) vznikne svalová buňka. Buňka je schopna reagovat na různé vnější signály, kdy využívá tzv. prvotní posly (zachytí se na buněčné membráně), kteří aktivují tzv. druhotné posly, což jsou vnitřní signály, které dosahují genů a zahájí příslušnou reakci. Tento postup je označován jako signální transdukce. Problémem je, že reakce není explicitní, jak by se dalo očekávat, tedy že určitý prvotní posel vyvolá vždy stejný soubor druhotných (vnitřních) poslů, jež zahájí vždy stejnou reakci na genu. Tak tomu ale není, neboť „různé signály mohou vést ke stejným výsledkům a stejné signály mohou vést k různým výsledkům.“ [Barbieri, 2006, str. 94] Jako příklad si můžeme uvést signál (acetylcholin), který vysílají nervy svalům. U kosterního svalstva vede tento signál ke stahům, u srdečního svalstva k relaxaci a jiné svaly jsou vůči nim netečné. [Barbieri, 2006, str. 93]

Co to ale znamená? Pro biology to byl oříšek, se kterým se vyrovnali tak, že vytvořili teorii buněčné paměti a organických kódů. Buněčná paměť je známá i laikovi a můžeme si ji představit na příkladu bakterií a antibiotik. Problémem současných způsobů léčby antibiotiky je v tom, že po léčbě se nepodaří zničit úplně celou populaci bakterií. Bakterie, které přežily, jsou vůči stejnému antibiotiku resistantnější – přizpůsobily se tehdy nehostinnému prostředí a pamatují si to. Pokud se dostane část bakterií do nového prostředí (člověka) a tento člověk již má nějaké bakterie stejného druhu v těle, potom dokážou „nově příchozí“ bakterie předat svoji zkušenost (konjugací) i původním bakteriím, a tím ještě více znesnadnit účinnost dříve podaných antibiotik.⁹⁴ Takto to funguje u jednobuněčného organismu, ovšem u mnohobuněčných organismů (člověka) funguje buněčná paměť u jednotlivých buněk zajímavěji. „Jejich chování ovládá nejenom genom, ale také jejich historie. [...] Během embryonálního vývoje se buňky nesmějí jenom odlišit, nýbrž musí rovněž odlišné vytrvat. [...] Rozdíly se udržují, protože si buňky nějak pamatují účinky předešlých vlivů a předávají je svým potomkům. [...] Buněčná paměť je podstatná pro vývoj i údržbu složitých vzorů

⁹³Jednoduchý determinismus.

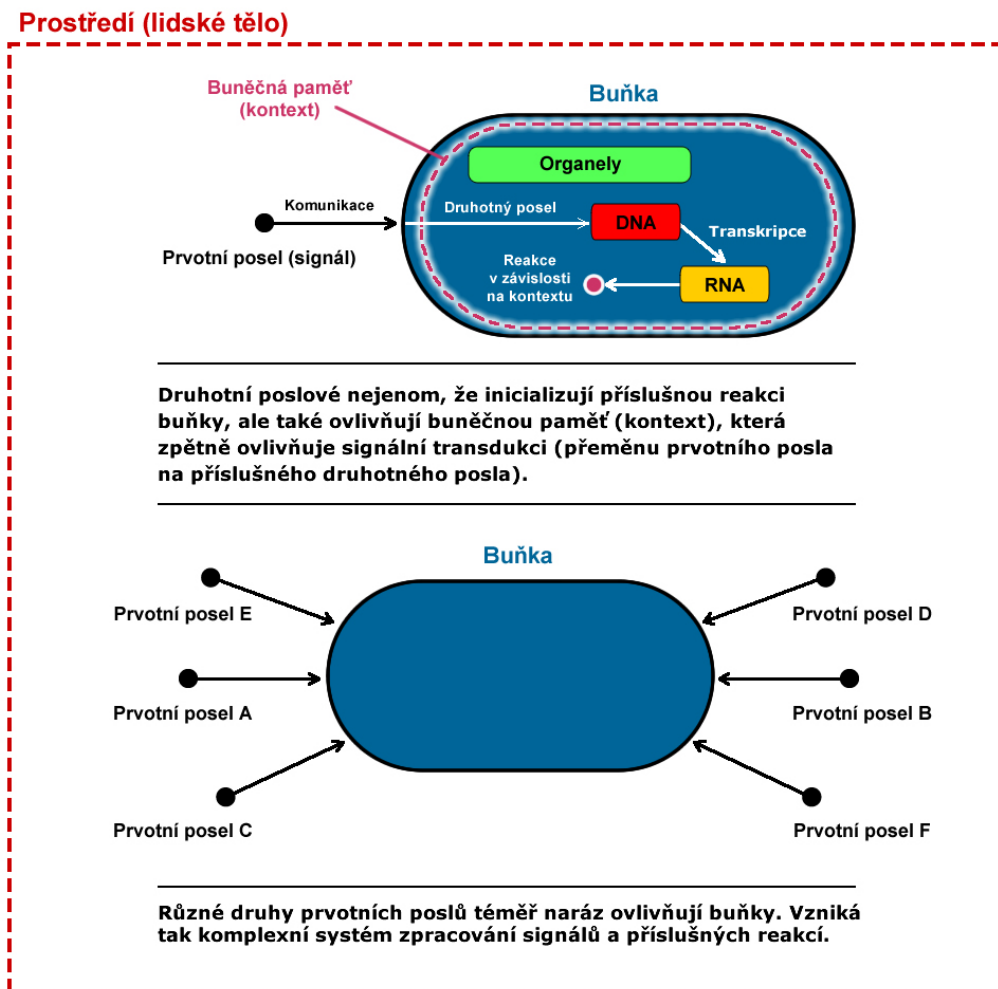
⁹⁴Mimochodem, i proto lékaři upozorňují pacienty, aby dobrali celé balení antibiotik, i když se cítí bez potíží. Také to generuje problémy ve spojitosti s úplnou resistencí některých bakterií na antibiotika, která se dříve běžně podávala, ale nyní jsou již neúčinná. Proto je třeba neustálý výzkum v této oblasti léčiv.

specializace.“ [Barbieri, 2006, str. 97, původní citace Alberts (1989) v knize Molekulární biologie buňky] Nejhmotatelnější důkaz, že buněčná paměť existuje, najdeme u člověka, který je složen z diferenciovaných buněk, přičemž jejich paměť je trvalá a stabilní po celý svůj život, život potomků, respektive život člověka.

Dalším pojmem je organický kód, který si lze představit jako soubor pravidel, která vytvářejí vztah mezi dvěma nezávislými světy. Informační struktury, jež vznikly díky pravidlům tvoření těchto struktur z individuálních znaků, přidávají význam v závislosti na interpretaci příslušné buňky.

Vraťme se nyní zpět k příkladu svalové buňky. Ta si pamatuje, že je svalovou buňkou, a dokonce díky interakci s okolním prostředím⁹⁵ v průběhu dosavadního života zná svůj aktuální kontext, lépe řečeno – uvnitř buňky vznikne určitý paměťový záznam, kterému můžeme říkat vnitřní kontext buňky. Ten se neustále mění podle toho, kde se buňka nachází (jaké signály přijímá). Buněčná paměť ovlivňuje, jaký význam pak buňka dá určitému signálu, čímž vybírá vždy jen určitá pravidla z organického kódu. „Účinky, které vnější signály na buňku mají, tedy nezávisí na energii a informaci, jež nesou, nýbrž jedinečně na významech, které jim buňky poskytují na základě pravidel, jež lze označit jako kódy pro přenos signálů. [...] Stručně řečeno, experimentální výsledky prokazují mimo pochybnost, že zevní signály instruktivní efekt postrádají. Buňky jich využívají k interpretaci tohoto světa, nikoli k jeho využívání.“ [Barbieri, 2006, str. 93]

⁹⁵Interakci s okolím si lze představit jako putování buňky či její části po jejím vzniku po těle, což se otiskává v její paměti. Představme si to na příkladu neuronu, který má inervovat určitý orgán. Po své specializaci (vzniku) se přemístí na své konečné místo (to má též určeno v buněčné paměti) a začne vysílat výběžek (neurit/axon), který hledá orgán, který má inervovat. Jak ho ale v tom obrovském těle najde? Orgány vysílají nervové růstové faktory, díky nimž axon najde příslušný orgán, pokud ovšem na tento druh molekul (signál) nenarazí, spáchá po určitém času, který má k inervaci, svoji sebevraždu (apoptóza). I z těchto důvodů jich vzniká nadbytek a přežijí jen ty, které se dostanou k cíli (nepotřebují žádné složité údaje o své poloze). Obdobné je to i u migrace buňky v těle, která reaguje na signální molekuly a podle toho mění i svůj vnitřní kontext jak interpretovat nejenom tyto signály. [Barbieri, 2006, str. 100 - 101]



Obrázek 2.5.: Schématický náčrt interakce buňky s prostředím.

