

České vysoké učení technické v Praze Fakulta elektrotechnická

Dizertační práce



listopad, 2008

Václav Svatoš

České vysoké učení technické v Praze

Fakulta elektrotechnická

Katedra kybernetiky

*Behaviorální řízení robotů s
hodnocením přínosu jednotlivce pro
kolektivní cíl*

Disertační práce

Václav Svatoš

Praha, *listopad 2008*

Doktorský studijní program: Elektrotechnika a informatika

Studijní obor: Umělá inteligence a biokybernetika

Školitel: *doc. Ing. Pavel Nahodil, CSc.*

Gerstnerova laboratoř, Katedra kybernetiky

Fakulta elektrotechnická, České vysoké učení technické v Praze

Karlovo náměstí 13, 121 35 Praha 2, Česká Republika

Telefon: +420 224357666, Fax: +420 224923677

WWW: <http://cyber.felk.cvut.cz>

Abstrakt

Disertační práce se zabývá behaviorálním řízením robotů a jejich přínosem pro cíl kolektivu. Hodnotí nejenom chování a řízení mobilního robota-jednotlivce ve svém prostředí, ale také toto řízení podřizuje požadovanému cíli v dalších robotů v kolektivu. Návrh přechází postupně od jednotlivých podmínek ke složitějším závěrům (bottom-up). Jádrem behaviorálního řízení robota mechanismem výběru akce (ASM). Pravidla, jimiž se aktuální akce odvozuje, pokrývají jak bezprostřední reakce tak i postupně narůstající intence – záměry (či také chtění) danou akci provést. Závěr z těchto předpokladů se pak interpretuje elementárními funkcemi robota – bazálními akcemi, v tomto případě pohybovými.

Návrh řídicího systému mobilního robota využívá hybridní architekturu. Pro koordinaci aktivity robotů v kolektivu se striktně využívá pouze několik omezených využívající vjemů z prostředí. Tyto podněty hybridní architektura zpracovává na úrovni reaktivní a intencio-nální. Ve výsledku hybridního ASM se bazální akce průběžně realizují za nejvýhodnějších podmínek. Mechanismus výběru je navržen tak, aby vybraná akce byla výhodná nejenom pro robota-jednotlivce, ale pro celý kolektiv robotů.

Možnosti řídicího systému z pohledu jednotlivce a kolektivu dokládá formalismus ekogramatického systému, v práci předvedený na skupině agentů-robotů. Jím podchycuji vývoj celého kolektivu včetně vlivů z prostředí, které nejsou přímo způsobeny jednotlivými roboty kolektivu. Aparát ekogramatického systému dále doplňují kritéria pro vyhodnocování kvality navrženého ASM. Integrálním kritériem hodnotím přínos robotů k dosažení cíle celého kolektivu, Na druhou stranu v diferenciálním kritériu dávám každému ze zapojených robotů metodiku, jak hodnotit vlastní chování vzhledem k individuálním cílům robota.

Použití kritérií jako nástroje pro adaptaci řídicího systému je ukázáno v simulacích. Výchozím zadáním simulace je seznam bazálních akcí a hodnoty kritéria v mezních případech – kdy bude situace z pohledu jednotlivce vyžadovat kritické ošetření, a zda daná situace signalizuje kolektivní či spíše individuální typ reakce. Takto nastavená kritéria pak umožňují v průběhu simulace postupně rozšířit pravidla na další použití bazálních akcí, tentokrát za nových podmínek. Navrhovaný mechanismus výběru akce jsem ověřil také v praktických experimentech s mobilními roboty; na pokrytí problému roboty a na úloze koordinovaného pohybu.

Abstract

The thesis analyzes power of collective task achievement by a behavioral control respecting robot individuals. Proposed control system architecture combines top-down commands (sub-goal assignment) with bottom-up instinctive responses. Force robot activity is built on bottom level from particular functions – basic actions. Such native task decomposition moves a problem handling toward to synchronization of distributed tasks when a robot performance is forced to one global group task achievement.

Core behavioral robot control is implemented in action selection mechanism (ASM). Designed hybrid ASM combines cues from several tiers: immediate reactions from a reactive tier, reflecting actual local context, and internal needs to goal achievement (intentions) coming from intentional tier.

A good analogy for intention enhancement has been identified in the eco-grammar framework. It gives a sufficient tool to estimate a theoretical power of a multi-domain hybrid

architecture incorporating capabilities of several physical robots. This framework has been completed by performance criterions as a feedback on occasional robot participation in group tasks. In the first criterion, an individual robot activity is viewed in context of robot action selection strategy. It is ready to maximize a robot profit on low-level action granularity, or satisfy its goal-achievement intentions. The second criterion rates one comprehensive efficiency in group behavior as a supervising view.

Practical implementation of a multi-robot platform strictly takes advantages of bottom-up synthesis. A simulation demonstrated a constrain approach determining a robot behavior in soccer team. Its stable behavior was evolved from sample cases – if correspondent behavior leads to collective/individual behavior or idle/critical handling. Then during indoor trials physical mobile robots utilized a minimal action sets to following and foraging behaviors in a group. Their well-oriented behavior forced from basic actions in an intention context.

Poděkování

Vzhledem k rozsahu je velmi obtížné jedním poděkováním ocenit podíl mého školitele Doc. Pavla Nahodila a jeho rodiny. Rád bych také poděkoval Prof. Vladimíru Maříkovi, který významnou měrou ovlivnil mé rozhodování v průběhu vypracování této práce.

Obsah

1	Úvod	1
1.1	Problémové okruhy předkládané práce	3
1.2	Cíle disertační práce	4
1.3	Struktura předložené práce	5
2	Současný stav a rozbor možných řešení	6
2.1	Analýza možných řešení přístupy Umělého života	7
2.1.1	Historické pozadí výzkumů Umělého života	7
2.1.2	Použité základní pojmy v kontextu Umělého života	7
2.1.3	Vymezení výchozích podmínek návrhu hybridní architektury	9
2.1.4	Typové Úlohy, řešící Problémy a Cíle přístupy behaviorálně	9
2.1.5	Pohled na entity – formalizace na několika úrovních	11
2.1.6	Modulární řešení versus standardizace	16
2.2	Paralelní řešení Typových úloh v multiagentním prostředí	17
2.2.1	Reaktivní model chování	18
2.2.2	Intencionální model chování	20
2.2.3	Sociální model chování	22
2.2.4	Integrace modelů chování – hybridní přístup	24
2.2.5	Důležitá fakta a závěry paralelního řešení Typových Úloh	27
2.3	Popis paralelních dějů gramatikou	28
2.3.1	Paralelní aktivace akcí a gramatika	28
2.3.2	Paralela mezi gramatickými termíny a entitami v Kolektivu	29
2.3.3	Paralelní odvození	31
2.3.4	Důležitá fakta a závěry Popisu paralelních dějů gramatikou	32
2.4	Shrnutí závěrů a důležitých faktů Druhé kapitoly	33
3	Návrh vlastního řešení a jeho teoretický popis	34
3.1	Bloková schémata reaktivního a intencionálního modelu	35
3.2	Zvolené řešení realizované hybridním ASM	37
3.2.1	Algoritmus vytváření hybridního ASM	38
3.3	Reaktivní část hybridní architektury ASM	40
3.3.1	Interpretace prostředí a akcí v pravidlech reakční báze	40
3.3.2	Práce se soubory chování	43
3.3.3	Reaktivní ASM	44
3.3.4	Význam použitých modulů pro zpřesnění reakce	48
3.4	Intencionální část hybridní architektury ASM	49
3.4.1	Blokové schéma ASM s reaktivní i intencionální částí	49

3.4.2	Volba akce s autostimulací	50
3.4.3	Zapomínání u intencionálního modelu	51
3.4.4	Intencionální funkce popsaná eko-gramatickým systémem	52
3.4.5	Senzuální a akční vazby	53
3.4.6	Intencionální ASM	55
3.5	Adaptivní část hybridní architektury ASM	59
3.5.1	Blokové schéma hybridní architektury s adaptací	59
3.5.2	Rozbor kritérií kvality navržené akce ASM	59
3.5.3	Zpřesnění reakce adaptací	62
3.5.4	Význam adaptivních kroků při změně chování	63
3.6	Souhrn podstatných znaků navrženého hybridního řízení	66
3.6.1	Význam analytické části mechanismu výběru akcí ASM	66
3.6.2	Významová superpozice vnitřních a externích podnětů	66
3.7	Shrnutí závěrů a důležitých faktů Třetí kapitoly	67
4	Experimenty	68
4.1	Experimenty – ověření návrhu simulacemi na modelech	68
4.1.1	Simulace útočení a obrany ve srovnávací úloze Keepaway	68
4.1.2	Kritéria hodnocení	70
4.1.3	Zapojení simulačního počítačového modelu	72
4.1.4	Průběh a výsledky simulace	73
4.1.5	Zhodnocení simulací chování na počítači	75
4.2	Experimenty – ověření návrhu na reálných robotech	76
4.2.1	Obecné aspekty realizace návrhu – kalibrace měřítka	76
4.2.2	Obecné aspekty realizace návrhu – Senzory robota	77
4.2.3	Obecné aspekty realizace návrhu – Volba bazálních akcí	78
4.2.4	Společné znaky experimentů na reálných robotech	79
4.3	Srovnávací experiment	80
4.3.1	Kritéria pro vyhodnocení experimentu	80
4.3.2	Průběh experimentu	81
4.3.3	Výsledek experimentálního ověření na robotech	84
4.4	Experimentální ověření autostimulace na robotech	86
4.4.1	Kritéria pro vyhodnocení experimentu	87
4.4.2	Průběh experimentu	87
4.4.3	Výsledek experimentálního ověření	90
4.5	Experimentální ověření interakcí EG systému na robotech	91
4.5.1	Kritéria pro vyhodnocení experimentu	92
4.5.2	Průběh experimentu	95
4.5.3	Výsledek experimentálního ověření	95
4.6	Shrnutí a důsledky realizovaných experimentů s roboty	97
4.6.1	Individuální hodnocení experimentů	99
4.6.2	Hodnocení experimentů na úrovni Kolektivu	100
4.6.3	Zhodnocení simulací	102
4.6.4	Porovnání návrhu s jinými úspěšnými přístupy v oboru	102
4.6.5	Další možné rozšíření navrhovaného mechanismu	105
4.7	Shrnutí závěrů a důležitých faktů čtvrté kapitoly	107

5 Zhodnocení a závěr	108
5.1 Přehled výsledků práce včetně původního přínosu doktoranda	108
5.2 Splnění cílů disertační práce	110
5.3 Závěry pro další rozvoj vědy nebo pro realizaci v praxi	112

Přílohy

A Fotodokumentace	113
B Realizace mechanismu výběru akce	115
B.1 Algorimizace experimentů	115
B.2 Pravidla chování v experimentální simulaci	116
B.3 Pravidla chování ve srovnávacím experimentu	118
B.4 Pravidla chování pro ověření autostimulace	121
B.5 Pravidla chování pro ověření interakcí EG systému	122
C Obsah přiloženého CD	123
Rejstřík	124
Literatura	126

Seznam obrázků

3.1	Zavedení pravidel v reaktivní části řízení	40
3.2	Použití reaktivního ASM	47
3.3	Hybridní architektura řízení	50
3.4	Schéma spojení reaktivních a intencionálních podnětů v EG systému . . .	53
3.5	Příklady pohybových schémat užitých v experimentech	57
3.6	Hybridní architektura řízení s adaptací	59
3.7	Idea použití kritérií výběru akce	60
3.8	Adaptace reakce proximitních čidel při objíždění rohu	64
4.1	Simulace použitelnosti kritérií na standardní srovnávací úloze Keepaway . .	69
4.2	Navržené kritéria pro posouzení simulovaného chování.	70
4.3	Výsledky simulace navrženého hybridního ASM:	73
4.4	Statistické vyhodnocení kritérií v jednotlivých epizodách simulace	74
4.5	Konstrukce robota	77
4.6	Schéma chování v úloze sledování stěn	81
4.7	Experimentální scéna v úloze sledování stěn	82
4.8	Srovnání různých odvození navigace v úloze sledování stěn	83
4.9	Schéma chování pro ověření autostimulace	86
4.10	Pohyb robotů při ověření autostimulace	88
4.11	Kvality řídicího systému s autostimulací	89
4.12	Schéma chování v úloze párování	91
4.13	Interakce experimentálního eko-gramatického systému	92
4.14	Kvalita chování v úloze světelné navigace	94
A.1	Ukázka z videosekvencí – sledování stěn	114
A.2	Ukázka z videosekvencí – koordinační úlohy	114
B.1	Obránci a jejich taktika (intence) v různých místech hracího pole	119

Seznam tabulek

2.1	Vztah mezi gramatickými termy a entitami Umělého života	30
3.1	Korespondence mezi gramatickými entitami a entitami adaptivního ASM .	41
3.2	Příklad pravidel pro reaktivní ASM	46
3.3	Význam vazebních matic užitých v intencionálním modelu	51
4.1	Případové odvození průběhu kritérií Obránce	71
4.2	Případové odvození průběhu δ kritéria útočníka	72
4.3	Redukce vjemu do vnitřní reprezentace eko-gramatického systému	77
4.4	Testovaná chování systému	79
4.5	Experimentální ASM ve srovnávacím experimentu	80
4.6	Význam kritérií ve srovnávací úloze	84
4.7	Experimentální ASM pro ověření vlastností autostimulace	86
4.8	Experimentální části EG systému	91
B.1	Tabulka bazálních akcí použitých v simulaci	117
B.2	Rozložení intencionálního modelu po adaptaci	118
B.3	Transformace pravidel do pohybového schématu	120
C.1	Videosekvence na CD	123
C.2	Fotografie na CD	123
C.3	Dokumenty na CD	123

Seznam původních algoritmů

1	Hybridní ASM	38
2	Výběr výlučné akce v reaktivním ASM	44
3	Autostimulace v mechanismu výběru akce	52
4	Určení maxima pohybového schématu	56
5	Pokusný krok	63

Matematická symbolika užitá v této práci

$\mathcal{C}ol$	kolonie
N	množina neterminálů komponenty
T	množina terminálů komponenty
V	abeceda komponenty, $V = N \cup T$
V^*	množina všech řetězců nad V
V^+	řetězce neprázdných slov nad abecedou V
w_E	řetězec prostředí
S	axiom; startovní symbol nebo řetězec
G	gramatika
$L(G)$	jazyk gramatiky
Σ	formální parametry (<i>vázané proměnné</i>)
$\mathfrak{R}^n$	posloupnost parametrů
m	modul $m = V \times \mathfrak{R}^n$
χ	slovo $\chi \in (V \times \mathfrak{R}^n)^+$
pred, lc, rc	předpoklad, pravý a levý kontext pravidla (pred/lc/rc $\in \chi^*$)
cond	podmínka pro aktivaci
$\mathbf{C}(\Sigma)$	logické výrazy v pravidle
$\mathbf{E}(\Sigma)$	aritmetické výrazy v pravidle
ω	axiom
ψ	akční zobrazení
R	akční pravidla
P	evoluční pravidla
φ	evoluční zobrazení
$\mathbf{A}$	matice autostimulací ve vnitřní reprezentaci
$\mathbf{B}$	matice mapování senzorických vstupů do vnitřní reprezentace
W	pravděpodobnostní faktor
δ	diferenciální kritérium
ϵ	integrální kritérium
$SENSSCH$	perceptuální schéma (ohodnocení vjemu)
$ACTSCH$	akční schéma (interní ohodnocení vybuzených akcí)
$MOTSCH$	pohybové schéma (vnitřní reprezentace problému)
SCH	obecně značení perceptuálního, akčního nebo pohybového schématu

typografie

chování-1	entita chování užitá v experimentech
$akce_1$	entita bazální akce užitá v experimentech

Kapitola 1

Úvod

V předkládané práci se zabývám **návrhem a ověřením hybridní architektury řízení** spolupracujících agentů. Zde jsou představované skupinou mobilních robotů **v reálných situacích i při jejich simulacích** na počítači. Pokouším se prokázat, že u použitých robotů lze efektivně využívat i poměrně jednoduché vjemy, dostupné přímo z prostředí. Od těchto podnětů, dále doplněných o záměry a chtění (intence), lze odvodit chování robotů. Z takové chování také vychází jistá forma společenského (sociálního) chování robotů. Hybridním mechanismem vyhodnocení (ASM) právě řešené situace a výběrem následného dílčího postupu (akce), založeném jak na **reaktivním**, tak **intencionálním principu** se podrobně zabývám ve stěžejní 3. kapitole práce.

Práce je tedy svým zaměřením z oblasti výzkumu a realizací „**umělé inteligence**“. Tato vědní disciplína je zaměřena na konstrukci systémů, které při své činnosti projevují takové chování, které bychom u člověka považovali za projev jeho inteligence. Má disertační práce se opírá i o přístupy a algoritmy „umělého života“, simulující mimo jiné samoorganizující se společenství biologických entit. **Umělý život** (ALife) přináší současně *analytický přístup tradiční biologie a přístup syntetický* [Langton, 1984]. V syntetickém pojetí se více vzdaluje studiu biologických fenoménů pomocí dekompozice živých organismů. Spíše se pokouší o vysvětlení společných vlastností systémů, které se jako živé organismy chovají. Z jednoduchých mikroskopických spolupracujících prvků se generuje takové chování na úrovni makroskopické, které je možno interpretovat jako projev života.

Základními pilíři teorie umělého života jsou: **samoorganizující princip** behaviorálního řízení „**zdola nahoru**“ (bottom-up) , **emergentní jevy** a **evoluční mechanismy**. I zde prezentovaný výzkum je zaměřen na chování jednodušších umělých entit. Podstatnou částí bude definování jejich vzájemných interakcí, stejně jako jejich interakcí s okolím. Je také známo, že vzájemné lokální spolupůsobení dílčích entit vyvolává na globální úrovni zcela nový jev – *emergenci*. Z dílčích elementárních chování a vztahů se (obvykle nečekaně a bez jakéhokoliv centrálního řízení) „vynořuje“ výsledné chování obvykle složitější. Skutečná *evoluce s otevřeným koncem* musí mít pro přežití nejúspěšnějších jedinců v dynamicky se měnícím prostředí emergentní kritériální funkci. S ní mimo jiné lze na chování například reaktivních agentů ukázat na aktivity, které lze na určité úrovni pohledu označit za projevy života [Kelemen a kol., 2001].

Z přírodních věd převzalo ALife celou řadu pojmů jako jsou genetická informace, DNA, populace, křížení, mutace, přirozený výběr, generační výměna, výběr a ohodnocování jedinců apod. I v přístupech ALife platí, že *genotyp* zahrnuje všechny genetické informace, uložené v „chromozómech“ robota-agenta. *Fenotyp* pak popisuje stavbu a fyziologické vlast-

nosti agenta. Zahrnuje kombinaci vlivů prostředí a genů, které se v něm skutečně projeví. Prostředí tedy ovlivňuje jedince od samého začátku a dále v průběhu jeho života, ať už simulovaného [Kadleček, Nahodil, 2001] či reálného [Nahodil a kol., 2004].

Výběr následného evolučního kroku u jednodušších simulátorů života vychází vesměs ze zadané **kriteriální funkce** (*Fitness Function*). Fitness organismu hodnotí jeho celkovou schopnost přežít v daném prostředí. Čím vyšší je Fitness, tím má organismus větší šanci stát se vzorem pro tvorbu nové generace a šířit tak svoji genetickou informaci dál (schopnost reprodukce). Tento operátor nám simuluje mechanismus přirozeného výběru v přírodě, kde přežívají jen nejsilnější jedinci.

Klíčovým principem umělého života je **adaptace a evoluční pohled**. Evoluci lze chápat i jako neustálé šlechtění genetické informace, které probíhá ve skocích. Důvody skokových změn zdůvodnili J. Feynemann a W.D. Hillis. Pokud k evoluci dochází v konkurenčním prostředí, je vždy mnohem rychlejší a účinnější. Po období stagnace, v němž se přeskupují schémata a navenek (fenotypově) se zdánlivě nic neděje, následuje vždy bouřlivá výměna generací. V nové generaci se prosadí schopnější jedinci na úkor těch méně úspěšných. Evoluce tak zároveň prosazuje robustní a stabilní řešení a postupuje při vynalézání jaksi „zdola nahoru“. Vesměs platí, že nový, složitější postup je kombinací několika jednodušších, ověřených a fungujících. Výsledné chování má zpravidla emergentní charakter [Havel, 2007]. Tento „samoorganizující“ přístup je v protikladu k „inženýrskému přístupu“ klasické umělé inteligence. Ta, ve snaze přenášet postupně něco ze schopností člověka na agenty – roboty, řeší problémy „shora dolů“. Často s rizikem, že naslepo [Langton, 1992].

Tradiční klasická (majoritní) umělá inteligence vychází z myšlenky reprezentace světa, zatímco (minoritní) přístup převzatý z výzkumů umělého života se zaměřuje přímo na akce v tomto světě.

Předkládaný návrh řídicího systému je z kategorie hybridních architektur [Svatoš, 2001], [Svatoš a kol., 2001], [Nahodil, Svatoš, 2007], [Svatoš, Nahodil, 2008b], které využívají výhod obou směrů.

1.1 Problémové okruhy předkládané práce

Práce, tak jak jsem ji nyní přepracoval, zahrnuje nejen výsledky mého výzkumu v oblasti aplikované robotiky a umělého života coby interního doktoranda na katedře kybernetiky počátkem desetiletí, ale i mnou nově navržené a alespoň v simulacích úspěšně odzkoušené a opublikované metody z posledních dvou let [Nahodil, Svatoš, 2007], [Svatoš, Nahodil, 2008a]. Z nich si nejvíce vážím svého příspěvku, který jsem mohl úspěšně přednést na 8. konferenci *Kognice a umělý život* (KUZ VIII) v květnu 2008 [Svatoš, Nahodil, 2008b]. Ten navázal zejména na můj někdejší výzkum *Behavioural Formation Management in Robotic Soccer* [Svatoš, 2001], který jsem prezentoval v hlavním plénu na renomované mezinárodní konferenci o umělém životě *Advances in Artificial Life* (ECAL 2001). Rovněž navázal na oceněný příspěvek *Imitation Principle to Navigation in Robot Group* [Svatoš a kol., 2001] na prestižní akci *IFAC Workshop on Mobile Robot Technology* v Soulu.