Když si představíme takového agenta, jenž by se choval podle těchto principů ze světa buněk, tak kam jej zařadíme? Někdo by mohl říci, že to je deliberativní agent: „Je zde jasně vidět složitý kognitivní systém, využívá přece paměť, kde má i svůj vnitřní model prostředí. Určitě ta buňka díky své specializaci má i svůj cíl. Dokonce jsou zde náznaky komunikace mezi buňkou a orgánem (tělem či prostředím kolem) skrze signály z prostředí.“ Zastánce nové umělé inteligence však může kontrovat, že to je evidentně reaktivní agent: „Může to být obdoba subsumpční architektury, nehledejme v tom nutně složitosti. Záleží jen, jaká reakce má přednost před jakou, což je dáno specializací buňky. Navíc, inteligentní chování agenta je determinováno prostředím, ve kterém se nachází.“ Stoupenec umělého života ovšem bude oběma oponovat, že

se zcela určitě jedná o automat. „Výše popsané chování se podobá automatu vybavenému pamětí, který má k dispozici sadu pravidel.“

Klíčovou otázkou je, jak „hardwarově“ funguje buněčná paměť? To je otázka, na kterou nelze s určitostí odpovědět. Osobně si myslím, že tento příklad, který by si dozajista zasloužil formálnější zápis a případně i nějaký zkušební model,⁹⁶ je někde napůl cesty, jak to často bývá. Každopádně bych jej zařadil spíše k té reaktivnější větvi směrem k umělému životu, i když to není zcela ten typ reakce, který bychom očekávali – neustále se mění, existuje zde cíl, jež je určen prostředím. Ovšem dokud nebudeme podrobněji znát mechanismy fungování buněčné paměti (forma ukládání a vybavování), nemůžeme zaujmout jasné stanovisko. Nyní pozorujeme pouze její důsledky.

Můžeme však odhlédnout od klasických přístupů a zkusit se na buňku podívat jako na enaktivní systém. Je třeba odhlédnout od hranic systému (buňky), přičemž percepce a aktivita systému jsou tak propojeny. [Havel, Mitášová, 2009] Poznávací schopnost systému je součástí prostředí, které svým konáním ovlivňuje, a tím ovlivňuje i sám sebe. „Samotný proces poznávání se tak stává nejen procesem aktivního působení na vnější prostředí, ale zároveň i aktivním sebevytvářením, sebemodifikací.“ [Burian, 2005] Ačkoliv enaktivní přístup je uplatňován především v souvislosti s lidskou kognicí, myslím si, že i v tomto případě je více než vhodný. A právě skrze tento nový pohled se otevírají další možnosti tvorby a zajímavých přístupů k umělé inteligenci, respektive umělému životu.

Tento příklad měl také ilustrovat důležitost kognitivních procesů, které mají zásadní vliv na chování, a to ať už se jedná o živou buňku nebo o agenta v rámci umělé inteligence. Záleží tedy jen na lidech, kteří vytvářejí takové inteligentní systémy, aby jim vetknuli tu nejvhodnější sadu receptorů a případně dalších možností, jak tyto signály zpracovat (interpretovat). V oblasti výzkumu a vývoje takových systémů pak záleží jen na hloubce našeho pohledu, který zohledníme dále při vytváření modelů. Možná i tento příklad je jiný pouze z hlediska podrobnosti modelu, jež popisuje toto chování.⁹⁷

⁹⁶To už by bylo téma na samostatnou práci.

⁹⁷Další zajímavé poznatky z oblasti minimálního života lze najít v knize [Kauffman, 2004] nebo v publikacích ANTONA MARKOŠE.

3. Pilotní projekty inteligentních dopravních systémů v ČR

V následující kapitole se nechci zabývat jen určitou aplikací umělé inteligence na teoretické úrovni. Dle mého názoru v kontextu, který byl v předchozích částech naznačen, by to bylo neuspokojivé. Knih, článků a studií zabývajících se zpravidla teoreticky vysvětlenými principy fungování a možném nasazení takových systémů je mnoho. Vhodnější by bylo podívat se přímo na reálné (zrealizované) nasazení, a to nejenom ve smyslu užití určitého přístupu umělé inteligence, ale i v náhledu na inteligentní systém jako na celek, včetně jeho propojení a interakci s uživatelem – tedy z pohledu, který zde byl představen v souvislosti s cílem umělé inteligence, respektive inteligentních systémů: pomáhat člověku a rozšiřovat tak jeho možnosti v celé šíři jeho činností.

Zaměřím se na inteligentní dopravní systémy¹ neboli dopravní telematiku. Cílem těchto systémů je zvýšení přepravního výkonu, efektivity dopravy, komfortu dopravy a vyšší bezpečnost dopravy. [ITS, 2005] Za tímto účelem jsem si vybral dva pilotní projekty z grantu Ministerstva dopravy a spojů ČR (MDS 802/210/112) s názvem „Účast České republiky v projektu GALILEO“, v rámci něhož se řešilo dohromady pět pilotních projektů [Šunkevič, 2007]:

1. Experimentální přijímač GNSS² – Fakulta elektrotechnická ČVUT v Praze
2. Řízení a zabezpečení železniční dopravy na nekoridorových tratích s využitím družicové navigace – AŽD Praha s.r.o.
3. **Monitorování a řízení pohybu pohyblivých objektů po pohybové ploše letiště pomocí GNSS** – Fakulta dopravní ČVUT v Praze
4. Optimalizace řízení silniční dopravy využitím družicových systémů – Eltodo EG, a.s.
5. **Informační systém pro přepravu nebezpečných věcí využívající systém GNSS** – Fakulta dopravní ČVUT v Praze

¹Intelligent Transport Systems, zkratka ITS.

²Zkratka za Global Navigation Satellite System.

Vybral jsem dva pilotní projekty, na kterých pracovala Fakulta dopravní ČVUT v Praze.³ V následujícím textu se zaměřím na základní popis a cíle projektů, budu je hodnotit z hlediska kognitivních procesů a interakce s člověkem.

Dnešní inteligentní systémy nemusejí nutně používat nejnovější výdobytky v oblastech informatiky nebo informačních technologií. Důležitější je onen cíl, který se snaží naplnit, tedy prakticky pomoci v dané oblasti člověku, rozšířit jeho možnosti či zvýšit efektivitu práce.

3.1. Monitorování a řízení pohybu pohyblivých objektů po pohybové ploše letiště pomocí GNSS

3.1.1. Základní informace

Projekt s pracovním názvem CaMNA⁴ byl řešen mezi roky 2001-2006. Cílem tohoto projektu bylo ověřit v reálných podmínkách letiště a jeho provozu fungování systému, který by usnadňoval práci operátorům při řízení leteckého provozu. Konkrétně jde o zvýšení bezpečnosti a plynulosti na provozních plochách letiště, pomocí sledování a řízení pohybu vozidel po letištní ploše.

Ačkoli byl projekt koncipován co nejobecněji (měl být univerzální z hlediska budoucího výběru technologií), přesto musel využívat některé z dostupných řešení. Cílem projektu bylo samozřejmě také doporučit vhodné technologie. V tomto případě se jednalo o určování polohy pomocí amerického systému GPS a OBU jednotky⁵ ve vozidlech vybavené technologií WiMAX⁶ pro přenos dat při pohybu.

Architekturu systému si lze dobře představit z blokového schématu systému v příloze A této práce. Přijímač GPS signálu zjistí polohu vozidla na letištní ploše a zašle ji pomocí komunikační technologie WiMAX serveru CaMNA, který údaje rozpošle i ostatním vozidlům. Údaje o poloze vozidel jsou okamžitě k dispozici nejen posádkám vozidel, ale také řízení letového provozu (dispečer) a systému A-SMGCS.⁷ Ze softwarového vybavení je třeba zmínit upravený navigační systém Dynavix, který je uzpůsoben pro využití jak na straně klienta, tak na straně dispečerského stanoviště.

³Jedná se o projekty popsány v závěrečných zprávách [Svítek a kol., 2006a] a [Svítek a kol., 2006b].

⁴Zkratka za Car Movement on the Airport.

⁵Jedná se o tzv. carPC, neboli počítač uložený ve vozidle spojený s displejem, který je umístěn v zorném poli řidiče. Zkratka znamená On Board Unit.

⁶Zkratka za Worldwide Interoperability for Microwave Access.

⁷Systém zajišťuje informace o poloze a identifikaci letadla či vozidla. Předchází včasným varováním možnosti kolize. Uvedená zkratka znamená Advanced Surface Movement Guidance and Control System.

V příloze B naleznete fotografie z dispečerského stanoviště řízení leteckého provozu a vybavení vozidel pohybujících se po letištní ploše.

3.1.2. Zhodnocení projektu z hlediska kognitivních procesů

Při pohledu na fungování tohoto dopravního systému můžeme uplatnit dvě různé hlediska. Buď budeme popisovat systém celistvě, to znamená, že se budeme dívat na letiště jako na uzavřené prostředí,⁸ kde jednotlivá vozidla jsou agenti prvního typu a letadla agenti druhého typu. Oba agenti jsou centrálně řízeni (dispečerským stanovištěm), ovšem vozidla mají tu výhodu, že mají větší samostatnost, což znamená, že mají stejné informace, které má k dispozici řízení leteckého provozu.

V tomto případě úlohou dopravního systému ve vozidle, které se pohybuje po letištní ploše, je předzpracovat vstupy z různých senzorů a ve srozumitelné podobě (vizuální a zvukové) je předložit řidiči obdobně, jako expertní systémy představí svůj závěr a je jen na uživateli, zda jej využije nebo ne. Letadla se plně řídí pokyny z řídicí věže, přičemž se jejich poloha na letištní ploše monitoruje pomocí systému A-SMGCS. Cílem agentů je, aby se nedostaly do kritické vzdálenosti vůči ostatním agentům.

Druhou úroveň pohledu je pohled na jednotlivá vozidla nebo dispečerské centrum a jejich možnosti, které jsou dány signály získanými ze senzorů. V případě vozidla získáváme signál přímo z GPS přijímače a nepřímo ze serveru CaMNA, který poskytuje informace o poloze, které mu zaslaly ostatní vozidla a výstup ze systému A-SMGCS o poloze letadel či dalších objektů na ploše letiště. Tyto signály jsou zpracovány navigačním systémem a vhodně prezentovány řidiči.

Největším problémem celého systému je aktuálnost určených poloh jednotlivých objektů. První zpoždění vzniká u GPS přijímače, který aktualizuje polohu přibližně každých 100 ms. Navíc daná informace o poloze je nepřesná, což závisí na konstelaci satelitů, které jsou v dané zeměpisné šířce viditelné. Pokud je navíc vozidlo v pohybu, nepřesnost se zvyšuje. Jestliže to shrneme, tak pokud už jednou za 100 ms získáme pozici vozidla, může se tato poloha lišit od reálné polohy v řádech metrů, což je dáno nepřesností GPS. Dalším problémem je přenos polohy vozidla na server CaMNA. Čím rychleji vozidlo jede, tím menší časový úsek je nutný pro přenos polohy, tak aby byla dostatečně aktuální. „Data starší než 200 ms o poloze objektu pohybujícího se rychlostí $120 \frac{km}{h}$ již nelze označit jako aktuální.“ [Svítek a kol., 2006b, str. 10] Právě z těchto důvodů bylo jako nejlepší řešení pro přenos údajů o poloze zvoleno

⁸Inteligentní systém pracuje ve virtuálním světě, který je zjednodušeným odrazem toho fyzického.

využití technologie WiMAX.⁹ Další zpoždění nastává při distribuci dat ostatním vozidlům. Tento problém lze částečně řešit zvýšením kritické vzdálenosti vůči ostatním agentům.

Rychlost příjmu a zpracování signálů je zásadní pro správný běh systému. Kromě samotného umělého systému je zde i další důležitý činitel, který určuje výsledné reakce jednotlivých agentů – člověk. Přičemž i zpracování vizuálních nebo sluchových podnětů člověkem přináší další zpoždění, i když časový úsek, za který si člověk uvědomí situaci, je kratší. Na dispečerském stanovišti je situace obdobná, ovšem zde je zpoždění větší, neboť povely agentům (řidičům a pilotům) jsou předávány vysílačkou v přirozeném jazyce.

Dalším problémem je spolehlivost systému, a to jak na softwarové, tak na hardwarové úrovni. Pokud dojde k selhání, může to mít katastrofální následky. Nalezneme zde alespoň částečnou redundanci u dispečerského stanoviště, kde má obsluha k dispozici i výstup z nezávislého systému A-SMGCS. Pokud tedy obsluha uvidí, že se objekt na jednom monitoru pohybuje a na druhém nikoli, nebo se pohybuje jiným směrem, je třeba reagovat. Varování by bylo vhodné zpracovat i na softwarové úrovni.

Následujícím návrhem na zlepšení by mohlo být vytvoření tzv. silného klienta ve vozidle. V části nazvané „Člověk jako stroj“ jsem na příkladu pana MacGregora ukazoval, jakým způsobem zajišťuje tělo člověka správnou vzpřímenost. Ve zkratce řečeno: Využívá několik oddělených systémů, ovšem ne ve stejné představě, jako klasicky nabízí spolehlivost systémů (tedy čistě redundancí). Je to souhra více a méně specializovaných systémů na danou činnost.¹⁰

Obdobně lze využít údaje z palubního počítače vozidla. Takové spojení si můžeme představit tak, že propojíme OBU jednotku s palubním počítačem automobilu. Můžeme tak získat okamžité informace o rychlosti vozidla, což je údaj, který může zpřesnit polohu vozidla při horším signálu GPS. Přesnost určení polohy se dá zvýšit párováním více obdobných technologií – například duální přijímač signálu GPS a GLO-NASS či GALILEO.

⁹Například u technologie GPRS či EDGE bylo zpoždění dané přenosem běžně v řádech sekund a docházelo k výpadkům především při tzv. handoveru – přechodu mezi BTS.

¹⁰Zpravidla je jeden specializovaný na daný úkol a ostatní provádějí kontrolu správnosti jako svoji vedlejší činnost. Jinak řečeno, signály z dalších systémů a jejich zpracování může pomoci k zajištění lepšího chodu, jež primárně zajišťuje systém na to specializovaný.

3.2. Informační systém pro přepravu nebezpečných věcí využívající systém GNSS

3.2.1. Základní informace

Projekt s pracovním názvem NEBNAK¹¹ byl řešen mezi roky 2001-2006. Cílem tohoto projektu bylo navrhnout obecnou architekturu pro systém na sledování nebezpečných nákladů, nastínit více možností, jak daný problém řešit při použití různých technologií.¹² Při přepravě nebezpečných nákladů se vychází z mezinárodní dohody známé pod zkratkou ADR,¹³ která sdružuje kritéria pro přepravu nebezpečných nákladů.

Systém se skládá ze serveru NEBNAK, na který se přes internet přihlašují firmy, jež se podílejí na přepravě nebezpečných nákladů. Založení přepravy probíhá přes webové rozhraní vyplněním příslušného nákladního listu na serveru NEBNAK. Server NEBNAK je pomocí technologie GPRS propojen s OBU jednotkou ve vozidle. Základní OBU jednotka je vybavena navigačním systémem využívající GPS, systémem eCall (tísňového volání) a tlačítkem „panic button“ pro případ nenadálého nebezpečí nebo poruchy vozidla. Dle potřeby a charakteru nákladu jsou připojeny další senzory, kupříkladu teplotní čidlo. Propojení OBU s perifériemi je podrobněji popsáno v příloze D.

Při výjezdu vozidla jsou serverem NEBNAK předány OBU jednotce vybavené navigací potřebná data tak, aby mohla vypočítat vhodnou trasu dle parametrů. Zároveň po celou dobu přepravy může dispečer přes NEBNAK server operativně měnit trasu vozidla. Řidič se pak drží pokynů navigačního systému. Blokové schéma systému včetně jeho propojení na systém eCall je uveden v příloze C.