V práci dále představený návrh *mého alternativního přístupu k řízení mobilních robotů* vychází z mých zkušeností, které jsem tehdy na katedře získal při řešení následujících problémů robotického fotbalu. Naznačená témata nejsou ani dnes triviální. Na jejich řešení se u nás i ve světě stále pracuje.

Problémy robotického fotbalu + získané zkušenosti

1. **Orientace robota v prostředí (herním poli).** Detailní analýza herního pole výrazně napomohla kvalitnějšímu návrhu přesně cílených akcí.
2. **Výběr akce s maximálním užitekem.** Vybrán *nedeterministický algoritmus*, protože kromě známých vstupních a omezujících podmínek je nutný i *odhad užítku akce v dalších krocích*.
3. **Koordinace akce v týmu.** Herní podmínky si vytváří tým jako celek; *koordinace je jen prostředkem* každé týmové akce, nikoliv jejím hlavním cílem.

Právě vyjmenované **Problémy** 1. — 3. robotického fotbalu byly pro mne inspirací při stanovování **Cílů** této práce (viz následující kapitola 1.2) a to i přesto, že já sám jsem výzkum v oblasti *Robocupu* posléze opustil. Se stejnými problémy se však potýká i celá řada dalších úloh, vyžadujících koordinaci. Ve své disertaci proto pracuji s **Typovými úlohami** koordinace, umožňujícími širší použitelnost mých závěrů.

Na uvedené **Problémy** se často budu odvolávat v hlavních kapitolách práce, zejména pak v kapitole Experimenty, dokládající uplatnění mého alternativního řešení hybridního řízení v reálných situacích. Navrhovanou architekturu jsem se pokusil pevně teoreticky zakotvit v Lindenmayerově gramatice, která je základem pro vytváření aktivit robotů.

1.2 Cíle disertační práce

Moje dosavadní výzkumná činnost byla dosud vždy úzce spojena s vědní disciplínou Umělého života a s mobilní robotikou. Stejně je tomu i s touto prací, v níž z předchozího využívám zkušeností s paralelizací a koordinací procesů. *Základním procesem v řízení mobilního robota* zde rozumím *dílčí - bazální akce*. Od nich (po vzoru reaktivních agentů) odvozuji pak mnohem složitější chování robota. *Chování nechť je zde plně určeno daným cílem*. Tato apriorní znalost cíle se rozkládá do mnohem jednodušších, vesměs deterministických pravidel.

Hlavním cílem mé práce je řízení a vzájemná koordinace pohybu mobilních robotů v rámci Kolektivu (definovaného v kapitole 2.1.2). Účelově se omezují na řízení dostupných laboratorních robotů. Ty dobře znám, protože na jejich vývoji jsem se sám aktivně podílel. Od počátku proto záměrně vymezuji cíl práce na řízení pohybu mobilního robota, který je *vybaven jen minimem jednoduchých senzorů*. Dostupná (silně redukováná) senzorická data budou v mém návrhu mimo jiné i překvapivě důležitým podnětem pro koordinaci aktivit robotů.

Za těchto podmínek, ve snaze pokusit se přispět odlišným úhlem pohledu k řešení výše vytypovaných **Problémů** si v této disertační práci stanovuji následující tři hlavní cíle:

Cíle disertační práce:

- Cíl 1: Analyzovat a navrhnout hybridní architekturu, opírající se jak o reaktivní, tak intencionální řízení, využívající podnětů a vjemů z prostředí ke koordinovaným aktivitám více robotů.
- Cíl 2: Navrhnout aparát, resp. i kritéria pro vyhodnocování kvality navrženého ASM (cílového chování – intencí robota), opírající se o dostupné vjemy.
- Cíl 3: Ověřit funkčnost a chování navržené architektury řízení. Chování robotů na základě vjemů namodelovat v experimentální simulaci, případně v reálném světě – ve společenství několika mobilních robotů.

Tyto cíle jsou směřované do oblasti mobilní robotiky a Umělého života. Cíle posuzuji s těmito záměry:

1. Návrh by měl odrazit mé dosavadní výzkumné práce v oblasti instinktivního řízení robotů. Jejich chování je apriorně určeno pouze zadaným cílem. Toto omezení se v realizaci rozkládá do běžně užívaných deterministických pravidel.
2. Po vzoru reaktivních agentů je vhodné vyjít z odvození složitějšího chování robota z bazálních akcí.
3. Do řešení určeného cíle se snažím zapojit více robotů. V případě společného cíle všem robotům uvažuji o kolektivu známých laboratorních robotů, na jejichž vývoji jsem se sám aktivně podílel.

1.3 Struktura předložené práce

V úvodní kapitole jsem vyšel z příkladu spolupráce robotů v robotickém fotbalu. Ukazují se na něm **Problémy** obecně spojené s řízením skupiny mobilních robotů. Jimi inspirován jsem si v kapitole 1.2 stanovil vlastní **Cíle** své disertační práce.

V dalších částech práce se věnuji nejprve návrhu konkrétních dílčích modulů, řešících zmíněné **Problémy**. Práci doplňuji experimentálním ověřením předpokladů. Výsledky celé série experimentů rozebírám a hodnotím v závěru celé práce.

Následovat bude kapitola ***Současný stav a rozbor možných řešení***. Zde přehledově připomenou způsoby řešení **Problémů** v předchozích deseti letech a ukázu, jak se tehdy navrhované metody vyvinuly do dnešního stavu. Předvedu, jak se jejich *reaktivní a behaviorální podstata* rozpadá do *neomodulárních elementů* Umělého života [Kelemen a kol., 2001]. Zvolené elementy pak algoritmizují za použití aparátu *Behaviorální robotiky*. V doméně Umělého života bude *logika následků odvozena gramatickým systémem*. Ukazují zde i důležité formalismy, které je dobré mít na paměti při návrhu syntézy **individuálního (sekvencního)** řešení a **kolektivního (paralelního)** řešení. Z nástrojů *Behaviorální robotiky* se zaměřím na efektivní skládání významově podobných závěrů v *pohybovém schématu* (*Motor-Schema*).

Ve třetí kapitole, nazvané ***Návrh řešení a jeho teoretický rozbor***, předložím svůj *vlastní návrh*, založený právě na *pohybových schématech a eko-gramatických systémech*. Zaměřím se na vlastnosti eko-gramatického systému, využitelné v mém návrhu hybridní architektury řídicího systému robota. Inspiruji se zejména *iterativním hledáním optimální reakce* a průběžným rozšiřováním poznatků o řešení na základě pokusů a omylů. Při teoretickém návrhu konkrétních dílčích modulů se vždy odvolávám na vytypované **Problémy**, respektive z nich vyplývající konkrétní **Cíle** mé práce.

Čtvrtá kapitola ***Experimenty*** je plně věnována experimentálnímu ověření teorie a navržených algoritmů; jak v simulacích, tak již přímo na reálných robotech. Podrobně popisuji a analyzuji zde úlohy navigace. Na nich postupně předvedu řešení dílčích **Problémů**, analogických těm z úvodu práce. Kapitola taktéž blíže specifikuje programové části experimentálních řídicích systémů.

Výsledky celé série experimentů rozebírám a hodnotím v 5. kapitole Zhodnocení a závěr. Zde podávám souhrnný přehled konkrétních výsledků mé práce, spolu se jmenovitým výčtem a formou faktického splnění stanovených Cílů (z kapitoly 1.2).

Zcela na závěr ještě nastiňuji možná rozpracování mnou nově navržených principů v dalších fázích výzkumu.

Práci jsem výrazně přepracoval ve snaze především o její *lepší a logicky srozumitelnější čitelnost*. Proto jsem ji doplnil nejen interaktivním seznamem obrázků, tabulek a pochopitelně i abecedně řazené prostudované a zde citované literatury, ale i přehledem zkratk a hlavně nově i jmenným rejstříkem.

Poznámka: Předkládaná práce je tentokrát psána za použití textového editoru **Microsoft Word 2003**, umožňující mimo jiné i důkladnější kontrolu pravopisu a korekci překlepů. Sazba je již dokončena standardním publikačním systémem L^AT_EX 2_ε.

Kapitola 2

Současný stav a rozbor možných řešení

Druhá kapitola přibližuje stav v úvahu připadajících přístupů různých autorů, zejména k různě pojatým architekturám mechanismu výběru akce. Po jejich důkladné analýze jsem z nich vycházel ve svém pozdějším návrhu řešení mých cílů. (Kapitola 3 a jejich ověření v kapitole 4)

Nejprve podávám stručný přehled dosavadních přístupů Umělého života a vysvětluji zde používané základní pojmy. Také se zde poměrně podrobně věnuji analýze v úvahu připadajících možností, jak pojmut Cíle stanovené v této dizertaci.

Poté se pokusím vymezit výchozí podmínky pro návrh hybridní architektury mechanismu výběru akce ASM. Vysvětlím a zavedu Typové úlohy, o něž se budu v kapitole 3 opírat při pozdějším finálním návrhu svého pojetí řešení Problémů a Cílů.

Dále vysvětluji obecné výhody paralelního řešení v multiagentním prostředí (Kolektivu). Podrobněji rozebírám možné pohledy různých autorů na problematiku začlenění zejména intencionálního modelu, k němu uvažuji integraci reaktivního a okrajově i sociálního modelu.

Velkou část této kapitoly věnuji *popisu paralelních dějů gramatikou*. Ukazují, jak různí autoři formalizují syntézu individuálního (*sekvenciálního*) řešení robota (agenta) jednotlivce a (*paralelního*) popisují projev takové syntézy v Kolektivu robotů.

2.1 Analýza možných řešení přístupy Umělého života

2.1.1 Historické pozadí výzkumů Umělého života

Zastánci klasické umělé inteligence, opírající se výhradně o znalosti, dlouho tvrdili, že intelekt a inteligence jsou zhruba tytéž pojmy a schopnost zdůvodňovat může být vyčleněna a implementována odděleně od ostatních módů myšlení. Umělou inteligenci chápali jako jednotný – unifikovatelný pochod, který lze implementovat přímo jako nějaký algoritmus – tedy sofistikovanou sekvenci logických kroků. Inteligentní systémy měly být prý proto navrhovány jako systémy, opírající se o znalosti (*Knowledge-Based*) zásadně jen od shora dolů (top-down). Už koncem 90. let se ale objevily zcela opačné názory. Podle nich nejlepším způsobem, jak navrhnout mechanismus, který myslí, už není snažit se nejdříve pochopit proces našeho myšlení. Daleko lepší je začít vyšetřováním struktury biologických systémů a chování bazálních mechanismů. Ty leží mnohem hlouběji než ty vědomé a jsou navíc evolucí ověřeny. Inteligence je vlastností každé populace. Navrhovat inteligentní systémy (Behavior-Based) se musí proto i od zdola nahoru (*bottom-up*), aby se jejich inteligence mohla vynořit (*to emerge*).

Takto nestandardně uvažující vědci vybudovali základy teorie umělého života (*ALife*) jako jakéhosi mezioborového mostu mezi evoluční biologii, etologií, sociologií a etikou na jedné straně a umělou inteligencí, počítačovou vědou a matematikou na straně druhé. Prvními „A-Lifers“, stojícími v určité opozici proti většinovým názorům byli biologové Langton a Bagleye. V architektuře počítačů to byl Hillis, v celulárních automatech Wolfram, v genetických algoritmech Holland a v programování Koza. Robotik Rodney A. Brooks z MIT kritizoval, že se v mobilní robotice uplatňuje jen klasický přístup k Umělé inteligenci, opírající se o model světa. „Svět sám je svojí nejvýstižnější reprezentací“ píše v roce 1990. Doporučil namísto „géniů, kteří nedokážou přejít místnost“ začít stavět umělé pavouky a brouky, reagující instinktivně – a vůbec začít od jednoduchého života. Inteligentní roboti tak dnes nefungují již jen podle strohých a neměnných matematických postupů, ale i na základě stále se vyvíjejících emocí. Mají své vlastní ego, potřeby, vrtochy. Učí se mimo jiné obdobnému chování, které je vlastní živým tvorům v různých prostředích a situacích [Nahodil, 2007].

Tato 2. kapitola podává přehled současného stavu a analýzu v úvahu přicházejících možností. Hned v jejím úvodu vysvětlím některé mnou dále používané pojmy. Přehled používané matematické symboliky a další zkratky jsem zařadil ještě před Obsah.

2.1.2 Použité základní pojmy v kontextu Umělého života

Předkládaná práce se opírá o základní pojmy běžné ve vědních disciplínách Umělá inteligence, Umělý život, Robotika. Pracuji v oblasti *Behaviorální robotiky*. Tato oblast zahrnuje jak odvozování chování na základě bezprostředního stavu rozpracovaného problému, tak na základě dílčích souvislostí, vypočítávaných a očekávaných v průběhu jeho řešení. V tomto smyslu se v ní aplikují etologické principy, vypočítávané v biologických formách života.

Základním pojmy v Behaviorální robotice je *Bottom-up* přístup, v němž návrh dalšího postupu vychází ze syntézy chování „od zdola nahoru“, a *Mechanismus výběru akce* – ASM, který řešenou situaci a Bottom-up přístupem ošetřuje.

Pokud se týče *agenta*, vycházím z *definice agenta – animáta podle Maesové* „... jako umělého živočicha, který může být simulován v počítači nebo včleněn do robota,

kde musí přežít a adaptovat se v měnícím se prostředí.“ Také proto se zde nezabývám detailnějším aparátem pro práci s agentem. Taková definice agenta *sjednocuje širokou oblast zainteresovaných disciplín od umělé inteligence a robotiky po evoluční biologii, etologii, sociologii a etiku* [Maes, 1990, Nahodil, 2007].

Ze základních prvních čtyř vlastností agenta podle [Wooldridge, 2002] bude pro mou práci převažující zejména jeho *reaktivnost*.

Při použití pojmu *Robot* se omezují jen na *Mobilní roboty* a jejich *reaktivní řídicí mechanismy pohybu*. všechny další, mnou odlišněji používané pojmy se snažím stručně charakterizovat vždy hned v místě jejich prvního výskytu.

Nyní zde uvedu ještě některé další, mnou **v práci často používané pojmy**:

- **Koordinace**: slouží zde především jako *prostředek k prosazení nadřazeného cíle*. *Procedura koordinace* sestává z percepce jednoho jevu více roboty a ze synchronní odezvy na něj.
- **Kolektiv**: rozumím zde jako *skupinu mobilních robotů se společným nadřazeným cílem*. Interakce jednotlivců s kolektivem je opět na základě jejich chování. Prof. Parkerová připomíná, že ***vjemy v kolektivu*** je třeba chápat ještě těsněji – na jednu stranu jednotlivec kolektivní *cíle přijímá a plně se jím podvoluje* (Parker proto používá termín ***Acquiesce***). Na druhou stranu jednotlivec může svojí *netrpělivou, unáhlenou* individuální akcí výrazně ovlivnit trendy kolektivní (Parker používá pro to termín ***Impatience***) [Parker, 1993].
- **Distribuované řešení**: upřednostňuji přirozené členění kolektivního problému na dílčí podúlohy. Nová podúloha vzniká okamžikem zapojení dílčích schopností robota do problému [Brooks, 1999].
- **Sociální inteligence**: je chápána *v souvislostech s etologickým pojetím*, stejně tak, jak je pozorována v přírodě u ptačího hejna, včel, mravenců atd. [Lorenz, 1993].
- **Lindenmayerovy systémy (L-systémy)** slouží jako nástroj na analýzu celého procesu morfogeneze. L-systémy ukazují na logiku uvnitř přírodního formování do určitého cílového stavu. Umožňují predikovat postupný růst informačního systému, jeho vývoj a specializaci v rámci organismu. V Lindenmayerových systémech je už každý element brán jako organismus [Lindenmayer, 1968a], sloužící za základ vyšším organismům.
- **Gramatika**: v práci je dominantní *nedeterministická gramatika*. Zhruba lze říct, že gramatika omezí možnosti systému shora. *Generativní síla gramatiky* přímo vymezuje výčet všech možných výsledků. Generativní síla tedy shora omezuje i stavy dosažitelné v kolektivu robotů [Kelemen, Kelemenová, 1992].

Právě výše uvedené pojmy Umělého života pro mě zároveň vymezily různé, mnohdy nestandardní pohledy na mnou v této práci prezentovaný hybridní řídicí systém. S jednotlivými pohledy seznámím čtenáře v následujících odstavcích.

2.1.3 Vymezení výchozích podmínek návrhu hybridní architektury

Již ve své práci „Behavioural Formation Management in Robotic Soccer“, přednesené na prestižním fóru ECAL 2001 *Advances in Artificial Life* jsem ukázal, že **skládáním poměrně velmi jednoduchých alternativ** lze docílit výsledného — překvapivě mnohdy i značně sofistikovaného — chování robota [Svatoš, 2001]. *Ve dvou úrovních* mohu postupně vysledovat

Kroky, vedoucí na optimální řešení předloženého Problému:

- Na *úrovni detailu* by měly být nejprve ošetřeny *lokální podněty*.
- Poté na *úrovni kolektivu* budu sledovat, zda a jak jeho jednotliví členové vykrývají celou problémovou oblast scény.

Při řízení robota se těsné propojení obou úrovní užívá v oblasti robotiky, založené na *implementaci ucelených chování*. Do této oblasti behaviorální robotiky (Behaviour Based Robotics) jsem nasměroval i svůj návrh řídicího systému. Zdůrazňuji požadavek úzké souvislosti kroků na obou úrovních. Účinek kroků na jedné úrovni se zprostředkovaně odrazí v odpovídajícím chování druhé úrovně.

Princip interakce na více znalostních úrovních přiblížím spolu se základními postuláty Umělého života. Jedním z nich je *decentralizace řešení v rámci kolektivu*. V první řadě bude vlastní lokální orientace robota hodně záviset na reprezentaci okolních stavů — fakt ve formě pravidel. Zaměřím se proto na *pravidlové systémy s minimální reprezentací*. Povětšinou si takový systém pro bezprostřední reakci vůbec nemusí držet kontext. Musí pouze rychle rozpoznat své vlastní neefektivní reakce a podle nich okamžitě změnit rozhodovací strategii. O to větší důraz se proto klade na dobrou interpretaci dílčích senzorických vjemů. Obdobně je třeba správně interpretovat na jednotlivých znalostních úrovních i cíl kolektivu.

Behaviorální pohled se zde jeví jako vhodný způsob hodnocení cílového problému. Umožňuje při řízení robota lokálně přizpůsobovat jeho taktiku, v závislosti na aktuálním stavu rozpracovaného řešení. *Integrace obecněji vnímaného chování spolu s dílčími reakcemi* je proto také mým dalším záměrem, jak řešit **Cíle**. Mnou navržený a dále v této práci popisovaný hybridní systém bude založen na interakcích detailní (nejnižší) úrovně znalostí s ohledem na vzájemné chování mezi jednotlivými členy kolektivu. Právě takto zde *vymezený výchozí stav hybridní architektury* je pro mě základem — výchozím předpokladem efektivního návrhu. Proto se dále zabývám především *zpřesněním reakce agenta – robota externími vlivy*. Každým zpřesněním může hybridní řešení napomoci právě v řešení zmíněných dílčích problémů. Svoji, pokud vím zcela *původní variantu hybridní architektury*, prezentovanou úspěšně mimo jiné i na Workshopu IFAC v Soulu [Svatoš a kol., 2001] pak podrobněji popíšu ve 4. kapitole.

2.1.4 Typové Úlohy, řešící Problémy a Cíle přístupy behaviorálně

Většina z **Problémů** a z nich stanovených **Cílů**, na které se v disertaci zaměřuji již od úvodní kapitoly, má dnes více (paralelních) řešení.

Typové Úlohy členěné podle přístupu k vytčeným Cílům disertační práce:

- A. pokud *chci zorientovat robota v herním poli*, tedy i v kontextu ostatních robotů, podívám se jak je postavena a jak se řeší typová úloha **Distribuce zájmových oblastí**,
- B. pokud *chci vlastní výběr akce zaměřit na maximální užitek*, podívám se, jak toto řeší typové úlohy **Oportunistického využití prostředí**,
- C. a pokud *chci koordinovat akce v týmu*, mohu přímo vycházet z poměrně rozsáhlé skupiny typových úloh **Koordinovaného pohybu**.

Nastíněné základní členění přístupu k **Cílům** se snažím dodržovat v celé další práci. Dále o něm hovořím jako o **Úlohách**.

Zde chci zdůraznit společný jmenovatel uvedených paralelizovatelných **Úloh** s přístupy Umělého života: jsou to iterativní procesy. Omezení na jeden krok výrazně redukuje výpočetní nároky pro řízení kolektivu robotů, ze své podstaty NP-úplné [Parker, 2000]. S počtem případů bude náročnost analýz řešení narůstat exponenciálně. Je proto výhodné vyhradit čas na komunikaci mezivýsledků v rámci kolektivu i po každém kroku – výsledkem kroku jsou nové předpoklady pro další paralelní zpracování.

A. Typová úloha Distribuce zájmových oblastí

V typové úloze Distribuce zájmových oblastí se jedná o přidělování základních (dále bazálních) akcí mezi konečně velká společenství mobilních robotů. Prof. Parker například zakládá řešení distribuce zájmových oblastí na nepřímé (implicitní) komunikaci mezi roboty [Parker, 2000]. Podle jejího postupu pokrytí problémové oblasti se roboty nejdříve rozptýlí po vymezené ploše této oblasti, tak aby vzdálenost mezi nimi byla maximální možná.

B. Typová úloha Oportunistického využití prostředí

Ačkoliv úloha **Oportunistického využití prostředí** vychází z deterministického zadání, nakonec i řešení tohoto problému lze řešit přístupy Umělého života. Úlohu dělá NP-úplnou započtení omezujících podmínek v realizaci, tedy hledání maximálních nezávislých množin. Reálné prostředí řešení ještě více zkomplikuje. V něm je efektivnější iterativně volit pouze dílčí kroky.

C. Typová úloha Koordinovaného pohybu

Úlohy **Koordinovaného pohybu** demonstrují pro jednoduchost na geometrickém shlukování ve dvourozměrném prostoru například s Eukleidovskou metrikou jako mírou nepodobnosti; protože pokud budou omezující podmínky geometrické, i na geometrickém shlukování bude vidět přínos, jaký pro koordinování akce v týmu dává znalost postavení jednotlivců ve shluku. A z pohledu Umělého života vidím ale také hodnocení opačné — totiž, že *chování jednotlivce je prostředím také omezeno*.