3.2.2. Zhodnocení projektu z hlediska kognitivních procesů

Stejně jako u předchozího případu se můžeme na projekt dívat z několika úhlů. Opět můžeme jednotlivá vozidla vnímat jako agenty, s tím rozdílem, že svět, ve kterém se pohybují, se zvětšil – jedná se o veškeré silniční plochy v ČR (potenciálně v Evropě), které jsou vhodné k dané přepravě. Cílem centrálně sledovaných a řízených agentů je udržovat mezi sebou dostatečné rozestupy, volit vhodné trasy k přepravě (s pomocí

¹¹Zkratka pro Nebezpečné náklady.

¹²V době řešení projektu ještě nebylo zcela vyjasněno zapojení přidružených technologií, jako je systém eCall, nebo povinnost vybavit automobily telematickými jednotkami dané Evropskou unií.

¹³Evropská dohoda o mezinárodní silniční přepravě nebezpečných věcí, v originále Accord européen relatif au transport international des marchandises dangereuses par route.

serveru NEBNAK a případně dalších online informacích) a operativně jednat při nepředvídaných situacích.

Když se podíváme na samostatnou OBU jednotku ve vozidle, tak může být ve formě lehkého nebo těžkého klienta. Lehký klient umožňuje komunikaci pomocí GPRS, polohu získává z GPS přijímače a je vybaven funkcemi pro eCall. Neobsahuje navigaci pro posádku vozidla, pouze umožňuje dispečerovi jej sledovat a případně se s posádkou spojit telefonicky. Zajímavější a v projektu více testovaný byl silný klient, jenž umožňuje posádce přímou navigaci a tím i větší autonomii, neboť mají přehled o aktuálním dění, které jim zasílá server NEBNAK. Zůstaňme tedy nadále u vyspělejšího těžkého klienta, který byl projektovou zprávou také doporučen.

Navigační program Dynavix využívá při vyhledávání trasy upravený algoritmus A* pro heuristické prohledávání. Program umí pracovat jak se statickými daty (mapové podklady společnosti TeleAtlas s průjezdními profily) tak s dynamickými, které získává online (RDS-TMC kanál, GPRS). Díky tomu, že je jednotka neustále spojena se serverem NEBNAK (online), tak je vhodnější využít technologii GPRS pro přenos aktuálních dopravních dat o nehodách, uzavírkách, kongesci nebo počasí. Získáme tak data z mnoha vnějších senzorů, což zvyšuje vhodnost výběru trasy pro přepravu a snižuje riziko nebezpečí. Čím více relevantních informací máme k dispozici, tím lépe může inteligentní systém plnit svou úlohu.

Je velmi obtížné najít v takovém systému, který je zaměřen na bezpečnost, něco, co je hodné vylepšení. Přesto chci poukázat na problém, který není akutní v kontextu České republiky, ale v případě jiných států. Jedná se o tunely – kritické úseky, které jsou časté především v horských státech (Rakousko či Itálie). Ať už se jedná o tunel krátký nebo dlouhý několik kilometrů, vždy po vjezdu do něj ztratíme výhled na oblohu, a tedy i GPS signál. Ačkoli tedy může fungovat komunikace přes GPRS, tak neznáme přesnou polohu vozidla v tunelu.¹⁴ Tento problém se dá opět řešit už výše zmíněným napojením na palubní jednotku vozidla a zjišťováním aktuální rychlosti. Tak můžeme zajistit, že budeme moci zjistit pozici vozidla i při přerušení GPS signálu.

Z hlediska interakce celého systému probíhá komunikace na dvou úrovních – s řidičem a na druhé straně s dispečerem. Inteligentní dopravní systém nám vizuálně představuje rychle srozumitelné informace o poloze jednotlivých vozidel (dispečer) nebo polohu a trasu cesty v případě řidiče. Řidič je navíc hlasově navigován na cíl

¹⁴A také i v úseku po vyjetí z tunelu, dokud opět GPS přijímač nezačne přijímat data.

cesty, takže se může plně koncentrovat na řízení vozidla. V případě nouze je k dispozici tlačítko „panic button“, po jehož zmáčknutí se pokusí dispečer spojit s řidičem vozidla, a pokud to není možné, tak okamžitě aktivuje záchranné složky.

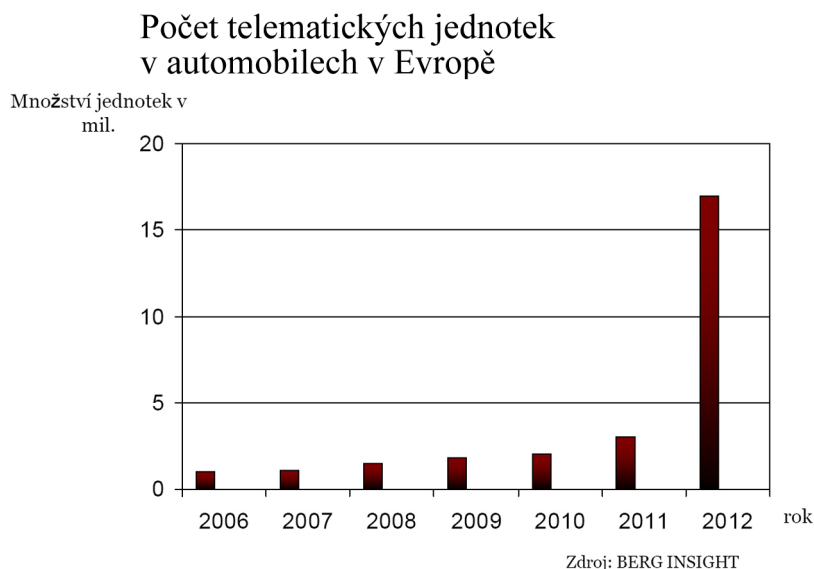
3.3. Směřování k inteligentnímu vozidlu

Zmíněné dva pilotní projekty jsou pouhým předstupněm směrem k tomu, jak má v budoucnu vypadat doprava využívající plně možnosti dopravní telematiky (inteligentních dopravních systémů). Cílem současných výzkumných aktivit nejenom v České republice je představit inteligentní dopravní vozidlo, které bude řidiči maximálně ulehčovat jeho řízení. První výsledky tohoto výzkumu už jsou v dnešní době aplikované, pro příklad uvedu: automatické brzdění při rychlém přiblížení, promítání informací do zorného pole řidiče (Head-up displej). Pracuje se také na automatickém rozpoznávání dopravních značek, rozpoznávání unavenosti řidiče (mikrospánků) a různé možnosti využití aktuálních dopravních dat v reálném provozu při součinnosti s navigacemi. Cílem je také zavedení centrálního řízení, aby se zvýšila efektivita a plynulost dopravy při minimální kongesci.¹⁵

Tyto iniciativy podporuje i Evropská unie, která uložila výrobcům automobilů povinnost v blízkém horizontu vybavit automobily telematickými jednotkami vybavené kromě jiného také systémem eCall. Mělo by se jednat o černé skříňky,¹⁶ které budou zachraňovat lidské životy. Počet takto vybavených vozidel by se měl zvyšovat. Pro představu uvedu graf publikovaný společností Telematix Software a.s. na březnové mezinárodní konferenci NavAge 2008 (převzato z prezentace).

¹⁵Systém řízení dopravy doporučí jednotlivým vozidlům trasy tak, aby je co nejrovnoměrněji rozprostřel na vozové komunikace.

¹⁶Které by měly být velmi odolné, podobně jako černé skříňky v letadlech.



Obrázek 3.1.: Rozvoj telematických jednotek v automobilech.

V praxi by to při havárii vozidla mohlo vypadat následovně: Okamžitě po nehodě odešle vozidlo (telematická jednotka) tísňové volání s tím, že uvede svou polohu, čas nehody, počet cestujících (senzory v sedadlech), přibližný rozsah poškození vozidla (různé senzory po celém autě), rychlost vozidla při srážce a typ vozidla. Pokud bude jednotka schopna další komunikace, tak může zaslat i další informace, které měla k dispozici těsně před srážkou (např. stav paliva v nádrži, druh nákladu). Na straně dispečerského centra, které přijme tísňové volání, dojde k aktivaci a předání důležitých údajů záchranným složkám, v případě hromadné nehody i vhodné koordinace vzhledem ke zjištěným informacím o nehodě. Také může v daném úseku odklonit dopravu a zajistit tím plynulost provozu.