Například u Balche je každá změna stavu ohodnocena cenou či přidělenou kapacitou. Jednotlivec pak volí svůj postup s ohledem na možnou kapacitu celého přechodu od počátečního stavu k cílovému stavu kolektivu. [Balch, 2000] ukázal, jak se geometrie pohybu celého kolektivu rozkládá do více paralelních trajektorií. Každý jednotlivec se projevuje

vlastní trajektorii. Dílčí trajektorie tak v lokálním kontextu interpretuje podmínky (pravidla) pro geometrii pohybu celého kolektivu.

2.1.5 Pohled na entity – formalizace na několika úrovních

Na širokém paralelismu (paralelním řešení dílčích akcí) je postaven celý Umělý život. Pravidla, pozorovaná v biologických (uhlíkových) formách, jsou pomocí přístupů Umělého života [Langton, 1992] nahrazena ekvivalenty informačními (nyní vesměs na silikonové bázi). Informační struktury takto získané se využívají zejména pro paralelní vyhodnocení systému ve vztahu k bezprostředním strukturám v jeho okolí.

Systém modifikuje prostředí intenzivními interakcemi. Zpětně se na stavu systému projevuje stav prostředí, modifikovaný dalšími paralelně pracujícími entitami.

Pojem entita – zde i v dalším nechť shrnuje *vlastnosti, které se v objektově orientovaném řešení uplatňují společně* – jsou spojeny se stejným objektem. Uvažování v entitách mně napomůže při analýze a následné reprezentaci problému v počítači.

Entity Umělého života se **formalizují obvykle na několika úrovních detailu**. Je to jistá obdoba postupné fokusace, známé z přístupů Umělého života. Například nejprve zaostřím svůj pohled na les, poté detailněji na strom a ještě hlouběji na listy stromů. v robotice je užitečné se zaměřit na tři úrovně detailu:

1. **Bazální akce.** Označuji tak elementární akce robota, a to na nejnižší úrovni dekompozice. Tyto elementární akce je nutné použít při řešení každého složitějšího problému. Bazální akce představuje modul, ve kterém se definuje rozhraní pro spuštění elementární aktivity robota. Rozhraní zahrnuje parametry takové aktivity (jak se spouští) a programový kód (jak se v konkrétním případě realizuje).
2. **Agent.** Opírám se definici agenta – animáta podle Maesové jako umělého živočicha, který může být simulován v počítači nebo včleněn do robota, kde musí přežít a adaptovat se v měnícím se prostředí. Také proto se zde nezabývám detailnějším aparátem pro práci s agentem. Taková definice agenta sjednocuje širokou oblast zainteresovaných disciplín od umělé inteligence a robotiky po evoluční biologii, etologii, sociologii a etiku [Maes, 1990], [Nahodil, 2007].
Z vlastností agenta, které uvádím dále [Wooldridge, 2002], bude pro mou práci dominantní především jeho reaktivita (podrobněji viz dále). Ve zde prezentovaných úlohách má agent při řešení problému roli aktéra. Aktér volí a spouští akci prostřednictvím své vlastní sady bazálních akcí.
3. **Kolektiv.** Tato entita představuje rámec pro vyjádření relací, které přirozeně plynou ze zapojení více agentů do (paralelního) řešení společného problému.

Mé řešení je zaměřeno především na Kolektiv (pohled na nejvyšší 3. úroveň). Pro návrh řešení je už třeba se zaměřit na vlastnosti potřebných konstrukčních prvků. Uplatnění agentů a bazálních akcí v kolektivu uvádím níže.

A. Pohled z úrovně bazální akce

V Umělém životě je nutné se věnovat samostatnému modulu bazální akce pro jeho mnohostranné použití. Jednotné rozhraní umožňuje modulární řazení bazálních akcí do jednoho uceleného řešení.

Vlastnosti bazální akce: *Bazální akce* představuje modul, který mění stav rozpracovanosti řešeného problému

1. jeho funkce je determinována vstupem, představuje v řešené úloze nejjednodušší generativní zařízení
2. znalost o ošetření události se ukládá jako relace mezi popisem události a bazální akcí.
3. dílčí znalost se v průběhu řešení obvykle využije v kontextu širším, než očekávaném.

B. Pohled z úrovně agenta

Entitou na výrazně vyšší úrovni je poté agent. Pro kontinuitu výkladu a úplnost zde uvádím jen definici inteligentního agenta (IA), jak ji formuloval Wooldridge v devadesátých letech minulého století [Wooldridge, Jennings, 1995]:

Agent obecně je hardwarový nebo mnohem častěji softwarový počítačový systém, který vykazuje řadu charakteristických vlastností, především **Vlastnosti inteligentního agenta 1.– 4.** Wooldridge tyto klíčové znaky, které odlišují inteligentního agenta od jiných softwarových aplikací v roce 2002 ještě rozšířil [Wooldridge, 2002]. Uvedu nyní jím formulované vlastnosti, které inteligentní agent může vykazovat:

Vlastnosti inteligentního agenta:

1. **Autonomnost (*Autonomy*)**. Inteligentní agent je schopen jednat samostatně a řídit své vnitřní stavy
2. **Reaktivita (*Reactivity*)**. Reaktivní agent vnímá interaktivně své okolní prostředí a reaguje včas na změny probíhající v něm, aby plnil cíle, pro něž byl navržen.
3. **Proaktivní přístup (*Pro-activeness*)**. Reakce agenta na podnět z okolí neurčuje pouze momentální stav, ale jeho chování je řízené cíli v tom smyslu, že agent si na základě své vlastní iniciativy vytváří cíl a může sám převzít iniciativu při řešení
4. **Sociálnost (*Social Ability*)**. Agent je schopen interakce s ostatními agenty (případně i s lidmi) prostřednictvím společného jazyka (pro komunikaci agentů) a má tak i možnost s nimi kooperovat. To se někdy nazývá i *komunikační schopnost agenta*, umožňující, aby agent dostával potřebné informace z různých zdrojů.
5. **Schopnost spolupráce (*Capacity for Cooperation*)**. Inteligentní agent pro dosažení cíle spolupracuje s ostatními agenty
6. **Schopnost zdůvodnění (*Capacity for Reasoning*)**. Inteligentní agent může vykazovat schopnost vyvozovat další závěry z aktuálních znalostí a vlastních zkušeností.
7. **Adaptivní chování (*Adaptive Behaviour*)**: tj. když se inteligentní agent učí nebo mění své chování na základě předchozích zkušeností.
8. **Důvěryhodnost (*Trustworthiness*)**: Uživatel agenta (*user*) si musí být vysoce jistý, že agent bude působit k jeho prospěchu a bude mu vždy poskytovat pravdivé informace.

Wooldridge ve své často citované zprávě [Wooldridge, 2002] zavádí některé další *občas diskutované vlastnosti a činnosti agentur* ve smyslu Minského definice [Minski, 1985]. Jsou to:

- **Mobilita (*Mobility*)**: schopnost agenta pohybovat se v rámci (např. elektronické) sítě.
- **Věrohodnost (*Veracity*)**: zda nebude její agent vědomě poskytovat nepravdivé informace.
- **Benevolence (*Benevolence*)**: v posouzení, zda a jak mnoho je užitečné držet se i těch cílů (conflicting goals), s nimiž se agenti dostávají do konfliktu.
- **Racionálnost (*Racionality*)**: zda agenti podřizují svoji aktivitu dosažení svých cílů a zda nebudou po rozvaze úmyslně zamezovat svým faktickým jednáním těm cílům, kterých se má (podle člověka) v dalším dosáhnout.
- **Učení/adaptace (*Learning/Adaptation*)**: Jedná se o vlastnost, kdy agenti v průběhu času svoji výkonnost zlepšují (obvykle učením a adaptací).

Pokud se týče *agenta*, vycházím z **definice agenta - animáta podle Maesové** „... jako umělého živočicha, který může být simulován v počítači nebo včleněn do robota, kde musí přežít a adaptovat se v měnícím se prostředí.“

Taková definice agenta sjednocuje širokou oblast zainteresovaných disciplín od umělé inteligence a robotiky po evoluční biologii, etologii a sociologii [Maes, 1990, Nahodil, 2007]. Také proto se zde detailnějším aparátem pro práci s agenty již nezabývám. Ze základních prvních čtyř vlastností agenta, viz Vlastnosti inteligentního agenta, bude pro mou práci převažující zejména jeho **reaktivnost**. Při použití pojmu *Robot* se omezují jen na *Mobilní roboty* a jejich reaktivní řídicí mechanismy pohybu.

Pohled na interakce a modularitu bazálních akcí

Použití *bazálních akcí pro reprezentaci základních funkcí*, respektive *agentů pro reprezentaci výkonných entit* (bazálních akcí a agentů), je voleno obvykle ve snaze o *maximální modularitu*. Záměrem je variabilní interpretace složité funkce za pomoci předem známých *výkonných entit*. Pro modularitu je proto důležitá znalost o interakcích výkonných entit a důsledku každé interakce na řešení. Výhodou *modularity* je redukce složitosti návrhu algoritmu. V případě modularity je celý *proces rozhodování postaven na syntéze výkonných entit*, mezi kterými probíhají:

- **Explicitní interakce**. Jsou zaimplementovány v mechanismu výběru akce agenta tak, aby se v řešení uplatnila *nejkvalitnější bazální akce*.
- **Implicitní interakce**. Agenti jsou schopni zprostředkovaně ovlivnit i proces rozhodování ostatních. Nejčastěji mění podmínky mechanismu výběru akce.

Dle [Kelemen a kol., 2001] za syntézou paralelní aktivity ani nemusí stát mechanismus výběru akce. Úměrně poměru implicitních interakcí proti explicitním se také chování Kolektivu jeví jako více či méně koordinované.

Prozatím se ve svých záměrech (intencích) omezujeme jen na pohyb robota v kartézských souřadnicích. (Nejedná se zde tedy ještě např. o myšlenkové pochody (záměry), kdy ale i zcela imobilní jedinec může svými závěry a návrhy dalších kroků zasáhnout do vývoje.

K realizaci paralelně postaveného modulárního řešení lze přistupovat:

- **Konekcionisticky.** V konekcionistickém systému *převažují implicitní interakce*. Na nich je celé řešení postaveno. *Sekvenční procedura* zastoupená elementem může být zcela *triviální*. Zato abychom ale mohli ve výsledku pozorovat vlastnosti explicitně neinicializované (neřiditelné), musí být *zapojen velký počet výkonných entit*.
- **Reaktivisticky.** Řešení problému se opět složí z *elementárních příspěvků*, tentokrát generovaných agenty – specialisty jen na omezený počet funkcí. Jednotlivci uplatní danou funkci pro řešení společného problému, jakmile k tomu mají příležitost. Svým zásahem mohou *připravit podmínky pro uplatnění ostatních specialistů*.

Modularita bazálních akcí a modularita agentů by se neměla explicitně směřovat. Tento problém si je třeba uvědomit hlavně v typových Úlohách Koordinovaného pohybu.

Explicitní interakce je třeba držet v odpovídajících úrovních detailu — na úrovni bazálních akcí a agentů. Díky *implicitním interakcím* je jedna a táž událost reflektována v obou úrovních detailu současně. Podrobněji jsou tyto interakce vysvětleny v následující části 2.1.5.

Z praktického hlediska je podstatné, že *systém je říditelný* (ve smyslu teorie řízení) většinou z úrovně bazálních akcí. Na této úrovni ošetříme detaily příčin aktuální události. Na úrovni Agentů a Kolektivu pak bude *pozorovatelné*, jak akce Agentů mění stav. Zpětně nový stav reflektují ostatní agenti při řízení, jak se dle svých omezených znalostí přiblížit k žádané hodnotě (k cíli). I společenství postaví hodnocení své strategie na drobném testu jednotlivců.

Výhodou řešení problému v několika úrovních detailu je redukce výpočetní složitosti algoritmu. Účelně se příčiny vyhodnocují pouze do hloubky pokryté v jedné úrovni detailu. Mechanismus výběru akce agenta **ASM** (*Action Selection Mechanism*) **pracuje nad úrovní bazálních akcí**, s konečným počtem jeho omezených znalostí.

Řešení úlohy *Koordinovaný pohyb* pak vyžaduje *heuristiku*. Relace mezi agenty jsou totiž v zásadě nedeterministické.

Pohled integrující neomodularitu a explicitní interakce

V kapitole 2.1.5 jsem zdůraznil, jak pro moji další práci budou důležité právě *explicitní interakce*, jimiž se v podstatě definuje odvození aktivity agentů-robotů z dostupných podnětů. Logika explicitních interakcí může být obecně zakotvena vhodným softwarem v *Mechanismu výběru akce* (ASM). Nyní uvedu, jak je takový mechanismus výběru akce realizován na jiných pracovištích, zabývajících se výzkumem Umělého života. Tak například Prof. Kelemen ukázal [Kelemen, 1989], že ke společné práci agentů nad jedním stejným (sdíleným) cílem Kolektivu není nezbytná hierarchická organizace. Běžný paralelní ASM opravdu nepotřebuje takovou hierarchii, ve které agentům řídící agent explicitně algoritmem, či heuristikami předurčí, co mají dělat.

Nezbytná *omezení cíle probíhají právě implicitními interakcemi* které vyplynou z *kommunikace mezi dvěma či více agenty – roboty. A nejenom mezi agenty, ale i mezi agenty*

a prostředím. Mechanismus ASM v Umělém životě tyto implicitní interakce reflektuje velmi jednoduše. Bere v úvahu jen to, co zná o daném problému lokálně. Tento lokální kontext vymezuje obecně data, podle kterých se ASM řídí. Mechanismus ASM obsahuje obvykle sadu pevně deterministicky daných příkazů, je schopen s nimi obsloužit i čistě náhodné události. *Význam datového řízení navíc není jen kvantitativní, ale i kvalitativní.* Lépe lze nalézt globální optimum řízení, protože se objevují stále nové, obecně náhodné podněty, které řešení nakonec vyvedou z kvalitativně horšího lokálního optima [Takadama a kol., 1999].

Důležitý je *v iterativním procesu přenos znalostí*. Proto také nezbytnou explicitní znalostí dále zůstává *zahrnutí příčin a následků do mechanismu výběru akce ASM*. Dobře se takové znalosti o *příčinných souvislostech* mezi entitami v ASM udržují jako například **atraktivní trendy**. Obdobně kvalitní výběrový mechanismus proaktivně předchází konfliktům – dotyčné souvislosti interpretuje jako **repulsivní trendy**. Repulsivní trendy jsou směry, které jsou pro další pokračování iterativního procesu nevýhodné. U kvalitního mechanismu ASM, který navrhla nedávno Prof. Chernova, blízká spolupracovnice Prof. Arkina je *znalost o atraktivních resp. repulsivních trendech uložena implicitně přímo v sekvenci následků jednotlivých chování* [Chernova, Veloso, 2007]. Pokud se entity *mechanismu výběru akce ASM* používají iterativně, vykazují neomodularitu. Tento pojem zavádí a vysvětluje Stein už v roce 1994 [Stein, 1994].

Neomodularita, kromě standardních vlastností modularity (variabilním složením složité funkce z výkonných entit), zapojuje ještě závěry nižší a vyšší úrovně. Na příkladu uváděném v kapitole 2.1.3) lze interakce z nižší na vyšší úroveň ukázat na růst listů. Zpětně ho pozorujeme na růstu celého stromu. Obdobně lze ukázat příklad závěrů (interakce) z vyšší na nižší – mimo jiné když růst lesa zastíňuje jednotlivé stromy.

Vlastnosti neomodularity postupně vyplývají vždy až v průběhu času [Stein, 1994]. Toto je vidět na tom, jak se tatáž situace průběžně odrazí v procesech, které jsou s neomodularitou spojovány.

Podle Stein jsou **nejčastější procesy spojované s neomodularitou**:

1. **Imaginace.** Imaginační doporučení posiluje informace z nižších úrovní detailu. Imaginace probíhá na pozadí s nižší prioritou. Agent založí své rozhodnutí na imaginaci až v okamžiku nejistého závěru.
2. **Relační reprezentace.** Implicitní vazby se snažíme co nejvíce přiblížit v relačním modelu — ve vztahu agenta k okolním významným objektům (další agent či statická překážka).
3. **Inkrementální adaptace.** Adaptace ASM se projevuje přeformulací výchozího chápání problémů podle aktuální situace.

Neomodulární řešení s výhodou využívá různorodost podmínek, na kterých se entita ASM může uplatnit. Celkovou kvalitu řízení mohou proto zlepšit i negativní výsledky interakcí. Přínos neomodulárního řízení, realizovaného v ASM, totiž významně zvyšují *konfliktní procesy*. Jimi se vymezují *repulsivní trendy* (viz výše). Tyto dodatečné negativní podněty musí ASM také plně respektovat.

2.1.6 Modulární řešení versus standardizace

V pohledu dvou znalostních úrovní lze explicitní a implicitní interakce omezit dvěma proti sobě postavenými kritérii. Stejná kritéria najdeme i v metodách optimalizace chodu hybridního systému. *Obecně bohužel při výzkumu v oblasti mobilní robotiky není měřitelným kritériem věnována dostatečná pozornost.* Nedostatek univerzálního kritéria je ještě výraznější ve hraniční oblasti mezi robotikou a Umělým životem. Můžeme tento stav ovšem chápat i jako jeden z přirozených důsledků nedeterminismu v Umělém životě.

Z mnou prostudovaných metrik, vhodných také pro ohodnocení různých řešení **Typových Úloh** mé kapitoly 2.1.4, jsem narazil na Donaldova poměrně letitá kritéria pro hodnocení koordinovaného pohybu [Donald a kol., 1997]. Jím definovaná kritéria jsou určena pro pohyb kooperujících agentů. Podstata hodnocení zmiňovanými kritérii je založena na distribuci.

Použití zmíněných invariantů nicméně podle mne není nejvhodnější pro případ této práce, protože:

- Na rozdíl od posuzování kooperace mi pro potřeby řešení dříve formulovaných **Problémů** postačí ohodnotit *kvality předvedené koordinace. Kooperace robotů ani není mým cílem.*
- Způsob zavedené distribuce u Donalda *nepostihuje možné emergentní jevy*, pro něž jsem si především zvolil přístup Umělého života (přístup směrem od *zdola–nahoru*).

Z těchto důvodů pro potřeby mé práce využiji vlastních kritérií v kapitole 3. *Donaldova kritéria mi poslouží pouze pro srovnání s nejčastějšími řešenými úlohami navigace ostatními výzkumnými pracovišti. V oblasti behaviorální robotiky není principiálně možné se snažit o standardizaci.*

Shrnu-li:

V této podkapitole 2.1 jsem porovnával *důsledky implicitní a explicitní interakce* mezi řídicími a výkonnými entitami. Oba dva typy interakcí by se v mém vlastním návrhu měly uplatnit. Můj návrh ASM je prioritně zaměřen na reaktivní přístup se zařazením intencionálního modelu, který umožní respektovat při výběru další bazální akce rovněž intence (záměry, chtění). Popisuji zde i přístupy, zavádějící do ASM adaptivitu. I to je mým záměrem.

Bazální akce jsou předmětem výběru akce na úrovni agenta. Aktivita jedince se ale současně promítnou do řešení, navrhovaného pro celý Kolektiv. Na podkladě prostudované literatury i svůj vlastní návrh proto povedu ve *třech zdůrazněných úrovních detailu*. Dále jsem zde zdůraznil praktické aspekty práce s entitami Umělého života: funkci mechanismu výběru dílčí akce a analogické nakládání s podněty a vnitřní reprezentací řešeného problému. Také ve svém návrhu chci *minimální reprezentaci* doplňovat dodatečnými podněty pro upřednostnění konkrétní akce. *Při inkrementální adaptaci pravidel pro odvození nejvhodnější akce už užiji neomodulární techniky.*

2.2 Paralelní řešení Typových úloh v multiagentním prostředí

Efekt paralelního řešení (parallelismu) zmíněných problémových úloh hodně závisí na tom, jak agenti mezi sebou interagují v reálných situacích. Dnes už jsou interakce mezi jednotlivými agenty propracované mnohem podrobněji v metodice multiagentního systému (MAS). Principy MAS jsou použitelné i na tak specifického agenta s vlastní fyzickou realizací, jakým je robot [Weiss, 1999]. Proto je možné MAS metodiku použít i ve zmíněných Typových Úlohách (v podkapitole 2.1.4) tak, že

- A. pokud jde o úlohy Distribuce zájmových oblastí, dobrých výsledků lze dosáhnout díky proaktivitě intencionálních agentů,
- B. úlohy Oportunistického využití prostředí jsou vlastní reaktivním agentům
- C. a úlohy Koordinovaného pohybu budou řešit sociální agenti.

Tohoto členění agentů do tří kategorií (intencionální, reaktivní a sociální), které zavedl [Wooldridge, Jennings, 1994], se budu držet při výkladu této kapitoly. Členění poslouží i jako základní orientace v paralelismu řešení Problémů použití mobilního robota v jakékoliv poziční dynamické hře. Na každém z nich se mohou uplatnit různé vlastnosti agenta. Jejich klady v následujícím krátkém shrnutí označuji (+), nevýhody (–).

A. Kde a jak bude agent intencionální

Intencionální agent je schopen ošetřit implicitní interakce tím, že do svého plánu zahrne další alternativy, ze kterých lze případně odvodit i požadovaný cíl. Naopak k těmto alternativám se při striktně vymezených pravomocích bude hodit podrobnější plán intencionálního agenta.

Shrnutí vlastností:

- (+) v plánu lze minimalizovat odhadovanou cenu postupu k cílovému stavu. Současně díky nově získaným informacím – vjemům z okolního prostředí lze plán průběžně přizpůsobit nové situaci.
- (–) Stejně tak je třeba aktualizovat model prostředí, ze kterého se plán odvozuje. Přitom jde jen o způsob, jak v Kolektivu zřetelná fakta přiblížit konkrétnímu Agentovi.

B. Kde a jak bude agent reaktivní

Reaktivní agent se dobře uplatňuje v úlohách, na jejichž řešení se podílí kolektiv s velkým počtem agentů [Arkin, 1997]. Co rozumím „velkým počtem“, závisí na aktuální situaci. Například v kapitole 4 ukážu situace, kdy postačí dokonce pouze čtyři roboti.