4. Závěr

„Dokážu si snadno představit, že příští velké dělení světa bude mezi lidmi, kteří touží žít jako lidské bytosti, a těmi, kdo chtějí žít jako stroje.“

WENDELL BERRY

V závěrečné části práce bych se trošku netradičně zabýval vyústěním problematiky inteligentních systémů a jejich rozvoje, což je neodmyslitelně spojeno se zdokonalováním jejich kognitivních procesů.

Poznávací procesy utvářejí vidění našeho světa a my tvorbou poznávacích procesů utváříme svět umělé inteligence, respektive inteligentních systémů, které pracují ve fyzickém prostředí, případně ve virtuálním prostředí. Tím také zpětně ovlivňujeme (rozšiřujeme) i své vlastní vidění světa. Inteligentní systémy se staly každodenní součástí našeho života, stejně jako virtuální světy vytvořené pro komunikaci s počítačem, které později přinesli alternativu k fyzickému světu, podobně jako dříve naše sny a představy.

Je to, jako by se nám se vzrůstajícími (zprostředkovanými) kognitivními schopnostmi (hnané rozvojem techniky) otevíral svět stále více, ale přesto nikdy ne zcela, jak se mohou někteří domnívat. To si bohužel často neuvědomujeme. Příkladem může být DR. P.,¹ který nepocíťoval, že by něco na vidění světa oproti běžnému člověku ztratil.

Vztah člověka a stroje se postupem času začíná měnit a čím dál více propojovat na všech úrovních. Běžný člověk tyto změny ani nevnímá, bere je jako normální součást života. Na scénu tak vstupují postmoderní filosofové a upozorňují nás na důsledky tohoto vývoje. Poukazují na příchod období posthumanismu, období charakteristické spojením (propojením) člověka a stroje (technologií), což je také označováno jako konec lidského a nástup nelidského. [Sim, 2003, str. 13]

Nemá cenu se ptát, jestli je to dobře nebo špatně, jako bychom si ještě mohli vybrat – už to nastalo a my se s tím musíme smířit a dále s tím pracovat. Mezi takové

¹Zmiňovaný v části Člověk jako stroj.

(nelidské) technologie patří i umělá inteligence a umělý život, které v očích některých proroků doby představují hrozbu v podobě pokročilejší životní formy, než jsme my. Kromě této technofobie, které toto směřování může budit, nalezneme i názory, jež naopak podporují zavádění nových technologií za účelem zlepšování lidského života.² Tento směr je označován jako transhumanismus.

Naděje i strach jsou při pohledu do budoucího směřování samozřejmostí. Je ovšem třeba si tento vývoj připustit a být si vědomi jeho hlavního činitele, což je „techno-vědní“ pokrok hnáný kapitalismem.

„Vědecké poznání je určitým vyjádřením názoru.“

JEAN-FRANÇOIS LYOTARD

Když se podíváme na samotný technický pokrok, který reprezentují především exaktní vědní disciplíny, tak nesmíme zapomenout na důležitý faktor, který je ovlivňuje, totiž na naše lidství.³ A právě toto lidství, ke kterému se v myšlenkách o tom, jaká bude ona nelidská budoucnost vypadat, obracíme, nám tuto budoucnost vlastně připravuje.

Věda se totiž může stát dobovou záležitostí, kdy vědci jeden po druhém akceptují a připouštějí obecné názorové prostředí. [Lukacs, 2009, str. 70] Tak tomu do značné míry je i u vzpomínané „strunové revoluce“ ve fyzice. Nejhuře kam to může zajít, naznačil už J. W. GOETHE (1749 - 1832), který ve svém rozhovoru s J. P. ECKERMANNEM (1792 - 1854) upozorňuje na lidskou ctižádostivost, kdy profesori (učenci či vědci), kteří navzdory opačným důkazům vůči své teorii dále trvají na svém (na své nepravdě), neboť své teorii vděčí za to, čím jsou a čeho dosáhli. Bez toho by museli začít od začátku. Jde jim tedy spíše o prosazení obecného názoru. [Lukacs, 2009, str. 76] Často nás tak věda zavádí do slepých uliček nejrůznějších konstruktů a experimentů daleko mimo oblast naší přímé zkušenosti.

Někteří jazykové tvrdí, že bychom se měli zabývat spíše člověkem⁴ a tím, co člověk může dále využívat, než zkoumat teorie operující s desítkami dimenzí nebo biologic-

² Ať už se jedná o technologie vylepšující mentální nebo fyzické aspekty člověka. Nebrání se tedy kupříkladu myšlence kyborga, pokud to pro člověka přinese něco kladného – zlepšení kvality života.

³ Lépe řečeno lidské jednání, to co člověka činí člověkem.

⁴ „Zajímá a znepokojuje mě přesvědčení, podle nějž se výsledky zkoumání primitivních organismů považují za samozřejmost a za automaticky aplikovatelné (ne-li již nadřazené) na naše poznání nejsložitějších organismů ve vesmíru.“ [Lukacs, 2009, str. 68] Včetně samotného vztahů mezi vědci, kteří přinášejí teorie a provádí výzkumy – čímž udávají vědě směr.

kými procesy uvnitř buňky. Určitě je v tom díl pravdy, leč zvědavost lidstva a zapálení pro pokrok je veliký. Přece bychom měli mít na paměti hlavní cíl (heslo) nejen v souvislosti s umělou inteligencí, zmiňovaný v první kapitole. Jakou formou toho vědec dosáhne, to už je jenom na něm.

Jednodušeji řečeno, směr, kterým se vydá naše nelidská budoucnost, je určen obecně přijímaným míněním, a to na jedné straně vědeckým a posléze i veřejným. Náзор vědců, jakým směrem se vydat, pokud je dostatečně výdělečný, se díky kapitalismem užívané mediální propagandě obratem dostává k veřejnosti, která tento směr bude vyžadovat.⁵ Neboli, směřování k nelidství je určováno paradoxně naším lidstvím. Důležitá je tedy cesta, kterou se vydáme, nikoli cíl (v dobrém či špatném smyslu), kterého můžeme dosáhnout.

*„Samotný život je jenom tenká vrstva barvy na naší planetě a my držíme
malířský štětec.“* DANIEL DENNETT

Postmoderní pohled odmítá nadřazenosti čisté racionality v procesu poznání. Ani člověk není zcela racionální. Ačkoli jsme byli posledních několik staletí ohromeni možnostmi, které racionalizace přinesla, přesto se obracíme k dalším přístupům k poznávání.⁶

Inspirujeme se životem, který je úzce spjat v globálním charakteru s naší existencí. Převzali jsme od přírody v jistém smyslu otěže budoucího vývoje života, a dokonce i sami sebe.

*„Individuum neznamena mnoho, ale žádné samo o sobě
není pouhým osamělým ostrovem; každé existuje ve struktuře
vztahů, které jsou nyní složitější a proměnlivější než kdy předtím.“*

JEAN-FRANÇOIS LYOTARD

⁵Mění se tak názor lidí na to, co je vhodné a co není. Příkladem může být sledování osob či aut. Když si to představíme v kontextu například 80. let 20. století, tak by to byla obecně odmítaná „služba“, s přihlédnutím k románu 1984. Ovšem nyní, kdy „ztráta soukromí“ v sobě ukrývá i další výhody (např. systém eCall, navigace nebo sociální sítě), je to naopak všeobecně přijímaná technologie, respektive směr.

⁶Často označováno jako kognitivní obrat.

Díky „předloze“, kterou nám život nabízí a možnosti tvorby, kterými disponuje člověk, směřujeme k novému chápání života i člověka samého. U člověka přichází období masovosti a nového pojetí individualizace člověka. Ten bude silně ovlivňován masovou kulturou, která mu nabídne svou verzi individualizace, kupříkladu ve formě alternativních světů. Naopak život, který jsme už schopni částečně přetvářet k obrazu svému, se čím dál více začne propojovat i s našimi dalšími výtvary. Také v tomto procesu, jenž nemusí vést jenom ke změně člověka v kybernetický organismus, budou hrát ústřední roli kognitivní procesy. Při tom je nutné tolik diskutované přiblížení i na straně kognitivních procesů umělých systémů a těch, které pozorujeme v přírodě. Dle mého názoru je to ještě dlouhá cesta, dost možná delší, než odhadují vědeckofantastické knihy na toto téma.

Terminologický slovník

Antinomie – Protiklad.

CAN sběrnice – Sběrnice určená ke komunikaci mezi řídicími jednotkami v automobilu. Zkratka pro Controller Area Network.

Dependabilita – Schopnost systému poskytovat službu, na kterou se dá spolehnout.

Dopravní telematika - Integruje informační a telekomunikační technologie s dopravním inženýrstvím.

Dualismus – Předpokládá, že hmota a mysl jsou dvě rozdílné kvality.

Duše – Široký pojem poukazující na kvalitu schovanou v lidské bytosti.
V křesťanství spojena s jedinečností a identitou osoby.

eCall – Jednotný evropský systém, který má pomoci motoristům v případě havárie tím, že odešle údaje o nehodě na evropskou linku 112.

Endocept – Vnitřní model světa v naší mysli.

Fenotyp – Soubor vlastností, které vznikají při působení genotypu a prostředí.
To, jak organizmus ve výsledku fyzicky vypadá.

GALILEO – viz GNSS

Genotyp – Soubor genetických informací organismu.

GLONASS – viz GNSS

GNSS – Globální družicový polohový systém, kterých existuje v současné době hned několik, přičemž nejpoužívanější je americký systém GPS. Další země, jež mají nebo připravují vlastní systémy: Rusko (GLONASS), Francie (DORIS), Evropská unie (GALILEO), Čína (Beidou-2).
Zkratka pro Global Navigation Satellite System.

GPRS – Mobilní datová služba pro uživatele sítí GSM.
Zkratka za General Packet Radio Service.

GPS – viz GNSS

ITS – viz Dopravní telematika

Komputacionismus – Teorie, která se domnívá, že mysl je funkce, kterou lze provádět nezávisle na fyzické struktuře.

Obdobně jako např. program pracující se symboly.

Konekcionismus – Teorie, která předpokládá, že mysl povstává z jednoduchých, vzájemně komplexně propojených jednotek.

Příkladem může být neuronová síť.

Kongesce – Dopravní zácpa.

Konstrukt mysli – Vytváří ji mysl na základě smyslových dat, která jsou jí dodávána z vnějšku, přičemž to, co je pro vzniklý obraz skutečně podstatné, jsou apriorní formy mysli, nikoliv vnější data.

Monismus – Vše povstává z jedné substance, což je protikladem k dualismu, respektive pluralismu.

Mysl – Soubor složitých procesů, jako vnímání, poznávání, cítění, rozumovost, pamatování a dalších.

Nástroj – Jednoduchá věc, která člověku rozšiřuje jeho základní schopnosti práce se světem.

OBU – Počítač ve vozidle, který může být vybaven displejem a dalšími periferiemi. Zkratka pro On Board Unit.

RDS-TMC – Systém doplňkových zpráv k FM vysílání, v tomto případě je služba TMC zaměřena na aktuální dopravní informace.
Zkratka pro Radio Data System a Traffic Message Channel.

Stroj – Zařízení sestavené člověkem, které přeměňuje jeden druh energie nebo síly v jiný.

Umělá inteligence – Obor, který se zabývá vytvářením strojů, které mají podobné (inteligentní) chování jako člověk.

Umělý život – Obor, který modeluje procesy pozorované u živých entit s cílem využít dále tyto poznatky (analogie) v praxi.

WiMAX – Bezdrátová technologie pro komunikaci. Zkratka pro Worldwide Interoperability for Microwave Access.

Seznam obrázků

1.1. Přehledné představení zmíněných dualistických koncepcí. . .	13
2.1. Nepřímá interakce člověka s prostředím.	39
2.2. Dva pohledy na virtuální svět.	45
2.3. Mozková aktivita opice při jednotlivých úkonech. Převzato z [Rizzolatti a kol., 2001].	59
2.4. Přehled různých úrovní programovacích jazyků.	65
2.5. Schématický náčrt interakce buňky s prostředím.	76
3.1. Rozvoj telematických jednotek v automobilech.	85

Seznam příloh

Příloha A – Blokové schéma navrženého systému v projektu CaMNA

Příloha B – Pohled na dispečerského stanoviště a vybavení vozidel
pohybujících se po letištní ploše

Příloha C – Blokové schéma vazby informačního systému pro podporu
přepravy nebezpečných věcí a systému eCall

Příloha D – Propojení OBU s perifériemi a Softwarová architektura
projektu NEBNAK

Kolofón

Tato práce byla vysázena v *L^AT_EX*u s využitím písma

New Century Schoolbook, Helvetica a Computer Modern Typewriter.

Finální verze dokumentu je z 4. května 2010.

Literatura

- [Akerkar, 2005] AKERKAR, Rajendra. *Introduction to Artificial Intelligence*. New Delhi (India): PHI Learning Private, 2005. 349 s. ISBN: 978-81-203286-48.
- [Bais, 2009] BAIS, Sander. *Rovnice. Symboly poznání*. 1. vyd. Praha: Dokořán, 2009. 96 s. ISBN 978-80-7363-228-1.
- [Barbieri, 2006] BARBIERI, Marcello. *Organické kódy: Úvod do sémantické biologie*. 1. vyd. Praha: Academia, 2006. 240 s. ISBN 80-200-1403-9.
- [Benenson a kol., 2004] BENENSON, Yaakov, GIL, Binyamin, BEN-DOR, Uri, ADAR, Rivka, SHAPIRO EHUD. *An autonomous molecular computer for logical control of gene expression*. [online]. Publikováno online 28. dubna 2004 [cit. 2010-01-30]. 6 s. Dostupný z WWW: <http://www.wisdom.weizmann.ac.il/~udi/papers/automoleculcomp_nat04.pdf>.
- [Berka, 2003] BERKA, Petr. *Současné trendy umělé inteligence*. [online]. 2003 [cit. 2009-09-24]. 6 s. Dostupný z WWW: <<http://sorry.vse.cz/~berka/docs/4iz430/P00-TrendyAI.pdf>>.
- [Blue Brain, 2010] *The Blue Brain Project*. [online]. Ecole Polytechnique Fédérale de Lausanne, 2010 [cit. 2010-03-20]. Dostupný z WWW: <<http://bluebrain.epfl.ch/>>.
- [Borges, 1999] BORGES, L. Jorge. *The analytical language of John Wilkins*. [online]. 1999 [cit. 2009-10-20]. Dostupný z WWW: <<http://www.alamut.com/subj/artiface/language/johnWilkins.html>>.
- [Bourbakis, 1992] BOURBAKIS, G. Nikolaos. *Artificial intelligence methods and applications*. World Scientific Publishing, 1992. 705 s. ISBN: 978-98-102105-71.
- [Burian, 2005] BURIAN, Jan. *Kognice kontra informace*. [online]. Praha: Vysoká škola ekonomická, 2005 [cit. 2010-04-26]. 12 s. Dostupný z WWW:

- <http://eldar.cz/honza/articles/burian_mysleni_modelu_final.doc>.
- [Carter, 2007] CARTER, Matt. *Minds and Computers: An Introduction to the Philosophy of Artificial Intelligence*. Edinburgh (UK): Edinburgh University Press, 2007. 222 s. ISBN: 978-0-7486-2098-2.
- [Čada, 1993] ČADA, Ondřej. *Operační systémy*. 1. vyd. Praha: Grada, 1993. 384 s. ISBN 80-85623-44-7.
- [Černý, Kočandrlová, 1998] ČERNÝ, Jaroslav, KOČANDRLOVÁ, Milada. *Konstruktivní geometrie*. 1. vyd. Praha: Vydavatelství ČVUT, 1998. 262 s. ISBN 80-01-01815-6.
- [Fišer, 2003] FIŠER, Jiří. *Principy operačních systémů I*. 1. vyd. Ústí nad Labem: PF UJEP, 2003. 99 s. ISBN 80-7044-505-X.
- [Feynman, 2007] FEYNMAN, P. Richard. *Šest snadných kapitol*. 1. vyd. Praha: AURORA, 2007. 164 s. ISBN 978-80-7299-090-0.
- [Gray, Singer, 1989] GRAY, M. Charles, SINGER, Wolf. *Stimulus-specific neuronal oscillations in orientation columns of cat visual cortex*. [online]. Max Planck Institute for Brain Research, 1989 [cit. 2010-04-05]. Dostupný z WWW: <<http://www.pnas.org/content/86/5/1698.full.pdf+html>>.
- [Habiballa, 2004] HABIBALLA, Hashim. *Umělá inteligence*. [online]. Ostrava: Ostravská Univerzita, 2004 [cit. 2009-09-29]. 81 s. Dostupný z WWW: <<http://www.volny.cz/habiballa/publ/umint.pdf>>.
- [Hajnal, 2005] HAJNAL, László. *Tělo a mysl – východiska a alternativy*. [online]. Katedra filosofie a dějin přírodních věd UK PřF, 2005 [cit. 2010-04-15]. 10 s. Dostupný z WWW: <<http://hilbert.chtf.stuba.sk/KUZV/download/kuzv-hajnal.pdf>>.
- [Hanušová a kol., 2006] HANUŠOVÁ, Marie, OUDOVÁ, Drahomíra, VOTAVA, Jiří. *Učení a paměť*. [online]. Institut vzdělávání a poradenství České zemědělské univerzity, 2006 [cit. 2010-04-11]. 159 s. Dostupný z WWW: <<http://etext.czu.cz/img/skripta/64/pamet-1.pdf>>.
- [Havel, Mitášová, 2009] HAVEL, M. Ivan, MITÁŠOVÁ Monika. *Percepce a tvorba agregátových objektů v prostoru*. [online]. 2009 [cit. 2010-04-26]. 11

- s. Dostupný z WWW: <<http://dai.fmph.uniba.sk/events/kuz2009/prispevky-pdf/havel-mitasova.pdf>>.
- [Haynes a kol., 2008] HAYNES, A. Karmella a kolektiv. *Engineering bacteria to solve the Burnt Pancake Problem*. [online]. 2008 [cit. 2009-11-10]. 12 s. Dostupný z WWW: <<http://www.jbioleng.org/content/pdf/1754-1611-2-8.pdf>>.
- [Hayward, Varela, 2009] HAYWARD, J. W., VARELA, F. J. *Mosty k porozumění: rozhovory předních vědců s dalajlamou o zkoumání lidské mysli*. 1. vyd. Praha: DharmaGaia, 2009. 332 s. ISBN 978-80-86685-83-0.
- [Heim, 1998] HEIM, Michael. *Virtual Realism*. 1. vyd. New York: Oxford University Press, 1998. 238 s. ISBN 978-0-19-513874-0.
- [Hodek, Šulc, 2009] HODEK, Petr, ŠULC, Miroslav. *Dynamika růstu a množení mikroorganismů*. [online]. 2009 [cit. 2009-11-14]. Dostupný z WWW: <<http://mikrobiologie.webzdarma.cz/5.pdf>>.
- [Horrocks, 2002] HORROCKS, Christopher. *Marshall McLuhan a virtualita*. 1. vyd. Praha: Triton, 2002. 76 s. ISBN 80-7254-269-9.
- [Ideje, 2009] *Robot s biologickým mozkiem*. In Ideje.cz [online]. Praha: Prague Media Group, 2009 [cit. 2009-11-29]. Dostupný z WWW: <<http://www.ideje.cz/cz/clanky/robot-s-biologickym-mozkem>>.
- [ICCI 2010] *The 9th IEEE International Conference on Cognitive Informatics (ICCI 2010)*. [online]. Tsinghua University, Beijing, China, 2010 [cit. 2010-03-18]. Dostupný z WWW: <<http://www.icci2010.edu.cn/>>.
- [Iser, 2009] ISER, Wolfgang. *Jak se dělá teorie*. 1. vyd. Praha: Karolinum, 2009. 246 s. ISBN 978-80-246-1672-8.
- [ITS, 2005] *Inteligentní dopravní systémy v České republice*. [online]. Ministerstvo dopravy ČR, 2005 [cit. 2010-04-22]. Dostupný z WWW: <<http://www.mdcr.cz/NR/rdonlyres/CEF8732F-19F1-43CB-9A37-1D299EF10D21/0/PublikaceITSMdcesky.pdf>>.
- [Jiroušek, 1995] JIROUŠEK, Radim. *Metody reprezentace a zpracování znalostí v umělé inteligenci*. Praha: Vysoká škola ekonomická, 1995. 103 s. ISBN 80-7079-701-0.
- [Kauffman, 2004] KAUFFMAN, Stuart. *Čtvrtý zákon – Cesty k obecné biologii*. 1. vyd. Praha: Paseka, 2004, 261 s. ISBN: 978-80-7185-636-3.

- [Kelemen, 1994] KELEMEN, Jozef. *Strojovia a agenty*. 1. vyd. Bratislava: Archa, 1994, 109 s. ISBN: 80-7115-089-4.
- [Kelemen, 1998] KELEMEN, Jozef. *Postmoderný stroj*. 1. vyd. Bratislava: F. R. & G., 1998, 128 s. ISBN: 80-85508-45-1.
- [Kelemen, 2010] KELEMEN, Jozef. *Myslenie a stroj*. 1. vyd. Bratislava: Kalligram, 2010, 388 s. ISBN: 978-80-8101-243-3.
- [Klíma, 2008] KLÍMA, Milan. *Záhady lidského těla*. 1. vyd. Praha: Ikar, 2008. 160 s. ISBN 978-80-249-1158-8.
- [Koukolík, 2010] KOUKOLÍK, František. *Lidství: Neuronální koreláty*. 1. vyd. Praha: Galen, 2010. 256 s. ISBN 978-80-7262-654-0.
- [Koukolík, 2006] KOUKOLÍK, František. *Sociální mozek*. 1. vyd. Praha: Karolinum, 2006. 270 s. ISBN 978-80-246-1242-3.
- [Krámský ed., 2009] KRÁMSKÝ, David (ed.). *Kognitivní věda dnes a zítra*. 1. vyd. Liberec: Nakladatelství Bor, 2009. 304 s. ISBN 978-80-86807-55-3.
- [Křemen, 2007] KŘEMEN, Jaromír. *Modely a systémy*. 1. vyd. Praha: Academia, 2007. 100 s. ISBN 978-80-200-1477-1.
- [Linhart, 1972] LINHART, Josef. *Proces a struktura lidského učení*. 2. vyd. Praha: Academia, 1972. 488 s.
- [Linhart, 1976] LINHART, Josef. *Činnost a poznávání*. 1. vyd. Praha: Academia, 1976. 576 s.
- [Lovgren, 2003] LOVGREN, Stefan. *Computer Made from DNA and Enzymes*. Nationalgeographic.com [online]. Washington, D.C.: National Geographic Society, 2003 [cit. 2010-01-30]. Dostupný z WWW: <http://news.nationalgeographic.com/news/2003/02/0224_030224_DNAcomputer.html>.
- [Lukacs, 2009] LUKACS, John. *Na konci věku*. 1. vyd. Praha: Academia, 2009. 152 s. ISBN 978-80-200-1781-9.
- [Mareš, 2006] MAREŠ, Milan. *Slova, která se hodí, aneb jak si povídat o matematice, kybernetice a informatice*. 1. vyd. Praha: Academia, 2006. 352 s. ISBN 80-200-1445-4.
- [Marková, 2007] MARKOVÁ, Ivana. *Dialogičnost a sociální reprezentace: Dynamika myslí*. 1. vyd. Praha: Academia, 2007. 284 s. ISBN 978-80-200-1542-6.