Společné instinktivní chování kolektivu vychází z autonomních reakcí [Ferber, 1996]. Zakódované autonomní reakce lze chápat i jako způsob reprezentace problému.

Shrnutí vlastností:

- (+) Celá znalost o aktuální situaci se zjednoduší na určení ceny reakce. Cíle agentů jsou už explicitně přednastavené v pravidlech jako terminální stavy. Zbývá je srovnat s implicitními fakty – vjemy z prostředí.
- (–) Reakční pravidla musí být dostatečně obecná. Natolik, aby pokryla i nespecifikované implicitní vazby mezi rozhodnutím jednotlivých agentů.

C. Kde bude agent sociální

Sociální agent může využít svých znalostí o modelu chování ostatních členů kolektivu. Fyzickému agentovi se taková znalost hodí při předvídání konfliktů — koordinace po vzájemném vyjednávání. Synchronizace využití společných zdrojů je také nejčastější důvod ke koordinaci aktivit fyzických agentů.

Shrnutí vlastností:

- (+) Protokolárně je výměna takových znalostí jednoznačně formalizovaná, pro komunikaci existuje celá řada protokolů.
- (–) Lokální interpretace sdílené znalosti už není deterministická. Jednu znalost si každý z agentů může interpretovat jinak.

Uvedené obecné vlastnosti přiblížím na příkladech přímo z mobilní robotiky. Podívejme se, nakolik multiagentní formalismy napomáhají v řešení Problému.

2.2.1 Reaktivní model chování

Předchozí hrubý přehled použití různých typů agentů kladl důraz na reprezentaci problému. V reaktivním modelu můžeme vzít přímo výběrový mechanismus, kde ASM vybírá ze sady přímých implikací podle vztahu:

$$Vjem \rightarrow Akce \quad : \quad váha \quad (2.1)$$

Výběrový mechanismus pak z celé sady vybírá kombinaci nejvhodnější právě pro danou příležitost. Jestliže pravidla používáme statisticky významně často, pak jejich váha odrazí četnost použití tohoto pravidla. To nastane za konkrétního předpokladu (ve vztahu (2.1) je to *Vjem*). Tímto vjemem je v kapitole 4 například detekce nárazu nárazníkem.

Proto se v robotice více než kde jinde při nastavení váhy pravidel reaktivního modelu chování využívají pravděpodobnostní přístupy. V robotice byly použity jak přímo Bayesovské metody odhadu vah [Kristensen, 1995], tak i fuzzy logika relací od *Vjemu* k *Akci*. Například v senzoricky popsaném světě RoboCupu je poměrně robustní fuzzy řízení založeno na fixaci vůči referenčním stavům [Staffiotti a kol., 1995].

U reaktivních agentů navíc uijeme dalších interakcí. Protože ale pracujeme s entitami Umělého života (v kapitole 2.1.5), v rámci aktuálního modelu bude možné paralelně generované aktivity i koordinovat. Stačí si k tomu uvědomit, že i složitější implikace jsou jen výsledkem spojení mnohem jednodušších akcí. Koordinace pak závisí na tom, jak agent dobře reflektuje implicitní interakce. Arkin [Arkin, 1989a] uvádí: „...váhy klíčových rozhodnutí je třeba náležitě podložit aktuální potřebou kolektivu.“ Praktické důsledky extenzivnějšího

skládání akcí do kolektivního chování [Goldberg, Matarić, 1997] nejsou též ničím novým. Předvedená koordinace je v takových případech založena na implicitních interakcích mezi reaktivními agenty.

Návrh reaktivního agenta je ovšem postaven nyní již na explicitní části ASM jednotlivého agenta. V reaktivním modelu se jedná o skládání aktivních pravidel (2.1). v podstatě jsou použitelné dvě nejčastější možnosti spojení bazálních akcí:

1. Ve **vrstvené architektuře** (*Subsumption Architecture*) jsou konkurenční akce buzeny paralelně, ne všechny úrovně akcí jsou však současně aktivní. Nesouhlasné akce se vzájemně převrstvují. Ve výsledku se pak projeví elementy majoritní právě v daném okamžiku.
2. V **pohybovém schématu**, o které se opírá můj návrh, dojde k superpozici vah dílčích pravidel. z tohoto souhrnného ohodnocení mechanismus ASM jednoznačně určí lokální rozhodnutí. Smyslem každého lokálního rozhodnutí je zvolit akci, kterou se robot přiblíží cílovému stavu. Běžně v každém takovém lokálním rozhodnutí robot postupuje ve směru gradientu ohodnocovaného pohybového schématu. v takovémto rozhodnutí se posuzují stavy robota pouze do vzdálenosti jednoho kroku.

A. Vrstvená architektura chování

V Brooksem navržené vrstvené architektuře [Brooks, 1991] se aktuálně přípustná chování mezi sebou vzájemně převrstvují. Vrstva je proto v pojetí vrstvené architektury (*Subsumption Architecture*) základní (bazální) entitou, odpovídající jen jednomu typu chování. Každá vrstva pak do celkového rozhodnutí přispívá sémantickými termy — a ve vrstvené architektuře se častěji než o bazálních (základních) akcích hovoří o signálech. Na nich se také intuitivněji vysvětlí principy skládání sémantických termů. Při sloučení signálu se používají dva funkční bloky — pro prioritní a aditivní kombinaci. Pro prioritní kombinaci zavedl Brooks inhibitory, upřednostňující signály specifitějšího chování před signály z vrstev podřízených. Pro posílení významově blízkých signálů – aditivní kombinaci používá navíc supresory, které omezí počet signálů tak, že spojují více signálů do jednoho signálu součtového. Výhoda spojení aditivního a prioritního rozhodování se projevuje právě v sémantické práci se signály. Potlačováním lze zajistit požadované chování v kritických situacích, kdy krizové řešení převrství chování, definované pro standardní situace. Naopak, pokud je možné k prioritnímu standardnímu postupu přidávat i současný test minoritních módů [Nahodil, Svatoš, 2007], lze v řídicích systémech na bázi vrstvené architektury testovat a potvrzovat alternativní hypotézy průběžně ještě během řešení problému.

Za povšimnutí též stojí i vztah mezi vrstvami a různými pohledy na entity či interakce obdobně. Vrstvená architektura prokazuje opodstatnění členění zavedeného v podkapitolách 2.1.5 a 2.1.3. V oddíle věnovaném integraci modelů 2.2.4 ukáží i možné redukce zvoleného pohledu do odpovídající vrstvy vrstvené architektury. A převrstvování více modelů chování už bude opět o operacích inhibice a suprese. I tak někteří autoři realizují explicitní interakce mezi agenty.

B. Pohybové schéma

Metodu distribuovaných pohybových schémat (*Motor Schema*) zavedl Arkin [Arkin, 1989b]. Tato metoda má gradientní základ v analogii k práci v potenciálovém poli [Khatib, 1986].

Oproti tomu, jak Khatib pracuje s mapou, pohybové schéma v Arkinově reaktivním modelu už striktně vychází z dostupných Vjemů (tak jak jsem je naznačil v implikaci (2.1). Zde je každý senzor zdrojem potenciálového pole. Pro tento senzor se v pohybovém schématu implementuje jeho vlastní perceptuální schéma *SENSSCH*. Obdobně Akci na druhé straně vztahu (2.1) odpovídá její vlastní akční schéma *ACTSCH*.

Při řízení pohybovým schématem nejsou pravidla typu (2.1) vyhodnocována samostatně, ale s respektováním dalších aktivních pravidel obdobného významu. Užití pohybového schématu (Obrázek 3.5. V podkapitole 3.4.6) vysvětluje podstatu mapování blízkých významů — čím jsou si podněty významově bližší, tím menší je vzdálenost jejich obrazů v pohybovém schématu. Výsledek superpozice významově blízkých závěrů (lokálních maxim) pak lze opět interpretovat jako váhu „nejžádanějšího“ pohybu v aktuálním místě, se směrnici $[x_{opt}, y_{opt}]$.

Ohodnotím-li informaci od senzorů, která hovoří o překážce, jako zápornou, a naopak při vhodných cílech senzorickou informaci označím jako kladnou, pak v bodě $[x_{opt}, y_{opt}]$, z něhož je dovozena směrnice pohybu z aktuální pozice, nabývá funkce pohybového schématu svého lokálního maxima.

Pohybové schéma organizuje znalost „vertikálně“. To znamená, že předpoklady (Vjemy) jednoznačně implikují důsledky (Akce) bez vazby na jiné předpoklady. Výhoda aplikované vertikální online organizace znalostí je inspirativní i pro tuto práci. Mechanismus pohybových schémat mi umožní postupně rozšiřovat škálu rozhodnutí. Dodatečným zdrojem potenciálového pole se například zesiluje vliv klíčových vjemů pro koordinaci (více o integraci externích podnětů uvádím v oddíle 2.2.4). Podnětem výsledného rozhodnutí navíc nemusí být striktně jen senzorický Vjem. V kapitole 3 budu například uvažovat, jak se v jednotném prostoru odrazí obsah komunikovaných sdělení. I v přijatých sděleních lze najít informace významově podobné vlastnímu reaktivnímu modelu, stejně tak jako od informace vlastního modelu zcela odlišné.

2.2.2 Intencionální model chování

Zatímco použití reaktivních modelů je primárně zaměřené na robotiku, intencionální modely už mají uplatnění mnohem širší. V předkládané práci se nicméně soustředí pouze na intencionální modely v robotice. Zde pod pojmem intencionální rozumím vlastnosti, kterými agent rozhoduje více proaktivně — zohledňuje i své vnitřní stavy, záměry, chtění.

Skládání záměrů – intencí lze realizovat:

- A. Neomodulárně. Intencionální model přímo odráží závěry reaktivního modelu. Podle zvoleného chování z nich intencionální model filtruje přípustné reakce na aktuální situaci.
- B. Autostimulací. K podmínce hraje velkou roli váha akce. Před spuštěním dané akce se posupně akumuluje v závislosti na bezprostředním významu akce. Autoři hovoří o aktivační energii
- C. Uvážením paralelního vývoje problému. Záměry lze odhadovat i z vyhodnocení situace na vyšší úrovni detailního pohledu (fokusace).

Mechanismus ASM v intencionálním modelu zohledňuje i vnitřní stavy robota. V podstatě tyto stavy odpovídají buď záměrům (intentions) robota k dosažení cíle, či pouhé deklaraci

potřeb (needs) dosáhnout jistého cíle. Záměr je vyjádřen plánem, jak se k takovému cíli sekvenčně přiblížit. Reakce agenta na podnět z okolí neurčuje tedy pouze momentální stav, ale jeho chování je řízené cíli v tom smyslu, že agent může sám převzít iniciativu při řešení.

Robotici, řešící problémy pod zorným úhlem Umělého života („ALifers“), navrhuji i nyní intencionální modely nejčastěji jako intencionální (záměrovou) nadstavbu reaktivního modelu chování robota.

A. Neomodulární skládání záměrů

Ve světovém výzkumu behaviorální robotiky je dlouhodobě rozšiřována architektura ve směru bottom-up. Model chování se v architektuře bottom-up proaktivně zlepšuje – k účelné sekvenci závěrů reakčního modelu se zpětnou vazbou doplňují ještě další, perspektivní variace na stejný problém na nejnižší úrovni. V konkrétních podmínkách se uplatní zřejmě výhody vrstvené architektury,

- rychlého prvotního zacílení,
- a za daných podmínek přijatelné směřování průběžných kroků.

Dodatečný faktor pro vytyčení směru poskytuje intencionální model. V něm se aktuální potřeby a záměry udržují. Intencionální model reprezentuje souslednost reakcí vedoucích k cíli, a to v explicitních interakcích.

Intencionální model již používá rozšířenější reprezentaci řešeného problému než reaktivní model. V robotice se do vnitřní reprezentace navíc doplňuje kontext, který není ve vjemu obsažen. Příklad, jak se dodatečný kontext interpretuje při iterativním hledání postupu k dosažení cíle, předložila [Matarić, 1994]. Hlavní atraktivní trend, jak splnit potřeby (*needs*), rozkládá Matarić ve smyslu neomodularity (připomenuté už v kapitole 2.1.5) do dílčích podnětů. Jimi průběžně (vesměs iterativně) splňuje předpoklady pro realizaci dílčích akcí. Jí navržený mechanismus pak v celkovém chování robota vykazuje logickou sekvenci použitých základních (bazálních) akcí.

B. Autostimulace

Další z možností vnitřní reprezentace *atraktivních trendů* je aktivační síť [Maes, 1991]. V ní Patty Maes využila bottom-up členění problému na dílčí podcíle. Jí navržená síť v podstatě vychází ze sítě neuronové, přičemž základní entitou sítě je bazální akce. Z pohledu interakcí (podkapitola 2.1.3) si entity bazálních akcí v síti, jak ji navrhla Prof. Maes, udržují časovou návaznost mezi jednotlivými aktivacemi. Přímosem jejího návrhu je možnost vážit příčinu, následek a potenciální riziko konfliktu. Každou takovou znalost o příčinách, následcích a možných konfliktech reprezentuje Maes v aktivační síti [Maes, 1991] odpovídající vazbou ve smyslu kapitoly 2.1.3.

C. Interpretace paralelních podnětů na vyšší úrovni

V obou uvedených výzkumech se v ASM zohledňují pouze podněty, které jsou pro splnění zadaného cíle důležité. v podstatě se jedná o interpretaci z etologie známého principu *syntézy a akumulace redundantních podnětů* [Lorenz, 1993]. Důležité podněty jsou zde získávány právě analogií takové syntézy. A stejně tak intencionální modely, založené na technických ekvivalentech poznatků etologie mají svoji vnitřní periodu. Ta se projeví v reakci

systému [Maes, 1990], kdy se na výstup ASM dostanou i minoritní módy chování. Toto je důsledek takzvané autostimulace.

Autostimulace je sebeovlivňování organismu žádoucími podněty. Co by sebereflexe vlastních, význam vyplňujících aktů a jejich řízení je autostimulace vlastní každému biosystému. Navází na Prof. Maes a použiji ji i já ve své práci.

Během autostimulace se podněty pro spuštění minoritního módu akumulují. Reakce pak někdy až překvapivě neodpovídá aktuálnímu stavu. Spouštěčem je (stejně jako v přírodě) až poslední mnohdy i nepatrný podnět, s jehož přispěním je překročena prahová hodnota.

Podobný efekt vykazují i *vnitřní spouštěče* v situacích, kdy robot není buzen žádnou senzorickou informací. Díky takovým spouštěčům ASM vybírá akce, kterými systém proaktivně vytváří podmínky pro následné přiblížení k cílovému stavu. Proaktivní podstata intencionálního modelu je často realizována právě takovým vnitřním spouštěčem. I tyto principy jsem ve své práci využil v mém původním návrhu hybridního agenta (zde v kapitole 3) a s úspěchem demonstroval jeho využitelnost v experimentech, dále popsaných v kapitole 4. A na již známých řešeních mohu doložit [Svatoš, Nahodil, 2008a], že jednodušší intencionální část je základem hybridního řešení ASM již relativně dlouhou dobu — připomínám návrh autorů [Bensaid, Mathieu, 1997]. Toto bude již tématem podkapitoly 2.2.4.

2.2.3 Sociální model chování

Sociální model se v přístupu Umělého života chápe jako nadstavba nad reaktivní či intencionální částí. Sociální model je nezbytným předpokladem pro analýzu chování kolektivu a následného vyvození akce jednotlivce, která je přínosem pro kolektiv. To může souviset s Cílem č.3 mé práce.

Pro predikci chování spolupracujících robotů si jednotlivec do svého sociálního modelu chování zanáší modely chování ostatních. Sociální model chování v Umělém životě reprezentuje chování ostatních robotů opět co nejjednodušeji, až do rozsahu minimální reprezentace, tak jak byla použita v intencionálním modelu v předchozí podkapitole 2.2.2. Z rozboru současného stavu jsem dospěl k názoru, že Sociální model chování nejčastěji dnes zahrnuje jak tzv. informace publikované (explicitní), tak i pozorované (implicitní).

- Explicitní interakcí lze získat přímo model chování druhého robota, tedy jeho reaktivní relace a intencionální minimální reprezentaci. Získaný model se následně spojí s robotovou vlastní reprezentací problému. Způsob zveřejnění jeho vlastního modelu chování v posledních letech přechází z „čistého“ ALife — bottom-up — k hybridnímu přístupu, kdy komunikace probíhá například v jisté hierarchii.
- Implicitní interakce vnímáme senzorycky. Zde Balch [Balch, Hybinette, 2000] ukazuje, že sociální model může být postaven i na lokálních znalostech. Podnětem k takovému řešení je chování pozorované v přírodním společenství – na příkladu hejna či mravenčí kolonie.

Implicitní interakce používám v experimentech *Sledování pohyblivého předmětu* (v kapitole 4.4) a *Koordinace světlem*. Robot tak vnímá chování druhého robota, identifikovatelného silným zdrojem světla (více v kapitole 4.5).

Kromě reprezentace aktuálního stavu ovšem sociální model navíc obsahuje odhad následných kroků. Bude se zde například modelovat, za jakých podmínek může využít výsledky akce i někdo další ze skupiny. Toto беру pouze v úvahu, v předložené práci těchto možností nevyužívám.

A. Explicitní interakce

Při použití explicitních interakcí získává agent informace o prostředí od druhého agenta už předem předzpracované (senzoricky či jinak). Každý takový okruh informací někteří autoři zahrnují do vnitřního sociálního modelu chování agenta s odpovídajícími vahami. Například v již klasickém modelu BDI modelu [Wooldridge, Jennings, 1995] lze předpoklad druhého agenta chápat jako fakt, který on považuje za platný (tj. belief). Druhý agent může navíc informovat příjemce informace (prvého agenta) jak o svém aktuálním cíli (tj. desire) včetně aktuálně vybrané dílčí akce. Od ní se očekává, že povede k dosažení zmíněného aktuálního cíle (tj. intention). V mé práci tytéž důsledky budou platit pro podmnožinu agentů — roboty.

Z pohledu Umělého života vidím analogii mezi lokální percepcí a intencionální složkou BDI modelu. „Belief“ složku je v případě takové analogie třeba brát s nižší vahou, protože informace, kterou agent získá percepcí, se nemůže snadno převést na jiného agenta. Zvláště pak ne v minimální reprezentaci. Tam ani není účelné hledat přímé zobrazení informace, kterou agent získává percepcí, na jiného agenta. Při minimální reprezentaci je už pro rozhodování agenta v kolektivu použitelnější desire složka aktuálních cílů (pro synchronizaci s ohledem na předpokládaný průběh chování druhého agenta). Nejvýznamnější je v procesu získávání znalostí od druhého agenta intencionální složka. V rozšíření minimální reprezentace může agent stejně jako s percepcí (Vjemy) nakládat s informací o intencích (záměrech). Když cizí agent poskytuje informaci o svých vnitřních stavech, významově to odpovídá klasifikaci podle Jenningse: záměrům (intentions) robota k dosažení cíle, či pouhé deklaraci potřeb (needs) dosáhnout určitý cíl.

B. Implicitní interakce

Implicitní interakce představují specifický způsob sociálního chování, vypozerovaný z přírody. Z etologie lze jako klasický uvést *sociální model chování hejna*. Implicitními interakcemi je v něm vnímání sady vzorů, jimiž hejno na zainteresované působí [Reynolds, 1994]. Lze je vnímat jak globálně, jako vzor chování celého hejna, tak lokálně jako návod – vzor – pro vyhodnocování vzdálenosti *jednoho jedince hejna od druhého*. Pro účely mé práce uvádím (a dále v kapitolách 3 a 4), že implicitní interakce probíhají komunikacemi následujících dvou typů:

1. **Koherence** – zde se jedná o uplatnění stejných pravidel u všech agentů. Pokud je takovým pravidlem například averze, pak v případě mravenců bylo vysledováno [Pasteels a kol., 1987], že takto jednoduchým pravidlem se v kolonii mravenců dosáhne jejich vyrojení po celé obranné ploše mraveniště.
2. **Sledování** – zde se jedná o sdílení aktuálního řešení problému podle důvěryhodného vzoru. V přírodě je příkladem značení cest za potravou [Holldobler a kol., 1974]. Cesta k cíli se mezi jedinci často udržuje řetězcem vlastních těl. Takový řetězec vzniká velice jednoduchou komunikací – taktilně či vizuálně.

Tento způsob implicitní komunikace je pro mne v této silně inspirativní. Sledování lze úspěšně využít už ve velmi malém kolektivu robotů, já jsem experimentoval se čtyřmi (posléze i reálnými) roboty. Aby se projevila koherence, je třeba mít v kolektivu již výrazně více robotů. Proto koherenci nebudu dále ve své práci demonstrovat, ani se jí jinak zabývat.

Přenos implicitních interakcí je zřetelně asynchronní. Alternativní biologické procesy jsou ovšem charakterizované vlastní periodou klidového buzení [Tinbergen, 1953]. Obdobné módy se objevují i v Umělém životě. Vlastní *periodu sociálního chování* vykazuje i kolektiv, sestavený z intencionálních agentů s významnou integrální složkou v rozhodování.

2.2.4 Integrace modelů chování – hybridní přístup

Při integraci sociálního modelu s reaktivním a intencionálním modelem je pro kolektivní úlohy ve smyslu Umělého života primárním problémem převedení znalostí ze sociálního modelu do dílčích aktivit robota. Způsob, jakým sociální model chování obohatí aktuální reaktivní či intencionální model chování robota, jsem naznačoval dosud jen rámcově v kapitolách 2.1.5 a 2.1.3. V tomto oddíle akcentuji integraci sociálního modelu v souvislosti se známými řešenými kolektivními úlohami. Začnu od reaktivního modelu chování, kde sociální model bude pouze reflektovat aktuální situaci z pohledu dalších členů kolektivu. Na použití intencionálního modelu chování je postaven můj nový přístup — návrh hybridního agenta, rozváděný podrobně v kapitole 3.

A. Integrace modelů chování převrstvováním

Vrstvení sociálních, intencionálních, či reaktivních hypotéz či signálů probíhá obecně stejně jako převrstvování bazálních aktivit jediného modelu. Neuvažují se předpoklady jejich vybuzení. Podstatné je určit, zda bazální akce jsou vzájemně komplementární. Pokud akce nelze spustit současně, spouští se z nich pouze majoritní akce. To bylo již podrobněji popsáno v podkapitole 2.2.1.

Hodnocení paralelního souběhu se získává těžko za běhu, snáze při návrhu. Integrace převrstvováním se totiž dobře rozšiřuje směrem od známého k méně známému. Tedy od funkcí, které zajistí základní chod robota. Syntéza chování, tak jak ji navrhnul Barnes [Barnes, 1996] se proto v kolektivu srozumitelněji odvozuje ve směru od vlastních poznatků k obecnější charakteristice problému.

Syntéza chování a její čtyři úrovně náhledu na problém (Barnes):

1. Úroveň aktéra obsluhuje vnitřní stavy. Chováním na této úrovni se robot udržuje v provozuschopném stavu. Fyzický robot si musí dobíjet baterie, doplňování energie je potřeba i nejnižší úrovně virtuálních agentů. Tato úroveň se nedotýká možného využití při řešení stanovených Problémů a dále ji proto nerozvádím.
2. Úroveň prostředí zahrnuje ta chování, která reagují ještě i na původní, nepředzpracovaná data ze senzorů. Například už na základě lokální detekce překážek je robot schopen bezkolizní navigace. Na této úrovni se už dotýkám Problému č.1 z úvodu této práce — orientace robota v herním poli — i když velice hrubě.
3. Úroveň kooperace vymezuje chování vázaná na konkrétní objekty. Na této úrovni jsou objekty popsány atributy lokální povahy. To znamená, že každý z robotů může tentýž objekt detekovat podle jiných hodnot těchto atributů. Proto do úrovně kooperace (jak

ji nazval Barnes) spadá i řešení Problému č. 2 — výběr akcí. Nad klasifikovanými objekty prostředí lze maximalizovat užitek dle některého z dříve uváděných kritérií. Stejně tak do atributů objektů patří i model, popisující chování robota na předešlé 2. úrovni — prostředí.

4. Úroveň úkolu se týká chování, odvozených od atributů globální podstaty. Chování této úrovně se synchronizuje na jednotlivých bodech kolektivního plánu, tedy jednoznačně identifikovatelných mezikroků. V kolektivním plánu při navigaci jde například o absolutní vymezení cílových oblastí, a jejich transformaci do relativní souřadné soustavy robota. Toto se nabízí jako jedno z možných řešení mého Problému č. 3 — koordinace akcí v týmu.

Tímto jsem zde shrnul poznatky, důležité pro splnění Cíle práce č. 1. a č. 2. Konkrétní řešení bude podrobněji rozebráno v kapitole 3 a 5.

Jak z mnou prostudované literatury, tak z vlastního výzkumu vyplynulo, interakce mezi objekty se projevují zejména ve 3. Úrovní kooperace a ve 4. Úrovní úkolu. V obou případech lze použít pouze implicitní interakce. Pro Úroveň kooperace [Goldberg, Matarić, 1997] to například znamená maximálně převést klasifikaci objektů na úroveň senzorů, a odpovídající sémantiku následně využít při určení majoritních a komplementárních signálů. Zde se vyplatí vyhodnocovat signály ekvivalentní aktuálnímu chování – pak je celá logika výběru akce pro dané chování v kooperativní úloze převoditelná na potlačování kontraproduktivních bazálních akcí. Na implicitních/explicitních interakcích coby indikátorech bazálního akcí je postavena úloha sledování [Parker, Touzet, 2000]. Parker píše, že pro řízeného robota se bazální akcí druhého robota rozumí nejen jeho oznámené aktuální chování, ale i jeho pohyby, indikované senzory (a to včetně veličin, převyšujících přednastavený práh). Všechna tato fakta mohou potlačit vykonávání kontraproduktivních akcí řízeného robota.

Senzorické detekce lze užít pro lokalizaci skupinového cíle na 4. Úrovní úkolu. Pouze je k označení cílových oblastí třeba použít veličin, které daný úkol charakterizují nezávisle na jedinci, ale v souhrnném hodnocení z pohledu kolektivu. Ve vyznačených oblastech se pak roboti svými lokálními pravidly navádí do shluků [Matarić, 1992]. Zmíněný efekt ovšem není jen záležitostí mohutného kolektivu. Převrstvováním celých bazálních chování [Parker, 1998] lze řídit i několik volně svázaných heterogenních jedinců. Výhoda je zřejmá, stejně jako při převrstvování modulů v rámci jednoho řídicího systému – při použití vrstveného modelu sociálního chování jsou roboti mezi sebou vzájemně zastupitelní. Parker použila pro řešení téže úlohy jak kolového robota, tak robota šestinohého, přičemž princip řízení byl u obou typů robotů stejný – oba tedy interpretovali tutéž 4. Úroveň úkolu. Parker také ukázala, že úloha je řešitelná i při výpadku některého ze spolupracujících robotů. V takovém případě je ale řešení logicky mnohem méně efektivní, oproti řešení s více roboty v kolektivu.

B. Integrace modelů chování superpozicí

Právě uvedená integrace modelů chování převrstvováním pracovala zejména se sémantikou bazálních akcí. Při využití superpozice lze dílčí podněty, přicházející z modelů chování ještě i vzájemně vážit a tak se i lépe v rámci řešeného problému adaptovat – přizpůsobit se jeho aktuálním podmínkám.

Pro ovážení chování robota se v literatuře dnes používá jednotná metrika. Jak spojit ohodnocení dosud nepoměřitelných akcí ukázal např. Balch [Balch, 1997] na jím zavedené

diverzitě kolektivu (odpovídá entropii).

Jinou možnost dává oceňování akce její „vzdáleností ve stavovém prostoru“ — vzdálenost se měří od bazální akce po cílový stav.

Obdobnou metriku zavádím i já [Svatoš, 2001], [Svatoš, Nahodil, 2008a] ve svém původním návrhu integrálního a diferenciálního kritéria. Jejich podrobnější popis následuje v kapitole 3.5.2.

V superpozici podnětů z jednotlivých modelů chování, obdobně jako ve fuzzy logice není zvolená akce jen důsledkem nejsilnějšího předpokladu. Akce je dále podložena (a modifikována) předpoklady podobného významu, ale nižší váhy. Takové podložení často plyne z implicitní interakce. S obdobným poměřováním chování se lze setkat již u reaktivního modelu. V pohybovém schématu (uvedeném v podkapitole 2.2.1) se analogicky skládají data pro aktuální krok z percepce do směrového vektoru. Jak ukázal Arkin [Arkin, 1992], v pohybovém schématu se mohou zvážít i implicitní interakce plynoucí ze sociálního modelu.

V kolektivních úlohách má superpozice tento přínos:

- Rozšiřuje záběr oblasti, reprezentované v perceptuálním schématu. Zatímco u detailní reprezentace prostředí se robot může spolehnout na své proximální senzory. Po konzultaci s druhými roboty (vzdálenými tak, že jsou mimo dosah jeho proximálních senzorů) si navíc může udělat i přehled o situaci v místech, které by svými senzory sám nedetekoval.
- Doplnuje dodatečné znalosti o společných cílech kolektivu. V pohybovém schématu lze využít i sémantický vztah mezi obrazem percepce a akce. Pohybové schéma tak v podstatě nabízí prostor pro integraci závěrů reaktivního řízení spolu s plánovačem na vyšší úrovni. Integraci plánovače intencionální úrovně předvedl [Gat, Dorais, 1994]. Gat a Dorais řadí reaktivní procedury na nižší úrovni tak, že v každé situaci hodnotí úspěšnost bazálních akcí pro dosažení dílčího cíle. Tedy zda daná bazální akce byla úspěšná či vedla k chybě — oddálení se od cíle. Sekvenční vrstva v nejistých situacích dále podle Gata volá nadřazenou vrstvu sofistikovaného plánovače, a jako radu dostane (za dané situace platný) globální cíl. Rada od plánovače není pro sekvenční odvození závazná. V případě rozkolu dává Gat možnost aktuální činnost přerušit.

Samostatným předpokladem po dílčí odvození akce může být také znalost vycházející ze sociálního modelu [Balch, Arkin, 1995b]. Balchova práce je pro mne velmi podnětná, protože on, stejně jako já pracoval se skupinou čtyř pohybujících se robotů. Aby jednotlivci fixoval svoji polohu vůči ostatním, zavedl pravidla formace. Pokud robot k pravidlům formace používá i situační pravidla (např. vyhýbání se překážkám), kolektiv se ve výsledku pohybuje po očekávané trajektorii. Jejich pohyb se podobá přesunu pravidelného obrazce (obdobně jako u hejna tažných ptáků). O žádném obrazci ovšem individuální pravidla nehovoří. Obrazec vzniká až vhodně provázanou superpozicí ve vhodné metrice. Evidentně platí, že čím lépe jsou bazální chování či dodatečné znalosti sémanticky přiřazeny do pohybového schématu, tím je i jejich kombinace citlivější oproti integraci převrstvování.

Z analýzy dostupné literatury v otázce použitelnosti převrstvování, respektive superpozice, jsem učinil obecně platný závěr: Pokud ze dvou kvalitních kandidátů na řešení (dvou bazálních akcí) vzejde kompromis, který je ještě kvalitnější než původní kandidáti, pak je superpozice přínosem. Pokud ne, pak je lepší pro sloučení reaktivního, intencionálního a

sociálního modelu chování použít převrstvování. Převrstvování mezi jednotlivými modely za běhu jsem však v práci dále nerealizoval, zaměřil jsem se na využití superpozice.

2.2.5 Důležitá fakta a závěry paralelního řešení Typových Úloh

V této podkapitole 2.2.5 jsem rozebíral modely chování multiagentního systému, které se dotýkají dříve zavedených Typových Úloh. Na hlavních zástupcích odpovídajících směrů behaviorální robotiky jsem na podkladě prostudované literatury doložil účel a vlastnosti reaktivního, intencionálního a sociálního modelu. Na stejném místě jsem připomněl i složitost vnitřní reprezentace problému v takovém modelu. Nakonec jsem připojil ukázkou projektů zaměřených na integraci uváděných modelů do jednotného mechanismu výběru akce.

V prvním oddíle u reaktivního chování zdůrazňuji metody pro vyhodnocení aktuálně přípustných reakcí. Na nejnižší úrovni detailu postačí podněty pro bazální akce vyhodnocovat superpozicí. Pokud je ovšem povaha lépe popsána na různých i vyšších úrovních detailu, je výhodnější prvotní záměry modifikovat převrstvováním.

V druhém oddíle jsem nastínil další současné možnosti korekce v intencionálním modelu chování. Jedná se zejména o dodatečné prosazování dlouhodobých záměrů vnitřními podněty a jejich vzájemnými cílem odůvodněnými autostimulacemi.

Ve třetím oddíle jsem připomněl implicitní interakce mezi roboty v kolektivu. Jedná se o vjemy, z nichž lze rekonstruovat cílově orientované jednání ostatních robotů v kolektivu. Stojí tak na pozadí jednoduchého sociálního modelu.

Ve čtvrté části podkapitoly 2.2.5, nazvané Integrace modelů chování, jsem pak ukázal možnosti použití základních principů skládání reaktivních podnětů i v případě zmiňovaných modelů chování do jednotného projevu autonomního robota. Převrstvování jsem doložil na příkladu skládání funkcí z různých problémově orientovaných úrovní. Příkladem superpozice bylo použití pohybových schémat při řízení robota v kolektivu.

2.3 Popis paralelních dějů gramatikou

Svá omezení má i Umělý život, který by se záměrně opíral jen o reaktivní model chování. Když už je totiž návrh postaven na minimální reprezentaci, pak by se tato reprezentace neměla zjednodušovat na úkor složitého výběrového mechanismu [Nahodil, Svatoš, 2002]. V práci se snažím o to, aby výkonný agent neměl složitější realizaci než stavový automat, resp. jeho rozhodování by nemělo být složitější než nedeterministická varianta stavového automatu. (Kromě informace o aktuálním stavu stavový automat nemá žádnou další paměť — přechod mezi jeho stavy je pevně určen — determinován symbolem čteným ze vstupu.)

Zde je na místě připomenout regulární gramatiku, coby v praxi nejčastější popis takového automatu. Existuje totiž celá metodika převodu reaktivního chování multiagentního systému na gramatický systém [Kelemen, Kelemenová, 1992]. Gramatikami se budu zabývat protože:

- při Distribuci zájmových oblastí budou základem gramatik agenta pravidla, (co odvodit z podnětů),
- pro Oportunistické využití prostředí se z takových pravidel vybere nejlépe hodnocené,
- a v úlohách Koordinovaného pohybu se využije přirozeného paralelismu gramatického odvození [Kelemen, Kelemenová, 1992].

U problému popsaného gramatikou lze nahlédnout i na další jeho vývoj. V první řadě je z užité gramatiky (zpětně podle sekvence pravidel) odvoditelné, nakolik je řešení problému v možnostech Kolektivu. U některých tříd gramatik lze také posoudit trendy popsaného problému. Možným závěrem vývoje na základě jednoduchých odvozovacích pravidel jsou stabilní struktury. Hovoříme pak o tom, že výsledek odvození je konvergentní. Příkladem tohoto vývoje může být výrazně se neměnicí struktura fraktálů a to ani po několika iteracích [Lindenmayer, 1968a, Mařík a kol., 2003, Mařík a kol., 2007].

2.3.1 Paralelní aktivace akcí a gramatika

Při rozboru stavu současných možností se zaměřuji i na gramatiku z pohledu návrháře reaktivního řídicího systému. Návrh i formalismus se zaměřuje na následující okruhy:

A. Podněty v gramatice a abeceda agenta

Referenční termíny v gramatice vymezují abecedu agenta. Abecedou se bude popisovat každý stavový přechod vyhodnocený v ASM. V následujícím oddíle, věnovaném vztahu komponent a prostředí, uvádím souvislost abecedy s entitami Umělého života. V kontextu prostředí je před ASM předkládán vždy nový dílčí problém (dekompozice problému). Do jeho gramatického odvození se už událost dostává pod odpovídajícím referenčním termínem. Komponenta heterogenního agenta, jenž ve své abecedě takový odpovídající termín obsahuje, tak detekovanou událost akceptuje za svůj dílčí cíl. Komponenta může poskytovat (sdílet) svoji funkci ostatním částem informačního systému a to prostřednictvím pevně vymezeného rozhraní. Obecně se jedná o ucelenou programovou část, zapojenou do celkové architektury informačního resp. řídicího systému použitého k řešení problému.

Obdobně lze z abecedy odvodit i možnost koordinace při řešení. Pokud se referenční termín vyskytne v abecedách více agentů zapojených do řešení, může dojít k paralelnímu

splnění dílčích cílů. U heterogenních agentů je častá také sekvenční varianta, kdy jeden agent produkuje termíny, které je schopen dále rozvíjet pouze jiný agent.

B. Přirozený paralelismus gramatických systémů

Jak ukázal prof. Brooks [Brooks, 1999], gramatika je také jeden z pohledů na paralelní aktivaci procesů. Ronny Brooks například celý kolektiv nastiňuje jako kolonii v gramatickém smyslu, a to i s jakýmsi jejím specifickým životem. Pojem Kolonie používá totiž pro gramatický popis nepřímé interakce stavových automatů. Stavový automat svojí aktivitou zapisuje do popisu prostředí nový stav tohoto společného prostředí Kolonie. Protože nový stav přečte jiný automat, dochází k oné nepřímé interakci. Život — umělý život dokládá Brooks větší generativní silou Kolonie v porovnání s původní generativní silou dílčích komponent. Použijeme-li pro Kolonii komponenty s regulární gramatikou, Kolonie generuje řešení, ke kterému by byla potřeba komponenta s ASM, založeném na kontextovém odvozování. Paralelním pokrytí problému tak např. Takadama položil základ i pro adaptaci čistě reaktivního systému [Takadama a kol., 1999] na řešení NP-těžkých problémů.

2.3.2 Paralela mezi gramatickými termíny a entitami v Kolektivu

V minimální reprezentaci jsou všechny možné stavy řešení problému (tedy i dílčí cíle) zakódovány ve vztazích mezi agentem a prostředím. Gramatika tyto entity ohraničuje jednoznačně. Aktivním prvkem je agent. Ten vůči prostředí představuje samostatnou komponentu. Na vlastní projev v prostředí má komponenta svá pravidla. Nejčastější příčinou aktivace bazální akce je existence význačných referenčních termínů v prostředí. Jiné příčiny se mohou být ve vnitřních stavových přechodech komponenty. Jevy pozorovatelné v prostředí pramení právě z těchto aktivních entit. Zavedením komponent s vhodnými pravidly máme možnost změnou funkce komponent podnítit žádoucí změnu chování Kolektivu. Na Kolektiv má vliv jen prostředí, a to je prvkem vesměs pasivním. Uvažujme jen uzavřené prostředí. Zevnitř ho mění pouze agenti svými zásahy. Korespondence – vztah mezi některými termíny – (entitami) Umělého života a gramatickými termíny je na další straně v tabulce 2.1.

V tabulce uvedené korespondence jsou svými závěry – paralelami gramatických odvození podnětné právě pro robotiku [Ferber, 1996]. Gramatika pokrývá celý proces postupné syntézy cílového stavu z mezivýsledků. U chování, sestaveného z bazálních akcí lze takto analyzovat mezivýsledek její a jak tento bude vnímán v následujícím kroku. Závěrem analýzy je mezivýsledek po n krocích.

Použitá pravidla vymezují tzv. tranzitivní uzávěr. Tranzitivním uzávěrem (tj. výčtem všech stavů, které lze postupně odvodit z vjemů a přípustných bazálních akcí) zúžíme v návrhu možný okruh výsledků, kterých robot může při použití uvažované gramatiky dosáhnout. Průběžnou analýzou vjemů a bazálních akcí zjišťujeme, nakolik je žádaný výsledek robotem reálně dosažitelný.

V naznačené korespondenci (paralele) se striktně dodržuje minimální reprezentace. Proto se v gramatice ani nerozlišuje mezi okamžitým stavem prostředí a vnitřním stavem komponenty. Také ASM ovlivňuje jak stav Agentů na nízké úrovni, tak i stav řešení celého problému na úrovni Kolektivu (kapitola 2.1.3). Aspekty integrace komponent na úrovni Kolektivu v mém návrhu v následující 4. kapitole. Oba dva pohledy v této práci využívám.

Značení	Gramatický term	Entita v kolektivu robotů (pro Umělý život)
N	množina neterminálů komponenty	seznam aktivních detektorů, vnitřních stavů
T	množina terminálů komponenty	seznam bazálních akcí s projevem navenek, výsledky aktivity
V	abeceda komponenty, $V = N \cup T$	senzorické vstupy, vnitřní stavy, výsledek aktivity
V^*	množina všech řetězců nad V – monoid generovaný abecedou	triviální řešení není záměrem – téměř výhradně se užívá V^+
V^+	řetězce neprázdných slov nad abecedou V	stavový prostor řešeného problému
w_E	řetězec prostředí	mapa/reprezentace okolního světa
S	axiom; startovní symbol nebo řetězec	iniciální podmínky řešení problému
G	gramatika	chování robota
$L(G)$	jazyk gramatiky	možnosti robota – výčet stavů prostředí, které může ovlivnit

Tabulka 2.1: Vztah mezi gramatickými termy a entitami Umělého života. Naznačená korespondence přibližuje závěry gramatických odvození, podnětné též pro robotiku. Gramatika pokrývá proces syntézy cílového stavu z mezivýsledků. Gramatikou lze analyzovat syntézu entit, které jsou nakonec částí výsledku. U chování z bazálních akcí lze takto analyzovat její výsledek, a jak je takový výsledek vnímán v následujícím kroku. Závěrem analýzy je mezivýsledek po n krocích. V tranzitivním uzávěru použitých pravidel, vjemů a bazálních akcí modelujeme, nakolik je žádaný výsledek robotem dosažitelný.

z pohledu nízké úrovně základ takové jednoduché komponenty reaktivního agenta zůstane na regulární gramatice, případně na Lindenmayerově systému [Lindenmayer, 1968a]:

A. Regulární gramatika

Ferber ukázal [Ferber, 1996], že shrnutím všech potřebných pravidel do regulární gramatiky se realizuje jednoduchý reaktivní systém dle přiřazení (2.1). Emergentní vlastnosti paralelně pracujících systémů se tak odvíjí od vcelku jednoduchých předpokladů.

Definice 2.3.1 Regulární gramatika

Chování reaktivní komponenty (agenty) je popsáno regulární gramatikou, tedy běžnou strukturou

$$G = (N, T, R, S),$$

se standardním významem N, T, S, V a prepisovacími pravidly R typu

$$A \rightarrow aB, A \rightarrow a, A \rightarrow \epsilon \quad \text{pro } A, B \in N, \quad a \in V,$$

kde ϵ značí prázdný (mazací) symbol.

□

Regulární gramatika není ovšem jediným stavebním prvkem, ze kterého lze postavit systém Umělého života [Sebestyén, Sosík, 2004].

B. Lindenmayerovy systémy

Jiný princip je u struktur, které narůstají ze základních elementů. Tam ztrácí význam rozdělení na neterminály a terminály. Rozvíjena tak může být většina závěrů. Záleží jen, zda v gramatice existuje pravidlo pro rozvoj získaného závěru. Takové jsou i Lindemayerovy systémy [Lindenmayer, 1968a]. Vzájemné přeměny modulů mají svá odvozovací pravidla (productions). Mezi každými dvěma okamžiky pozorování se podle nich mění celá struktura — všechny její moduly nahradí moduly předepsané odvozovacími pravidly.

Bez omezení na obecnosti se já v této práci soustředím na lineární strukturu výsledku, čemuž odpovídají bezkontextové Lindenmayerovy systémy:

Definice 2.3.2 *Bezkontextové L-systémy (0L-systémy)*

Řetězcový 0L-systém je uspořádaná trojice $G = (V, P, \omega)$, kde k abecedě V nově používám

- ω – axiom – neprázdné slovo $\omega \in V^+$
- P – množina odvozovacích pravidel $a \rightarrow \chi$
- a – přepisovaný symbol
- χ – levá strana pravidla $\chi \in V^*$.

□

Hovoříme o nich také jako o L-systémech bez interakce (interactionless) či o 0L-systémech [Lindenmayer, 1968a]. Zajímavější, a trochu komplikovanější, jsou pak kontextové Lindenmayerovy systémy. Budou uváděné v příští kapitole.

2.3.3 Paralelní odvození

Jednotlivé stavební prvky – entity v Kolektivu robotů (viz Tabulka 2.1) povětšinou realizují jednoduché gramatiky. Ostatně:

- na regulárních gramatikách je postavena Kolonie,
- a na Lindenmayerových systémech se staví složitější eko-gramatické systémy.