- [Mařík a kol., 1993] MAŘÍK, Vladimír, ŠTĚPÁNKOVÁ, Olga, LAŽANSKÝ, Jiří a kolektiv. *Umělá inteligence (1)*. 1. vyd. Praha: Academia, 1993. 264 s. ISBN 80-200-0496-3.
- [Mařík a kol., 1997] MAŘÍK, Vladimír, ŠTĚPÁNKOVÁ, Olga, LAŽANSKÝ, Jiří a kolektiv. *Umělá inteligence (2)*. 1. vyd. Praha: Academia, 1997. 374 s. ISBN 80-200-0504-8.
- [Mařík a kol., 2001] MAŘÍK, Vladimír, ŠTĚPÁNKOVÁ, Olga, LAŽANSKÝ, Jiří a kolektiv. *Umělá inteligence (3)*. 1. vyd. Praha: Academia, 2001. 328 s. ISBN 80-200-0472-6.
- [Mařík a kol., 2003] MAŘÍK, Vladimír, ŠTĚPÁNKOVÁ, Olga, LAŽANSKÝ, Jiří a kolektiv. *Umělá inteligence (4)*. 1. vyd. Praha: Academia, 2003. 476 s. ISBN 80-200-1044-0.
- [Mařík a kol., 2007] MAŘÍK, Vladimír, ŠTĚPÁNKOVÁ, Olga, LAŽANSKÝ, Jiří a kolektiv. *Umělá inteligence (5)*. 1. vyd. Praha: Academia, 2007. 544 s. ISBN 80-200-1470-2.
- [Mollon, 2000] MOLLON, Phil. *Freud a syndrom falešné paměti*. 1. vyd. Praha: Triton, 2000. 80 s. ISBN 80-7254-145-5.
- [Pavlík, 2004] PAVLÍK, Ján. *F. A. Hayek a teorie spontánního řádu*. 1. vyd. Praha: Professional Publishing, 2004. 805 s. ISBN 978-80-86419-57-6.
- [Pěchouček, 2005] PĚCHOUČEK, Michal. *Úvod do filosofie umělé inteligence*. [online]. 2005 [cit. 2009-11-07]. Dostupný z WWW: <<http://cyber.felk.cvut.cz/gerstner/teaching/zui/kui-phil.htm>>.
- [Peregrin, 2008] PEREGRIN, Jaroslav. *Filozofie pro normální lidi*. 1. vyd. Praha: Dokořán, 2008. 142 s. ISBN 978-80-7363-192-5.
- [Petrů, 2007] PETRŮ, Marek. *Fyziologie mysli: úvod do kognitivní vědy*. 1. vyd. Praha: Triton, 2007. 392 s. ISBN 978-80-7254-969-6.
- [Petříček1997] PETŘÍČEK, Miroslav. *Úvod do současné filozofie*. 4. vyd. Praha: Herrmann & synové, 1997. 180 s.
- [Petříček, 2009] PETŘÍČEK, Miroslav. *Myšlení obrazem*. 1. vyd. Praha: Herrmann & synové, 2009. 202 s. ISBN 978-80-87054-18-5.
- [Pstružina, 1994] PSTRUŽINA, Karel. *Etudy o mozku a myšlení*. [online]. 1994 [cit. 2010-02-24]. Dostupný z WWW: <<http://nb.vse.cz/~pstruzin/monog/etudy.htm>>.

- [Pstružina, 1995] PSTRUŽINA, Karel. *Kognitivní vědy a ontologická diference. Svět snů*. E-LOGOS [online]. 1995 [cit. 2010-03-08]. Dostupný z WWW: <<http://nb.vse.cz/kfil/elogos/mind/ont-dif.htm>>. ISSN 1211-0442.
- [Pstružina, 1998] PSTRUŽINA, Karel. *Svět poznávání: k filozofickým základům kognitivní vědy*. 1. vyd. Olomouc: Nakladatelství Olomouc, 1998. 184 s. ISBN 978-80-7182-074-1.
- [Pstružina, 2005] PSTRUŽINA, Karel. *Pojednání o lidském myšlení I*. 1. vyd. Praha: Ekopress, 2005. 234 s. ISBN 978-80-86119-89-0.
- [Rizzolatti a kol., 2001] RIZZOLATTI, Giacomo, FOGASSI, Leonardo, GALLESE, Vittorio. *Neurophysiological mechanisms underlying the understanding and imitation of action*. [online]. Università degli Studi di Parma, 2001 [cit. 2010-03-30]. Dostupný z WWW: <<http://www.unipr.it/arpa/mirror/pubs/pdf/Rizzolatti-Fogassi%202001.pdf>>.
- [Rusina, 2004] RUSINA, Robert. *Paměť a její poruchy*. *Neurologia pre prax*. 3/2004, 4, s. 200-201. Dostupný také z WWW: <<http://www.solen.cz/pdfs/neu/2004/04/04.pdf>>. ISSN 1335-9592.
- [Říčan, 2007] ŘÍČAN, Pavel. *Psychologie osobnosti: Obor v pohybu*. 1. vyd. Praha: Grada Publishing, 2007. 200 s. ISBN 978-80-247-1174-4.
- [Říhová, Vítek, 1989] ŘÍHOVÁ, Zora, VÍTEK, Miloš. *Kybernetika a teorie systémů*. 1. vyd. Praha: Vysoká škola ekonomická, 1989. 118 s.
- [Sacks, 2008] SACKS, Oliver. *Muž, který si pletl manželku s kloboukem*. 2. vyd. Praha: dybbuk, 2008. 256 s. ISBN 978-80-86862-59-0.
- [Sacks, 2009a] SACKS, Oliver. *Antropoložka na Marsu*. 2. vyd. Praha: dybbuk, 2009. 288 s. ISBN 978-80-86862-81-1.
- [Sacks, 2009b] SACKS, Oliver. *Musicophilia*. 1. vyd. Praha: dybbuk, 2009. 376 s. ISBN 978-80-86862-92-7.
- [Searle, 1994] SEARLE, J. R. *Mysl, mozek a věda*. 1. vyd. Praha: Mladá fronta, 1994. 136 s. ISBN 80-204-0509-7.
- [Sim, 2003] SIM, Stuart. *Lyotard a nelidské*. 1. vyd. Praha: Triton, 2003. 70 s. ISBN 80-7254-370-9.

- [Smrčka, 2002] SMRČKA, Pavel. *Fraktální a multifraktální analýza variability srdečního rytmu v extrémních stavech lidského organismu*. Disertační práce, ČVUT Praha, 2002.
- [Smolin, 2009] SMOLIN, Lee. *Fyzika v potížích*. 1. vyd. Praha: Argo/Dokořán, 2009. 380 s. ISBN 978-80-7363-207-6.
- [Svítek a kol., 2006a] SVÍTEK, Miroslav a kolektiv. *Informační systém pro podporu přepravy nebezpečných věcí využívající systém GNSS*. Závěrečná zpráva pilotního projektu. Účast České republiky v projektu GALILEO. Grant MDS 802/210/112. Fakulta dopravní ČVUT v Praze, 2006. 119 s.
- [Svítek a kol., 2006b] SVÍTEK, Miroslav a kolektiv. *Monitorování a řízení pohybu pohyblivých objektů po pohybové ploše letiště pomocí GNSS*. Závěrečná zpráva pilotního projektu. Účast České republiky v projektu GALILEO. Grant MDS 802/210/112. Fakulta dopravní ČVUT v Praze, 2006. 136 s.
- [Šunkevič, 2007] ŠUNKEVIČ, Martin. *Projekt ministerstva dopravy a spojů "Účast České republiky v projektu GALILEO"*. [online]. Česká kosmická kancelář, 2007 [cit. 2010-04-16]. Dostupný z WWW: <<http://www.czechspace.cz/cs/galileo/aktivita-cr/narodni-projekty/ucast-cr-v-projektu-galileo>>.
- [Veselý, 2005] VESELÝ, Arnošt. *Inteligentní systémy a neuronové sítě*. [online]. Praha: Česká zemědělská univerzita, 2005 [cit. 2010-03-21]. 5 s. Dostupný z WWW: <<http://www.agris.cz/etc/textforwarder.php?iType=2&iId=137104&PHPSESSID=3e>>.
- [Wiedermann, 2004] WIEDERMANN, Jiří. *Spojení samoorganizace s výpočty: minimální život v moři umělých molekul*. [online]. Praha: Ústav informatiky AV ČR, 2004 [cit. 2010-04-26]. 15 s. Dostupný z WWW: <<http://www.cs.cas.cz/semweb/download.php?file=07-10-wiedermann-spojzeni-samoorganizace&type=pdf>>.
- [Zeman, 1964] ZEMAN, Jiří. *Kybernetika a moderní věda*. 1. vyd. Praha: NPL, 1964. 108 s.
- [Zeman, 1978] ZEMAN, Jiří. *Teorie odrazu a kybernetika*. 1. vyd. Praha: Academia, 1978. 252 s.