S využitím definice [Kelemen, Kelemenová, 1992] zde upřesním pojem Kolonie:

Definice 2.3.3 *Kolonie*

Kolonie je uspořádaná trojice

$$Col = (G, V, T),$$

kde je

- G - konečná množina regulárních gramatik popisujících chování jednotlivých komponent kolonie (2.3.1) $R = \{R_1, R_2, \dots, R_n\}$
- V - abeceda kolonie
- T - terminální symboly kolonie.

□

Jednotlivé gramatiky Kolonie popisují chování odpovídajících agentů v multiagentním systému. Posuzujeme v nich stavové změny – modifikaci řetězce prostředí akcemi v jednotlivých odvozovacích krocích. Jak dokázal Kelemen, čistě jen na základě znalosti o množině

komponent S nelze rozhodnout, zda daná Kolonie je konvergentní, divergentní, nebo stabilizovaná [Kelemen a kol., 2001].

Já jsem se zaměřil na jednu z vlastností Kolonie, kterou uvádí Definice 2.3.3, totiž, že Kolonie operuje v konečném počtu vnitřních stavů. I já budu vycházet z gramatického systému, kde bude chování Kolonie postihnuto (popsáno) pomocí konečné abecedy a pravidel na ni aplikovaných [Svatoš, Nahodil, 2008a]. Oproti gramatickému popisu bude má charakteristika Kolonie rozšířena o nové vazby. Prof. Kelemen ve svém popisu Kolonie na úrovni společenství sjednocuje jednotlivé regulární popisy komponent [Kelemen, Kelemenová, 1992] dle operací nad společným prostředím. Popis chování celku analogicky i u mne bude vycházet z přijaté a vzájemně sdílené abecedy, pravidla jsou vyjádřena množinou gramatik [Kelemen, 1993].

2.3.4 Důležitá fakta a závěry Popisu paralelních dějů gramatikou

Na abstrakci objektů za neterminály jsem ukázal gramatiku pro popis generativního zařízení senzorio-akčních rozhodnutí. Ekvivalentní symboly jsem přiblížil i v následném formalismu paralelních gramatických systémů. v uvedeném průřezu gramatických systémů jsem se zaměřil na symboliku popisu průběhu chování agentů. Paralelní zpracování v ní přešlo na několik generátorů událostí v prostředí. Regulární gramatiky, Kolonie a Lindemayerovy systémy připomenuté v tomto oddíle mi poslouží k úvahám nad vlastnostmi eko-gramatického systému. v rámci této 2. kapitoly „Současný stav a rozbor možných řešení“ jsem se zabýval také možnostmi implementace senzorických vjemů a vnitřních stavů robota do výsledného postupu řešení. Problematiku takovéto potřebné syntézy základních – bazálních akcí jsem se v této kapitole zabýval ze dvou obecných úhlů pohledu:

- Na modelu chování robota jako agenta ve velmi jednoduchém multiagentním systému.
- Na gramatice, která popisuje logiku odvození Akce ze Vjemu (senzorického nebo vnitřního podnětu).

Oba tyto procedurální přístupy k syntéze bazálních akcí ve svém návrhu využiji pro vyhodnocení kvality mechanismu výběru akce ASM. Oba pohledy jsou procedurální. Ukazují postup skládání bazálních entit chování před vstupem ASM. Na modelu chování jsem ukázal možnosti syntézy vhodných akcí metodou superpozice. Druhou inspirativní metodou je převrstvování celých akčních bloků čistě na základě aktuálních podnětů. Obě metody umožňují postupně řešit komplexnější problém bez nutnosti jeho významnější reprezentace.

2.4 Shrnutí závěrů a důležitých faktů Druhé kapitoly

Celá druhá kapitola se věnuje analýze současného stavu a rozboru možných řešení mnou hned v úvodu nastíněných Problémů. Ukazují v ní užitečné alternativy, jak Problémy a Cíle práce lze řešit jednodušeji a operativněji přístupy Umělého života. Ty jsou založeny na samoorganizujícím principu behaviorálního řízení „zdola nahoru“ (bottom-up), emergentních jevech a evolučních mechanismech. Již základní princip odvozování „zdola-nahoru“ považuji za účinný prostředek k dosažení Cílů, které jsem si v práci stanovil.

Uvedený přehled současných možností mi slouží jako odrazový můstek mého vlastního návrhu hybridního řídicího systému mobilního robota. Okruh analyzovaných problémů jsem si zúžil pouze na několik typických úloh. Omezují se přitom pouze na skutečně nezbytné interakce mezi základními akčními entitami.

Shrnu-li:

V první části jsem si ke stanoveným Cílům zvolil tři Typové Úlohy. Na zvolené „neomodulární formování výsledku“ zavádím pohled ve třech úrovních detailu („směrem od mikro do makrosvěta“). Na Úlohách jsem pak ve druhé podkapitole ukázal možnosti řešení na nejvyšší 3. úrovni — na úrovni Kolektivu. V následujících kapitolách práce pak budu své konkrétní řešení rozebírat detailněji — už s vazbou na jednotlivé Cíle.

Ve druhé podkapitole ukazují princip vytváření kolektivního řešení na multiagentním systému. Problémům odpovídají tři principiální modely chování agenta. z reaktivního modelu chování se obecně získává přehled aktuálních senzorických podnětů a přímých akčních důsledků. Pomocí intencionálního modelu se do rozhodování nejčastěji zavádějí důsledky agentovy aktivity v časovém kontextu. Integraci těchto modelů do výsledného iterativního rozhodování lze provést ve stejné hloubce rozpracovanosti. Proto i vnitřní reprezentace problému může být ve všech modelech popisována jednotným způsobem. Do svého ASM integruji dílčí závěry z jednotlivých modelů převrstvováním a superpozicí.

V závěru jsem se zaměřil na exaktní popis nezbytné vnitřní reprezentace. Ukázal jsem ho v úzké paralele mezi uváděnými předpoklady k řešení Problémů a zavedeným gramatickým systémem. v literatuře jsem našel nejjednodušší analogie mezi reaktivními agenty a regulárními gramatikami. Paralelou ke Kolektivu jednoduchých reaktivních agentů je gramatická Kolonie. Obdobný závěr dokládám s použitím složitějších gramatických systémů na konci kapitoly.

Kapitola 3

Návrh vlastního řešení a jeho teoretický popis

Tato třetí kapitola analyzuje teoretické funkce vhodného řídicího systému. Posléze vybírá finální varianty jednotlivých bloků předkládaného návrhu hybridní architektury řídicího systému. Vybraná varianta umožňuje jednoduché situace řešit čistě *reaktivně*. Při řešení složitějších situací se pak využívá její *intencionální model chování*. z reaktivního modelu řízení robota odvodím nejdříve *přípustné akce*. Jádrem kapitoly je superpozice ohodnocení dílčích akcí. Bazálními entitami v navržené architektuře jsou proto *parametrizované akce*.

Architektura sleduje dílčí interakce dobře popsané eko-gramatickým systémem. V tomto formalismu jednotliví roboti vybírají své akce sekvenčně. V Kolektivu se pak akce jednotlivých robotů skládají paralelně. V rámci jednoho robota se obdobně k reaktivnímu modelu postupně připojuje intencionální část spojenou s vlastním mechanismem výběru akce. Intencionální model mi umožní vhodněji zahrnout mnohem více předpokladů spuštění přípustných akcí. Nové a původní je zde rozšíření základní idey mechanismu výběru akce. ASM zde doplňuji ještě o kritéria, umožňujícími rozpoznat kvality formovaného řešení.

3.1 Bloková schémata reaktivního a intencionálního modelu

Navrhovanou hybridní architekturu modifikuji jednoduchým, ale účinným *adaptivním mechanismem*. Užívám učení bez učitele. V průběhu řešení metodou pokusu a omylu koriguji dříve nastavené hodnocení bazálních entit. Závěrem přibližuji obecné možnosti a přednosti mnou navrženého mechanismu výběru akce v několika vybraných typových úlohách.

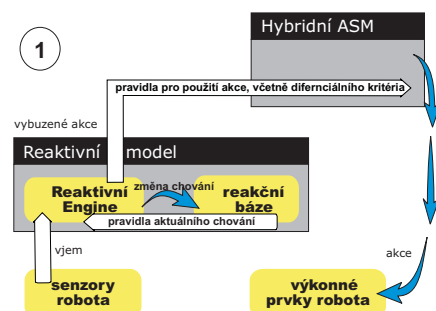
Můj způsob řešení **Problémů** z úvodu plně odráží možnosti cílové platformy mobilních robotů. K řízení čtyř mobilních robotů jsem zvolil iterativní metodiky Umělého života jako způsob, kterým lze řešit problém i na základě kusých informací. Já v duchu přístupu Umělého života hlavně *hledám dílčí kroky*, jenž se hodí *pro řešení aktuálního problému* v souladu s udanými pravidly. Takových kandidátů generují mechanismy Umělého života bezpočet. Já nad nimi rozvíjím svůj osobitý analytický přístup, jak vybrat nejlepší variantu. Tam už se nezaobírám sémantikou, pracuji pouze s kvalitou posuzovaných alternativ. Podstatu mého původního řešení vysvětlím právě v této kapitole.

V této kapitole se opírám o jisté třídy vjemů a bazálních akcí, tak, jak byly zavedeny v kapitole 2.2.1. Pomocí těchto entit minimalizuji celou reprezentaci stavu problému až na úroveň událostního řízení. Namísto snahy o precizní interpretaci kusého vjemu robot reaguje na dostupné příznaky – hrubé senzorické vjemy. Konkrétní senzory, které jsem využíval ve svých experimentech, jsou rozebrány v tabulce 4.3. Z výběru akce se snažím eliminovat závěry, které jsou postaveny na příliš problémově specifických znalostech. *Nechtěl jsem používat pro přesné závěry neodpovídajících předpokladů resp. nepřesných vjemů*. Naopak chci ukázat, jak je dokáže zastoupit „hrubá síla“ metod Umělého života. Časté použití jednoduché metody tak paradoxně vede k lepšímu dosažení cíle, než s použitím složitějšího sofistikovaného algoritmu.

Události nastavuji pro **reaktivní a intencionální model**. Za této přednost kompozice považuji modularitu. *Oba použité modely spadají do vyhodnocení jediného scénáře ASM*. Při splnění logických podmínek, se do ASM dostává hrubý výběr přípustných aktivit (bazálních akcí). Následně ASM k volbě jedné akce přípustných aktivit zohledňuje podklady z intencionálního modelu. Jimi se posouvají zejména parametry zvolené akce. Takto modifikované rozhodnutí odpovídá optimu hodnocení bazálních akcí v intencionálním modelu. Přehledově se v kapitole 3.1 zaměřuji na **tři oddělené části řídicího systému**:

(1)

V *reaktivním modelu* vystihuji sémantiku aktivit (tak jak jsou odvozovány v ASM) jejich gramatikou, která zjednodušeně ukazuje na příčiny aktivit. A to nejen v odkazu na *vjem* (dle reaktivního modelu), ale i na vnitřní stav (dle intencionálního modelu). Zde proto nejdříve ukážu korespondenci mezi ASM a gramatikou, a poté i faktické důsledky její realizaci ve skupině reálných robotů. Pro pokrytí obou dvou modelů se inspiroji třídou ekologických systémů. Oproti základnímu Lindenmayerovu systému užívám moduly v kontextu, jedno pravidlo v předpokladech kombinuje více lokálních podnětů do jednoho uceleného vzoru.



(2)

Návrh *intencionálního modelu* využívá vlastností zdůrazněných v předchozí kapitole. Jeho hlavním znakem je paměť, která umožňuje využít autostimulací. Oproti naznačenému principu dle [Maes, 1991] kombinuji autostimulaci se zapomínáním. Zapomínání je připraveno hlavně pro následující rozšíření — adaptaci intencionálního modelu chování metodou pokusu a omylu.

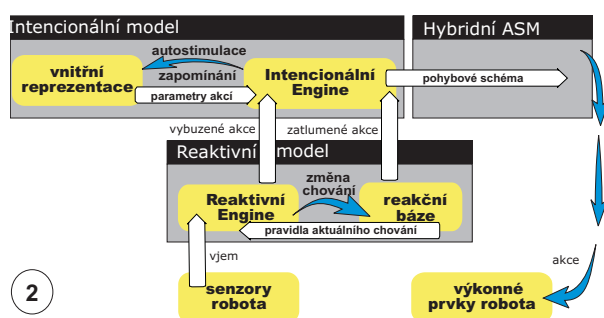
Moji další novinkou oproti uváděné podstatě intencionálního modelu je použití pohybového schématu. Výhody jeho parametrizace se ocení zejména při adaptaci.

(3)

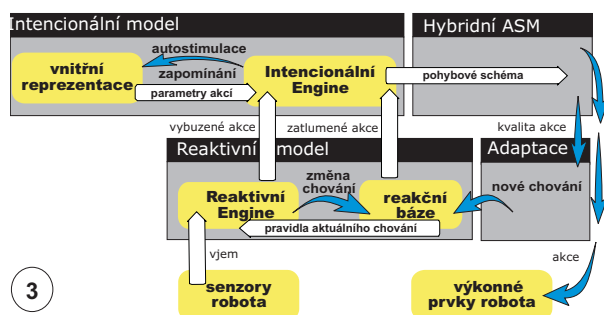
Intencionální model je následně rozšířen o *adaptaci*. Řídicí systém je schopen se adaptovat samostatně, bez učitele. *Adaptace je v mém návrhu postavena na vhodně navrženém kritériu diferenciální a integrální kvality řešení*. Je také plně podřízena parametrickému nakládání s akcemi. *Prostřednictvím pohybového schématu akce parametrizují*, a s optimálními parametry je i spouštím. A pokud navíc parametry nastavím adaptivně, odrazí vybraná akce aktuální situaci ještě věrněji. Během takovýchto testovacích kroků se uplatní jedno ze zavedených kritérií: **Pokud se bude optimalizovat užitek agenta**, uplatňuje se *diferenciální kritérium* δ . **Pokud se optimalizuje v rámci kolektivu**, využívám *integrální kritérium* ϵ .

Nyní ještě shrnu pojmy, jenž se v této 3 kapitole objevují nově. Používám je pro přiblížení mého vlastního návrhu v následujícím významu:

- **Paměť systému:** Rozšiřuje prostý reaktivní systém na hybridní. Paměť udržuje vnitřní reprezentaci problému. Je v ní možné přechodně držet předchozí průběh podnětů v hodnocení akce. *Pro pozitivní hodnocení* užívám *atraktor* (též atraktivní vzor), *pro negativní hodnocení repulsor*.
- **Hybridní řízení:** Integruje podněty z více modelů chování. V paměti systému nalezneme z reaktivního modelu vjemy, z intencionálního modelu průběh pozitivního ohodnocení akce a ze sociálního modelu atributy objektů zcela mimo dosah percepce. Zajímavé je, že přirozený hybridní systém je i celý kolektiv. Vývoj prostředí závisí na stavu agentů, a evoluce agentů závisí na stavu prostředí.
- **Vzor:** Kombinace podnětů (senzorických podnětů či vnitřních stavů), jenž představují nutnou podmínku pro spuštění akce.
- **Chování:** Chování bude mít i svou interpretaci sadou reaktivních pravidel. Jaká sada pravidel se použije, bude záviset na aktuálním chování.
- **Stavový prostor:** Stavovým prostorem zde rozumím prostor, definovaný v Teorii řízení. V něm atraktivní vzor nemusí mít pouze geometrický význam.



2



3

3.2 Zvolené řešení realizované hybridním ASM

V celé této kapitole přibližuji má *specifika řízení pohybu robotů v Kolektivu*. Pro řízení jeho komponent, hybridních agentů, včlením – integruji reaktivní a intencionální model chování. Reaktivní model je víceméně dán základními pohybovými rutinami robota. V *intencionálním modelu* pro tato základní chování počítám pracovní bod, v němž se vybraná akce využije v nejvyšší kvalitě. S ohledem na závěry předchozí kapitoly to pro mne znamená, že

- při *Orientaci robota v herním poli* se entita, krom sémantiky, bude vyznačovat i kvalitou vjemu,
- při *Výběru akce s maximálním užtkem* se hrubý sémantický výběr akce zlepší, pokud akce bude volána s optimálními parametry,
- a při *Koordinaci akce v týmu* vymezit nejenom sémantiku explicitní komunikace, ale i poměr započtených cizích podnětů.

Vlastnosti pravidlového systému posuzuji odděleně od optimálního buzení akce. *Podstatou mého návrhu hybridního mechanismu výběru akce – ASM je **buzení akcí podle pohybového schématu***. Zahrnuje v sobě i **pravidlovou část**. Nejprve podle reaktivního modelu hrubě vymezím aktivity podmínkami pro jejich realizaci v okolí robota vymezím vše, co robot může dělat. Z nich potom podle intencionálního modelu jemněji určím přímo směr nejbližšího postupu. Může vyjít i na základě kompromisu mezi několika přípustnými oblastmi.

A. Sémantika a kvalitativní uvažování o vjemu

Prvotní v mém ASM je lokální vjem. Reaktivní model k vjemu v reakční bázi udržuje vazbu na akce. Reakční bázi plním při inicializaci systému. Další pravidla se do modelu zavádí adaptací, kdy se k dané kombinaci senzorů nově přiřadí ověřená reakce.

Intencionální model navíc ke každé unikátní kombinaci senzorů a buzené akce udržuje metrické atributy — *vzdálenost a azimut vjemu*, či *míru základních chování za aktuální a předešlé situace*. Využívám vjem i jako kvalitativní zpětnou vazbu k provedené akci (slovy gramatiky multiagentního systému: každá realizace odvodí nový stav prostředí, v němž operuje), a pohybové schéma adaptuji metodou pokusů a omylů.

B. Buzení akce v jejím efektivním pracovním bodě

Intencionální prvek se u akcí projevuje v jejich hodnocení. Ve vnitřní reprezentaci se hodnotí, nakolik akce přispívá k uspokojení stanoveného cílového stavu. Při aktuálním hodnocení každé akce se díky mechanismům *autostimulace* a zapomínání zohlední i hodnoty předchozích akcí. V následující kapitole se zaměřím na parametry k uchování takových hodnot. Ty uplatním v pohybovém schématu. Na podkladě takového *parametrizovaného pohybového schématu* získám gradient. Ve směru jeho nejvyššího růstu zahájí robot zvolenou akci.

C. Využití explicitní komunikace

V kolektivu robotů pak tentýž hybridní systém zpracovává i koordinaci mezi agenty. Je v něm třeba pouze zajistit, aby správně interpretoval implicitní a explicitní vazby mezi členy kolektivu. Už v reaktivním modelu robot vnímá z prostředí implicitní vazby – podněty vypovídající o aktivitě ostatních. V závěru této kapitoly se navíc zamyslím nad komunikací robotů. Předmětem takové explicitní vazby budou parametry vnitřní reprezentace. Tyto parametry obohatí intencionální model příjemce. Robot s nimi bude moci ve svém odvození vážit i vliv zatím předpokládaných objektů mimo dosah svých senzorů.

3.2.1 Algoritmus vytváření hybridního ASM

Mé navržené řešení je založeno na sekvenčním skládání řešení z bazálních akcí. V hybridní architektuře ASM *kombinuje dva dříve uváděné principy superpozice a převrstvování*.

- Sekvence řešení se vytváří průběžně při jeho realizaci. Robot v daném okamžiku provádí právě jednu bazální akci.
- pravidla pro použití akce jsou vlastní každému chování. Chování tedy představuje signál pro převrstvování novou sadou bazálních akcí.

Pro přehlednost uvádím **Algoritmus 1** činnosti hybridního Mechanismu výběru akcí ASM v jednotlivých krocích.

Algoritmus 1 Hybridní ASM

- 1: {inicializace reaktivního modelu}
 - všechny akce v modelu převed' do stavu «zatlumená»
 - inicializuj *pravidla aktuálního chování*
 {následuje aktualizace reaktivního modelu:}
 - 2: **pokud** jsou k dispozici vázané parametry individuálního *pravidla aktuálního chování*,
 - 3: vyhodnoť pravidlo – pravdivé je za udaných podmínek a kontextu,
 - 4: **jinak**
 - 5: pokračuj bod 2 pro další individuální *pravidlo aktuálního chování*
 - 6: akce, které plynou z pravdivých pravidel, převed' do stavu «vybuzena»
 - 7: {změny v tendencích v intencionálním modelu}
 - váha vybuzené akce se zvyšuje o váhu pravidla, z něhož vyplývá
 - váha zatlumené akce přispívá ke zvýšení váhy příčinné akce
 - váha přechodných senzorických podnětů se sníží zapomínáním
 - 8: {volba akce v intencionálním modelu}
 - najdi maximální váhu přes všechny vybuzené akce
 - urči akci, včetně jejich optimálních volných parametrů
 {a dále interpretace závěru z intencionálního modelu:}
 - 9: **pokud** zvítězí pravidlo na změnu chování,
 - 10: změň chování a zopakuj výběr akce od bodu 1
 - 11: **jinak**
 - 12: proved' akci
 - 13: sniž váhu pro výběr téže akce v následujícím kroku
-

A. Chování a jeho interpretace

Chování chápu jako ucelený soubor pravidel opírající se o bazální akce. O tom, která akce je vybuzena, případně zatlumena, rozhodují nejen aktuální podmínky v prostředí řešeného problému, ale také pravidla daného chování.

B. Sekvenční chování

Hlavní kroky hybridního mechanismu výběru akce vystihuje **Algoritmus 1**. V podstatě je hlavním podnětem pro aktivitu agenta reaktivní model. V takovém případě intencionální model filtruje na jednu jedinou variantu. v okamžicích, kdy je reaktivní model v klidu, intencionální model proaktivně spouští akce, které nevyžadují přímý senzorický podnět.

Zvolené chování pak bude i nejčastější podmínkou v detailních pravidlech. Samotný fakt chování tak zužuje výběr akcí. Výběr proběhne mezi akcemi, jenž dané chování realizují. V navrhovaném mechanismu stejně jako ostatní autoři **zapojují akce na principu vzájemného převrstvování**:

- V ASM se potlačuje průchod odvození, jenž s aktuálním chováním nesouvisí.
- Změna chování je inicializována signálem. Jedním takovým pravidlem lze vybudit celou sadu následných závěrů.

Oba dva kroky jsou naznačeny i v **Algoritmu 1**. Z pohledu chování mechanismus výběru akce ASM pracuje pochopitelně vždy než v jednom z konečné množiny chování. Podle definovaných pravidel pak může přejít do nového stavu, charakteristického novým, právě aktivovaným chováním. *Vymezení nové vnitřní podmínky* pro použití vhodnější alternativy (jiné akce) implementují také jako *zavádění nového chování*. V takovém případě *reflektují vzor*, který vystihuje *důvody pro použití alternativního řešení*.

Shrnu-li:

V této podkapitole 3.2.1 jsem *zavedl strukturu hybridní architektury řídicího systému*. Tuto strukturu budu kompletovat z modulů podrobněji popsanych v podkapitolách následujících. V samostatné pasáži jsem popsal svoji základní ideu (Algoritmus 1) hybridního ASM. včetně popisu včlenění reaktivních a intencionálních procedur do tohoto algoritmu. V souhrnu: vybuzené akce určuje reaktivní model. Intencionální model tuto rámcovou filtraci rozvine s ohledem na další vliv podnětů na možná jiné vybuzené akce, jim významově podobné. S intencionálním modelem tak vytvářím již hybridní ASM. Nově zapojený mechanismus výběru akce kromě nejvhodnější akce také určuje s jakými parametry je nejvhodnější vybranou akci třeba použít (za aktuálního stavu řešení).

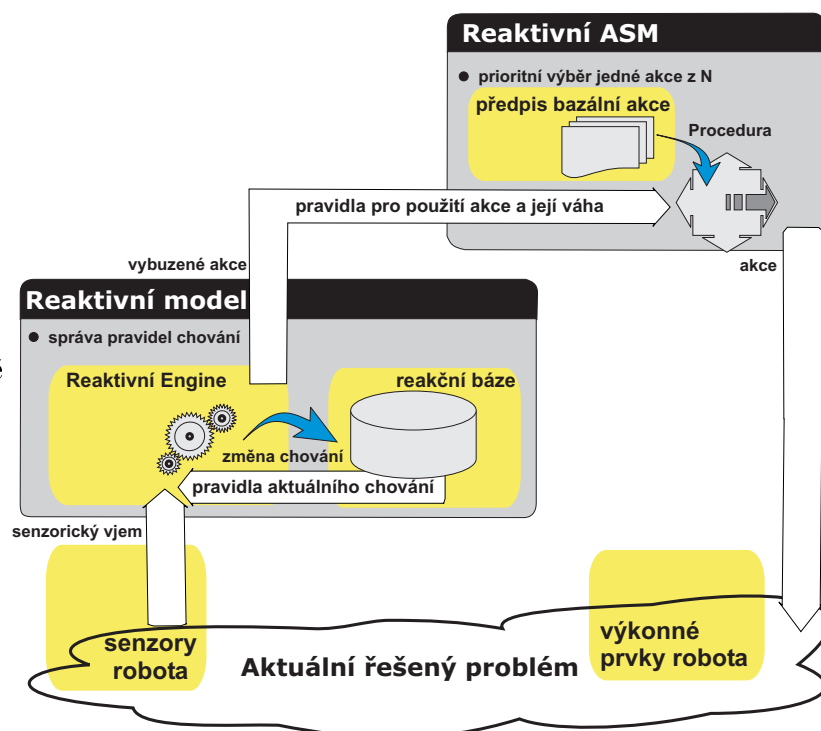
Realizaci nastíněných procedur detailněji rozeberu ještě v následujících kapitolách.

3.3 Reaktivní část hybridní architektury ASM

Mnou zvolená *architektura řídicího systému* je *modulární*, postupně rozšiřitelná od jednoduchých komponent pro reaktivní ošetření aktuální situace až po složitější sémantické vyhodnocení řešení v delším časovém intervalu.

Nejjednodušší samostatně pracující částí navrhovaného řídicího systému je *reaktivní model*, *deterministicky vyhodnocovaný v reaktivním ASM*. V tomto spojení je na výběr několik způsobů, jak reaktivní ASM bude kombinovat fyzickou reakci z vybuzených bazálních akcí reaktivního modelu. Autonomní funkci reaktivního modelu demonstrují na prioritním ASM a poté rozšířím dalšími principy hybridního ASM.

S ohledem na další kroky návrhu proto nejprve uvádím schématické propojení základních bloků reaktivního modelu na **obrázku 3.1**. Zde je ukázán i způsob zavedení pravidel pro použití jednotlivých akcí. *Jádro vyhodnocení reaktivního modelu spočívá v reaktivní inferenci bazálních akcí*. Při ní se zpracovávají aktuální senzorické vstupy ve vztahu k pravidlům aktuálního chování. Pravidla aktuálního chování poskytuje reakční báze. Dále popsany *prioritní způsob aktivace akcí* pak zajistí, aby v daném okamžiku byla užita akce, která plyne z pravidla ohodnoceného nejvyšším pravděpodobnostním faktorem.



Obrázek 3.1: Zavedení pravidel v reaktivní části řízení. Samostatný chod reaktivní části architektury zajišťuje konečný automat. Stavem reaktivního modelu je chování. Podle chování se z reakční báze volí pro datovou část řízení *pravidla aktuálního chování*. Pro jejich kombinaci se *senzorickým vjemem* *Reaktivní Engine* inicializuje *Reaktivní ASM* seznamem reakcí. Reakcí na aktuální senzorický vjem je buď přímo dílčí akce robota, nebo pouze změna rozhodovací strategie – chování. Reaktivní ASM pak podle nejvýše hodnocené reakce pak inicializuje další změnu v rozpracovaném řešení: akce mění stav prostředí, a v případě chování se mění vnitřní stav automatu.

3.3.1 Interpretace prostředí a akcí v pravidlech reakční báze

Na začátku návrhu reaktivního modelu řídicího systému rekapituluji *dva základní kroky návrhu minimální reprezentace reakční báze*. Pro vymezení reaktivní části budu definovat:

- **základních entity reakční báze**, tedy *abecedy minimální reprezentace problému*, použité při řešení,
- **vazby mezi entitami**, s přihlédnutím k jejich možnosti parametrizace.

Těmto krokům se budu věnovat v popisu experimentů v kapitole 4. Jejich formalismus ovšem mohu rozvést už nyní, v kontextu s dále uváděnými teoretickými předpoklady hybridní architektury řídicího systému.

A. Reprezentace entit hybridního agenta

V hybridním systému půjde *nově* především o *jemnější výběr akce*. Podmínky pro *hrubý výběr akce* lze popsat gramatickými pravidly (viz kapitola 2.3.2). O *hybridních agentech* tak uvažuji v paralelách s gramatickým systémem; a shrnuji je v **tabulce 3.1**. Požaduji, aby se *jednotlivé iterace autonomních přístupů dílčích řešení uplatnily přímo na řešení problému celého kolektivu*. To vyplývá z korespondence (z určité paralely) s Lindenmayerovým systémem. I jeho dílčí řešení se průběžně rozvíjí v paralelním odvození.

značení	gramatický term	entita/operace v kolektivu robotů s hybridním ASM
V	abeceda	sada bazálních – základních akcí (jednotlivce, jejich sjednocení pak pro kolektiv)
Σ	formální parametry (<i>vázané proměnné</i>)	seznam mandatorních parametrů – kvalit, které je nutné zvážit při aplikaci pravidel
$\mathbb{R}^n$	posloupnost parametrů	aktuální hodnota uvažovaných kvalit, jenž ovlivní realizaci bazálního – základního chování – např. pohybový vektor
m	modul $m = V \times \mathbb{R}^n$	vjem nebo akce, zpřesněné dodatečnými parametry
χ	slovo $\chi \in (V \times \mathbb{R}^n)^+$	aktuálně předvedené chování v prostředí, tj. výsledky aktivity kolektivu
pred, lc	předpoklad a levý kontext pravidla ($\in \chi^*$)	kombinace senzorů, které identifikují specifický případ
cond	podmínka pro aktivaci	chování
succ	závěr pravidla	vybuzená akce
$C(\Sigma)$	Logické výrazy	reaktivní část řízení
$E(\Sigma)$	aritmetické výrazy	intencionální část řízení
ω	axiom	popis počátečního stavu problému
ψ	akční zobrazení	ASM jádro (prioritní nebo pravděpodobnostní)
R	akční pravidla	sémantika akce v popisu problému
P	evoluční pravidla	ohodnocení vjemu (perceptuální schéma)
φ	evoluční zobrazení	sémantika vjemu
W	pravděpodobnostní faktor	váha pravidla

Tabulka 3.1: Paralela parametrizovaných gramatických termů a entit hybridního a adaptivního ASM. V předloženém výčtu rozšiřuji korespondence čistě reaktivního modelu z Tabulky 2.1 o transformace intencionální podstaty. Nově v tabulce zmiňuji parametrizaci pravidel, ekvivalentní programovým procedurám. Dalším logickým znakem mého návrhu je aktivační podmínka, která zatlumí či povolí podnět. Užívám též obecnější, nedeterministické zobrazení – akčním zobrazením se akce projeví, evolučním zobrazením se projev akce vnímá v následujících krocích.

V předloženém výčtu *rozšiřuji korespondence původně čistě reaktivního modelu z Tabulky 2.1* (kapitola 2.3.2) *o transformace intencionální podstaty*. Nově v tabulce

zmiňuji parametrizaci pravidel, ekvivalentní programovým procedurám. Dalším typickým znakem mého návrhu je **aktivační podmínka**, která *zatlumí* či *povolí* *podnět*.

Způsob integrace (zapojení) paralelních bloků ASM do **Kolektivního řešení** je právě v **tabulce 3.1**. Opírám se o hybridní agenty. *ASM každého agenta pracuje s vlastní sadou modulů.*

Nejdůležitější vlastnosti:

- Optimální pracovní bod akčního modulu vymezují *volnými proměnnými*. V jejich dimenzi hodnotím aktuální užitečnost modulu z lokálního pohledu.
- Přes senzorický modul se do ASM může dostat odraz integrálních kvalit Kolektivem řešeného problému. K tomu užívám schopnosti *P1L-systémů* akceptovat jak aritmetické výrazy, tak i rozmístění podnětů pro další akci. Například vnímaný směr podnětů a jejich intenzitu dobře vymezuje polohu robota vůči Kolektivem vymezeným atraktivním oblastem při gradientní navigaci. Při sledování je naopak integrálně důležitější kontext.
- Standardně zavedený pravděpodobnostní faktor interpretuji nejčastěji jako prioritu při gradientním výběru. Původního významu náhodného výběru užívám při navigaci více robotů v úloze v kapitole 4.

B. Parametrizace, kontext a podmínka v použitém ASM

V reaktivní části mého návrhu hybridního řízení pravidla seskupuji do jednotlivých chování. Jelikož neorganizuji pravidla jen událostně, mohu razantně filtrovat podněty přímo podle situace a výsledku uvažovat pouze o vybuzených akcích. Jejich seznam předurčuje aktuální možnosti v intencionálním modelu hybridního agenta.

Jiná sada podnětů pochází z doplňku vybuzených akcí do reaktivního modelu robota. Hovořím o zatlumených akcích. Ty autostimulací podle [Maes, 1991] v intencionálním modelu dále podporují z již vybuzených akcí zejména takové, co aktivně připraví podmínky ke spuštění zatlumených. Takové preference příčin požadovaných stavů se odráží i v *prioritách* ASM. Co to znamená: Na mechanismu výběru akce ASM v intencionálním modelu zvážit přínos akce vzhledem k jejímu účelu, zvolit nejprínosnější akci a určit její optimální pracovní bod.

Reaktivní model poskytuje pro rozhodování více atributů, než jen aktuální platnost odvození ze senzorů na přímo vyplývající reakce. Mým záměrem je proto takové dodatečné atributy reakce podržet v průběhu celého vyhodnocení jako přímou součást pravidla. Pravidlo rozšiřuji na popis celého faktu: nyní *v pravidle rozhoduje nejen aktuální kontext*, ale i *dodatečná podmínka* či *stochastická pravděpodobnost* (při použití víceznačných pravidel). Ve zbytku této práce je budu značit v notaci

$$\text{id} : \text{context} : \text{cond} \mapsto \text{succ} : W \quad (3.1)$$

V pravidle pod částí **context** stále rozumím gramatické termy, tak jak byly ostatně uvedeny i v **tabulce 3.1**. Tamtéž byly zavedeny i podmínky **cond**, závěry **succ** a pravděpodobnostní faktor **W**. Jednotlivá pravidla zde označuji identifikátorem **id**.

3.3.2 Práce se soubory chování

Podmínka cond v notaci pravidla (3.1) popsaného v **tabulce 3.1** určuje, zda bude toto pravidlo v daném okamžiku použito. V mém modelu je většina pravidel vztažena k určitému chování a platí pouze pro ně. Takovou přímou explicitní interakcí vymezují pro každého agenta použití pravidel mnohem nižší úrovně. *Chování nadřazují akcím na stejné – jednotné úrovni*. O složitější hierarchii mezi chováními už neuvažuji. Od logických sekvencí a přechodů mezi chování například bude vycházet i mé přizpůsobení reakční báze pro řešenou úlohu. Analytická dekompozice problému do chování bude více zřejmá z popisu jednotlivých experimentů. Procesy v každé úloze budou na začátku popsány ve stavech použitelných chování a stavových přechodech mezi nimi.

Jedním ze způsobů, jak určit *vzor pro inicializaci stavového přechodu na nové chování, je adaptace*. Vzory zavedené adaptací přímo odrážejí podmínky, které se adaptací měly řešit. Adaptace systému bude ukázána v následujících podkapitolách.

A. Parametrizace a kontext v gramatickém vyjádření

Nejbližší gramatickou podobnost s mnou navrženým algoritmem shledávám u Lindenmayerova systému. I zmíněný L-systém vystihuje a do požadovaného řešení zahrnuje rozvoj aktivit, generovaných několika paralelně pracujícími roboty ve společenství. A aby byla reakce dostatečně autonomní, i mnou zaváděná pravidla by měla splňovat náležitosti **Lindenmayerova systému rozšířeného o kontext**, takzvaný **P1L-systém**. V případě naznačeného algoritmu jde zejména o volné a vázané parametry. Odpovídající parametrickou definici systému opět přebírám z [Lindenmayer, 1968b].

Definice 3.3.1 Kontextové parametrické L-systémy (P1L-systémy)

K abecedě V mějme

- množinu všech (konečných) posloupností parametrů $\mathfrak{R}^* = \{a_i^{(n)}\} \ a_i \in R$
- množinu formálních parametrů Σ
- množinu logických výrazů $C(\Sigma)$ správně utvořených z formálních parametrů
- množinu aritmetických výrazů $E(\Sigma)$ správně utvořených z formálních parametrů
- množinu všech neprázdných parametrických slov nad abecedou V $(V \times \mathfrak{R}^*)^+$

Pak parametrický L-systém je uspořádaná čtveřice

$$G = (V, \Sigma, \omega, P)$$

S významem

- ω - axiom - neprázdné slovo $\omega \in (V \times \mathfrak{R}^*)^+$
- P - množina přepisovacích pravidel $(V \times \Sigma^*) \times C(\Sigma) \times (V \times E(\Sigma))^*$

□

Interpretace P1L-systému zůstává přirozeně paralelní. *V jednom odvozovacím kroku se nezávisle na sobě mohou všechna pravidla podílet na modifikaci (přepisu) předloženého popisu situace*. Je zřejmé, že se takto paralelně do robota dostane *reprezentace stavu prostředí* (v gramatických termtech evolučního pravidla). Proti paralelní realizaci dílčích aktivit robota už stojí nejenom konstrukční omezení. Například v mém návrhu algoritmu paralelní (další) aktivity jednoho robota vylučují exkluzivním výběrem akce.

Dále je pro analýzu optimálního postupu za aktuálních podmínek důležitá nová dimenze, *parametrizace Lindenmayerova systému*. V parametrickém L-systému lze popsat celou třídu případů jedním odvozovacím mechanismem. Výsledek jejich aktivit konverguje zejména při jednotném odvozovacím mechanismu ASM.

3.3.3 Reaktivní ASM

Počáteční úvaha nad praktickým zapojením modulů do Lindenmayerova systému je v korespondenci na mechanismus výběru akce individuálního robota. Na chování robota zde přiblížím sekvenčním přepis; a ve stejném principu robot inkrementálně generuje svoji novou reakci. V navrhované architektuře vlastní logiku iterativního vyhodnocení takových podkladů definuji v reaktivním ASM. Poslouží mi k řízení pohybu robota podle pravidel dosud rozváděného Lindenmayerova systému. Taková pravidla zde přímo určují způsob stabilizace pohybu robota. Návrh reaktivního ASM stavím na odvození bezprostředního pohybu robota z modulů systému tak, aby se stabilizovala relace robota vůči okolním objektům. **V praktické realizaci reaktivního ASM používám:**

- vázaných parametrů při přesném zacílení reakce,
- váhu pro použití akce při volbě z více alternativních reakcí,
- a časově omezené platnosti akce.

Tyto aspekty odráží logika mnou navrhovaného Mechanismu výběru akcí ASM. Podstatu Výběru konkrétní – výlučné akce vystihuje následující Algoritmus 2:

Algoritmus 2 Výběr výlučné akce v reaktivním ASM

Předpoklady: prováděná akce, seznam vybuzených akcí (+vázané parametry a priority)

- 1: z vybuzených akcí vyber akci nejvyšší priority
 - 2: **pokud** akce je parametrizovaná, {je třeba uvážit vázané parametry}
 - 3: volání akce s vázanými parametry, akce může být přerušena pouze vyšší prioritou
 - 4: **jinak** {zpětná vazba pro reakci na obecný podnět}
 - 5: spustí akci, jakákoliv změna rozhodnutí akci přeruší
-

Reaktivní mechanismus výběru akce byl navržen úmyslně s maximální jednoduchostí. Podpora reaktivního modelu ovšem nepokrývá mezní stavy mimo předdefinované situace. Tak při malém počtu předdefinovaných reakčních pravidel není žádoucí dílčí závěry pravidel ještě skládat v superpozici. Vždy se z několika přípustných variant projeví výlučně jedna bazální akce.

A. Použití vázaných parametrů při přesném zacílení reakce

V reaktivním modelu má význam uvažovat o přesném zacílení pouze u čidel s dostatečnou rozlišovací schopností. Nepřesnosti, které plynou v závěru vyhodnocení vjemu z čidel s nízkou rozlišovací schopností, řeším až při vyhodnocení minimální reprezentace problému v intencionálním modelu — tedy v podkapitole 3.4. Nyní tedy předpokládám jen *čidla s vysokou rozlišovací schopností*, která jsou použitelná ve všech pravidlech aktuálního chování:

1. *U pravidel s vázanými parametry* se akce odvozuje přímo z měřené hodnoty.
2. *V neparametrizovaných pravidlech* se takové čidlo projeví pokud je jeho měřená hodnota nenulová.

B. Váha pro použití akce při volbě z více alternativních reakcí

V návrhu se *omezují na převrstvování akcí*. Volba bazální akce je podpořena přímo reaktivním modelem. Ten dodává pro ASM seznam vybuzených akcí s jejich ohodnocením. Fakticky je takové ohodnocení dáno pravděpodobnostním faktorem W aktuálně platných pravidel. Tento faktor přímo určuje i váhu akce, používám proto pro ni opět označení W . Jak ukazuje obrázek 3.2(c), *všechny ostatní akce převrstvuje (přebíje) bazální akce s nejvyšší vahou*.

C. Časově omezená platnost vybrané akce

S přesným zacílením je úzce spjata i doba platnosti následné reakce robota. Na stabilizaci pohybu se podstatně projeví doba trvání bazální akce. Tato doba je opět *odvozena od rozlišovací schopnosti použitých čidel*:

1. **Při nízké rozlišovací schopnosti** volím *těsnou zpětnou vazbu*. Rozhodnutí následuje periodicky s každým vyhodnocení senzorů. Platnost akce v závislosti na stavu čidel s nízkou rozlišovací schopností je ukázán i na průběhu akční sekvence v reakci na infračidla na obrázku 3.2(a)
2. **Při vysoké rozlišovací schopnosti** dávám v reaktivním ASM *důraz na ovládání*. Jeden krok takové reakce je naznačen na **obrázku 3.2(b)**. Senzory se zde vyhodnocují ve stejné periodě, avšak ASM vyhodnocuje tento měřený stav vždy po několika měřeních. Do té doby se ASM vyhodnocuje pouze pravidla bez vázaných parametrů, neboť vzhledem k přesnější specifikaci akce lze předpokládat i delší dobu pro její užití. Získaný čas je věnován na filtraci měřených hodnot.

V obou případech nastává *změna aktuální bazální akce* nejenom 1) *po ukončení doby její platnosti*, ale také v okamžiku, kdy 2) *je vybuzena akce s ještě vyšší vahou*.

D. Praktické aspekty reaktivního mechanismu výběru akce ASM

Diskutovaný reaktivní ASM ilustruji na následujících dvou příkladech. Odděleně posoudím reakce robota s rozdílnou časovou platností dílčího závěru (s rozdílnou dobou trvání poslední vybrané akce). V prvním příkladu v *reakci robota na proximitní čidla* dostáváme *trajektorii složenou z krátkých kroků*. Druhým příkladem je *souvislejší reakce robota, založená na vektoru gradientu*. Vektor se v toto příkladu odvozuje z hodnot, získaných ze světelných čidel.

Uvažme postupně obě varianty pro jednoho a téhož robota. Jednotnou abecedu V , předepisující chování robota, stejně jako její užití ve formátu P1L-systému shrnuji v Tabulce 3.2 jako příklad definice pravidel pro reaktivní Mechanismus výběru akce ASM.

Axiomy řešených případů vyplývají ze situace na dílčích obrázcích 3.2(a), 3.2(b), 3.2(c), 3.2(d). Logické výrazy $C(\Sigma)$ zde nyní nefigurují. Podrobněji nyní k jednotlivým příkladům:

Reakce v hrubém rozlišení

Odvození reakce z proximitních čidel je značně omezeno jejich minimální rozlišovací schopností. Proximitní čidla popisují bezprostřední okolí robota ve dvou stavech. Robot se dozvídá pouze o detekci překážky, a jeho reakce musí dostatečně odrážet už její změnu. Okamžik detekce překážky proximitními čidly v čelním směru například pro robota signalizuje potenciální riziko kolize se stěnou. Reaktivní ASM aktualizuje zvolenou akci po každém vyhodnocení senzoru. Bezprostředně poté, co je zaznamenána změna hodnoty senzoru, reaktivní model aktualizuje seznam vybuzených akcí včetně jejich vah. Pokud reaktivní ASM vyhodnotí v tomto seznamu akci s vyšší vahou, namísto dosavadní bazální akce aktivuje. Výběr nejvhodnějšího pravidla chování sledování zdi podle jeho váhy je patrný z grafu na obrázku 3.2(c). Na průběhu váhy se většinou uplatní pohyb vpřed A_s . Jeho váha je v případě výpadku postranní detekce infračidlo I_{11} převrstvena vahou otáčení A_s . Obdobně se podle pravidel P robot otáčí napravo při detekci překážky infračidlem I_{16} v čelním směru.

položka P1L	hodnota	poznámka
abeceda V	$A_s, A_j, A_v, A_z, I_1 \div I_{16}, \mathbb{C}$	A =akce, I = infračidla, $\mathbb{C}$ =světlo
formální parametry Σ	$\alpha, \quad \alpha \in (-\pi; \pi)$	směr gradientu
aritmetické výrazy $E(\Sigma)$	$\alpha = \arctan \frac{C_s - C_j}{C_v - C_z}$	odvození gradientu z bazálních složek intenzity světla
pravidla P	1: $I_{16} \mapsto A_v$: 0.7 2: $I_{11} \mapsto A_z$: 0.6 3: $I_x \mapsto A_s$: 0.5 4: $\mathbb{C}(\alpha) \mapsto A_s(\alpha)$: 0.5	pravidla 1÷3 vymezují chování robota při sledování zdi, pravidlo 4 udává směr při jízdě za světlem

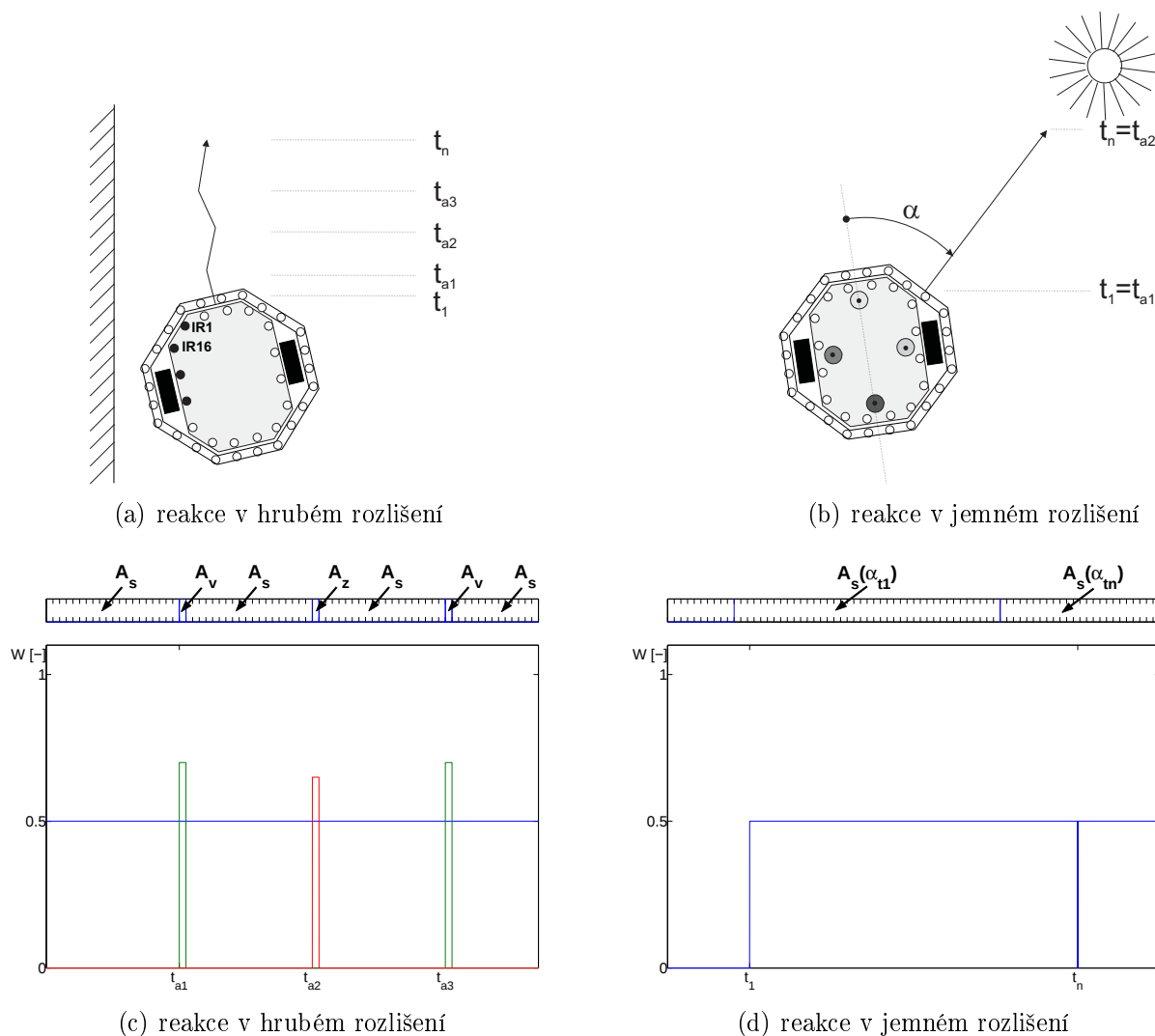
Tabulka 3.2: Příklad definice pravidel ve formátu P1-L pro reaktivní ASM. Abeceda dále uváděných příkladů je postavena na proximitních čidlech $I_1 \div I_{16}$, filtrovaném gradientu světla $\mathbb{C}$ a bazálních pohybových akcí – jízdě vpřed A_s , vzad A_j , a otáčení doprava A_v , respektive doleva A_z . Jízdu vpřed lze odvodit bezparametricky, nebo s parametrem $A_z(\alpha)$. V takovém závěru se robot nejprve otočí o úhel α a pokračuje jízdou vpřed. Formální parametr α je dán gradientem světla a určen ze hodnot intenzity světla měřených v pohybových bázích s – j – v – z .

Pravidla se pak definují pro moduly tvořené z abecedy a formálních (vázaných) parametrů. Jízda vpřed je v pravidlech 3 a 4 zastoupena dvěma moduly: A_s v neparametrické variantě udává pouze jízdu vpřed, zatímco při parametrické variantě $A_s(\alpha)$ se jízda vpřed nejprve modifikuje pootočením podle připojeného vázaného parametru α .

Reakce v jemném rozlišení

Na příkladu světelné navigace jsem zvolil specifické uchovávání váhy parametrizované akce $A_s(\alpha)$. Průběh přesně zacílené akce umožňuje více výpočetní kapacity věnovat na předpracování jejího vázaného parametru α . Ve smyslu aritmetických výrazů z definice P1L-systému tak ASM odvozuje směr pro další ovládání pohybu robota. Po dobu předzpracování $t_{a1} < t < t_{a2}$ proto nevyhodnocuje reálnou váhu akce $A_s(\alpha)$, ale stále předpokládá její konstantní váhu na poslední známé hodnotě. Váha se anulují až po dokončení přesně zacíleného pohybu. Ihned poté se odvozuje nová bazální akce. Může to být opět jízda vpřed. v detailu na obrázku 3.2(b) po ní následuje jízda vpřed s novým – aktuálním vázaným parametrem $A_s(\alpha')$.

Akce „Přesně zacílený pohyb“ může být přerušena také asynchronně, jakmile jiné senzory zaregistrují změnu stavu a jeho ošetření má vyšší váhu, než právě rozpracovaná akce $A_s(\alpha)$. Pokud se v průběhu předzpracování odvodí akce s vyšší vahou, $A_s(\alpha)$ se ukončí a následuje spuštění této nové akce s vyšší vahou.



Obrázek 3.2: Použití reaktivního ASM. Na obrázcích ilustruji práci s reaktivním ASM ve dvou případech. Na obrázku (a) je načrtnuta reakce na vjem s nízkou rozlišovací schopností. Na příkladu jízdy podél stěny takovou informaci dodávají proximitní čidla. Vybuzená čidla jsou na obrázku označena plnými body •. Obrázek (b) popisuje reakci na vjem ze světelných čidel ☉, úroveň šedi ilustruje senzorickou hodnotu.

Zpracování měřených hodnot probíhá periodicky v časech $t_1..t_n$. Ke změně reakce dochází buď při registraci změny senzory (v obou případech) anebo po dokončení přesně vymezeného kroku (příklad b.). Na následujících grafech je nejprve v časové ose vyznačena perioda měření senzorických hodnot. V tomto měřítku je vyznačena doba trvání jednotlivých bazálních akcí. K použitým akcím následující graf ilustruje, s jakou vahou je v daném okamžiku akce zahrnuta do reaktivního ASM. Na obrázku (c) je tak plnou modrou čarou vyznačen průběh váhy pro bazální pohyb vpřed A_s , zelenou čárkovanou otáčení doprava A_v a červenou tečkovanou otáčení doleva A_z .

Na obrázku (d) je uvažována pouze váha parametrizované jízdy vpřed $A_s(\alpha)$. Ostatní akce jsou v tomto příkladě nulové. Váha $A_s(\alpha)$ se uchovává po dobu $(t_{a1}; t_{a2})$. Poté ASM starou hodnotu anulují a podle předdefinovaných pravidel určí novou váhu. V případě dalšího výběru parametrizované akce se nová váha opět uchovává až do okamžiku jejího dokončení.

Toto přesné vymezení pohybu umožňuje ovládat robota v delších krocích. Čas jízdy se využívá pro pro zpřesnění měření – úhel α se určuje ze sumy hodnot měřeného senzoru $\alpha = f(\sum_{t_1}^{t_n} E(t))$.

3.3.4 Význam použitých modulů pro zpřesnění reakce

V této kapitole jsem začal s představováním navrhovaného algoritmu. Vyšel jsem z reaktivního modelu. V něm šlo o pravidlové odvození z kategorie „hrubého“ odvození. Z reaktivního modelu budu dále uvažovat *výčet přípustných akcí*, podle okolností „vybuditelných“.

Vybuzené akce mohou už od úrovně reaktivního modelu vyhodnocovat parametricky. Parametrizace je pro mě způsob, jak lze adaptovat znalosti na nové podmínky. V odvození jsem v P1L-schématu uvažoval jeho vázané parametry pro určení vybuzené akce. Volné parametry P1L-systému jsou pro upřesnění. Parametrický systém dále popisuje chování komponenty (hybridního agenta), který je už schopen autonomně reagovat na podněty, a navíc dokáže zvolené pravidlo realizovat v různé kvalitě. Aby taková kvalita byla ta nejlepší, to bude záležet na kvalitách mechanismu výběru akce. Práce s intencionálním modelem je náplní následující kapitoly 3.4.

Snažím se zkombinovat dříve vytypované perspektivní *mechanismy syntézy akcí*. Zapojuji je do svého jednotného hybridního mechanismu výběru akce. Podstatu hybridního ASM definuji právě v této podkapitole 3.3. Nejprve v návrhu *sleduji logiku převrstvování*. V navrhované architektuře převrstvování ošetřuji na úrovni reaktivního modelu chování. Podílu reaktivního odvození na celém chování byla vlastně věnována celá tato citovaná podkapitola.

Shrnu-li:

V prvním části podkapitoly 3.3 jsem po vzoru přímé analogie mezi gramatickým systémem a reaktivním ASM usadil svůj návrh do *souvislosti s parametrizovanými iterativními moduly Lindenmayerových systémů*.

Omezení na Lindenmayerovy systémy přímo vymezuje jednoduché atributy odvození zásahu robota: Předpokladem je *váha* a *podmínka zásahu*. Předmět zásahu je upřesněn v jednostranném kontextu. Navržený *Algoritmus hybridního řízení* z těchto atributů *odvozuje vybuzené akce a jejich pracovní bod*.

Ve druhé části pak ukazují důsledky použití iterativních modulů v paralelním řešení. Do funkce robotů se v navrhovaném řídicím systému promítají i jejich vnitřní stavy. Zpočátku *uvažují pouze o váhách vybuzených a zatlumených akcí*. Později rozšířím vnitřní stavy i o krátkodobě pamatované podněty. Již nyní lze ale vložením paměti významně doplnit reálné podněty, jenž robot získává percepcí. Kolektiv robotů s vnitřní pamětí taktéž splňuje předpoklady eko-gramatického systému. Robot v něm dostává nejen vyšší uplatnění pro své aktivity. Významně se rozšiřuje i dynamika chování robota

Ve třetí podkapitole zavádím reaktivní ASM jako samostatnou komponentu pro převedení závěrů reaktivního modelu na akční komponenty. Selektivní znakem pro výběr akce, a tedy i akční komponenty, je *pravděpodobnostní faktor pravidla*. V mém návrhu reaktivní ASM reálně vybírá vždy právě jednu akci. *Vybraná akce vyplyne ze závěru pravidla s nejvyšším pravděpodobnostním faktorem*.

V tomto souhrnném pododdíle se ještě jednou vracím k *paralele mezi parametrizovanými moduly a akčními komponentami*. Zdůrazňuji, že **parametrizace umožňuje pokrýt více případů jedním specializovaným pravidlem**. V tomto duchu budu s pravidly pracovat i v intencionálním modelu. Princip intencionálního modelu rozvádím v následující podkapitole.

3.4 Intencionální část hybridní architektury ASM

3.4.1 Blokové schéma ASM s reaktivní i intencionální částí

Z naznačeného algoritmu intencionální část určuje rozhodování celého hybridního systému od okamžiku, kdy již známe kandidáty na další iteraci v řešení. V tomto oddíle bude výklad dotyčné části Algoritmu 1 (od kroku 7) sledovat jednotlivé fáze, jak má intencionální nadstavba vznikat. Opět v nich identifikuji Problémy z úvodu této práce

- pro orientaci robota v prostředí vnitřní reprezentaci stavím na analyticky zpracovávaném pohybovém schématu,
- pro maximální užitek takovou reprezentaci zpřesňuji v pokusných krocích,
- a v mém pojetí dochází ke koordinaci akcí v týmu v okamžicích, kdy hybridní agenti komunikují o svých pohybových schématech.

Intencionální model prakticky ve všech těchto případech zjednodušuje superpozici v rámci chování. Krom podnětů konkrétní chování zde určuje i autostimulaci závislých akcí (v Algoritmu 1 k ní dochází v druhé části kroku 7). Tím zaručuji, že superpozice proběhne v jednotném významu. Navrhuji ASM, které tento intencionální model interpretuje exklusivně — akce je určena mechanismem *vítěz-bere-vše*(WTA).

A. Analytické vyhodnocení pohybového schématu

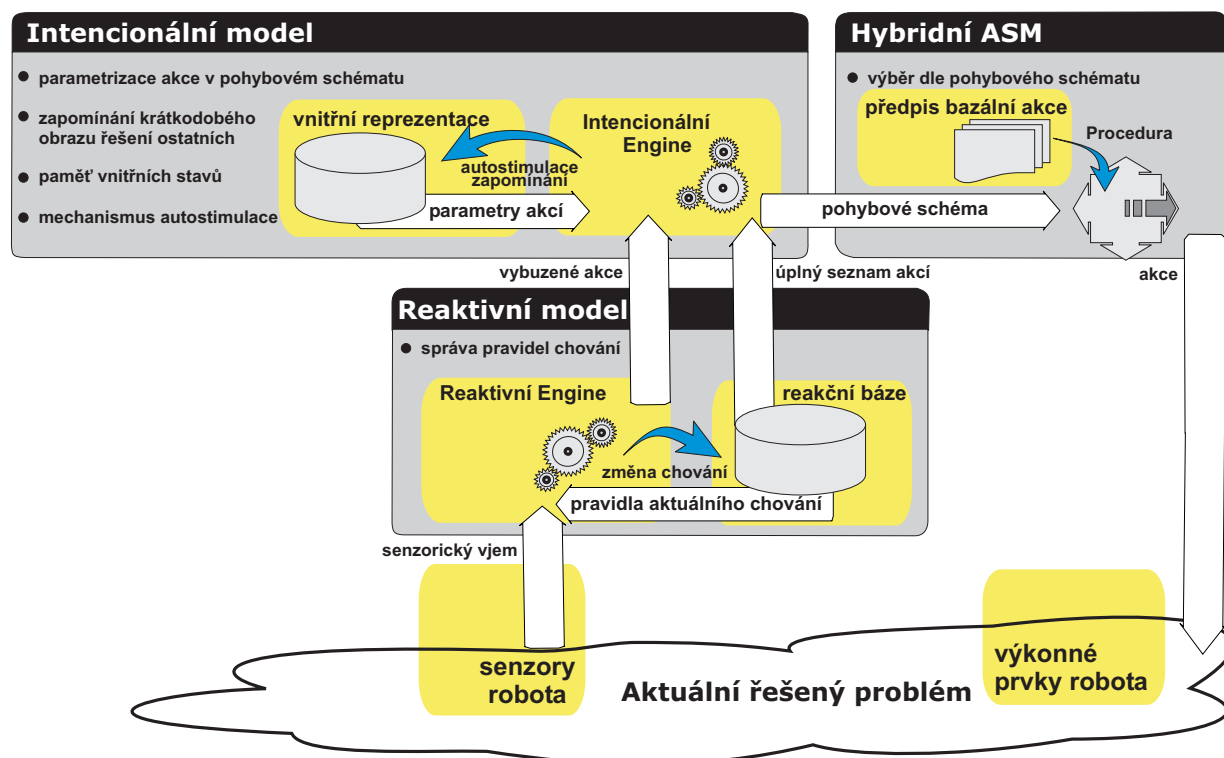
Ve výběru mé ASM zohledňuje doplňující pravidla; a z nich je vybrána akce lokálně optimální. Pro optimalizaci intencionální model z pravidel přebírá především pravděpodobnostní faktor. Posuzuji ho v celém pohybovém schématu, v dimenzích prakticky užitých při navigaci robotů v aréně. Tak už v této kapitole posuzuji v pohybovém schématu dva volné parametry, kartézské pohybové báze $[x, y]$, a přistupuji k jejich iterativnímu vyhodnocení.

B. Pokusný krok

V další fázi jsem se zaměřil na možnosti adaptace navrženého pohybového schématu. Po chybné reakci zařazuji automatické vyhodnocení chyby. Robot nejprve volí opravnou akci. Už v takové volbě figuruje zkušenost, který směr je nyní nevhodný. Ve zvoleném směru pak vykoná pokusný krok. Pokud se ukáže oprávněnost takového kroku, ukládá se směr pokusného kroku jako nové pravidlo.

C. Aspekty komunikace

V pohybovém schématu se z principu pracuje lokálně. Paměť, kterou nám intencionální model nabízí, se ovšem dá využít i pro započtení vzdálených podnětů, o kterých se robot dozvídá zprostředkovaně. V intencionálním modelu proto vidím významnou možnost pro rozšíření pohybového schématu. Především robot si vytváří obraz o situaci mimo dosah jeho percepce. V pokročilejší variantě navrhuji brát s obdobnou vahou i sdělené záměry druhého robota. Významové přiblížení dvou záměrů se extrapoluje podle směru pohybu obou robotů.



Obrázek 3.3: Hybridní architektura řízení rozšiřující reaktivní část o intencionální model chování. V hybridní architektuře kombinují reaktivního a intencionálního modelu chování. Reaktivní ASM nahrazují novým hybridním ASM, poskytujícím přesnější zacílení zvolené akce – doplňkové parametry pro buzení akce se určují z pohybového schématu. Na vybrané akci se tak projeví přidaná hodnota zavedeného intencionálního modelu – zahrnutí významově podobných podnětů, jejich vzájemná autostimulace a zapomínání krátkodobých poznatků.

V hybridní architektuře kombinují reaktivní a intencionální model chování. Reaktivní ASM nahrazují novým hybridním ASM, poskytujícím přesnější zacílení zvolené akce — doplňkové parametry pro buzení akce se určují z pohybového schématu. Na vybrané akci se tak projeví přidaná hodnota zavedeného intencionálního modelu – zahrnutí významově podobných podnětů, jejich vzájemná autostimulace a zapomínání krátkodobých poznatků.

3.4.2 Volba akce s autostimulací

Pro naznačenou stavovou reprezentaci uvažuji přímou vazbu na vstupy, stejně jako na další vnitřní stavy. Tyto vazby zavádím jako lineární, a tak s použitím známého stavového popisu lineárního systému je lze zapsat v rovnicích autostimulace

$$\begin{aligned} X_{sens}(t) &= \mathbf{B} \quad u(t) \quad + X_{sens}(0) \\ X_{act}(t) &= \mathbf{A} \quad X_{act}(t-1) \end{aligned} \quad (3.2)$$

Mohu tak popsat s jakou vahou W excituje konkrétní stav $[S_1, S_2]$ prostoru vnitřních akcí. Zdůrazňuji, že stav zde identifikuji třemi složkami – (S_1, S_2, W) . Proto nepracuji se stavovým vektorem, ale maticí rozměru $m \times 3$. Význam a dimenze všech matic uváděných v rovnicích (3.2) shrnuji v tabulce 3.3, kde m značí počet senzorů a n počet vnitřně reprezentovaných akcí. Fakticky pak rovnice (3.2) zahrnuje do váhy rozhodnutí dva aspekty, a

značení	dimenze	vlastnosti
X_{sens}	$m \times 3$	poloha jádra namapovaných senzorů S_1, S_2, W
X_{act}	$m \times 3$	poloha jádra akcí ve vnitřní reprezentaci S_1, S_2, W
A	$n \times n \times 3$	předpis, jak akce stimuluje své kauzální předpoklady – jiné akce, jenž je k ní ještě nutné vykonat. Stimuluje se pouze váhová složka, tedy $A_{ijk} = 0$ pro $i = j$ nebo $k \neq 3$
B	$m \times 3$	koefficienty mapování hodnoty senzoru $u_{(i)}$ $S_{1(i)}, S_{2(i)}, W_{(i)}$. Ve smyslu definice eko-gramatického systému 3.4.1 takto realizují evoluční zobrazení φ . Matice je diagonální – $B_{ijk} = 0$ pro $i \neq j$
$X_{sens}(0)$	$m \times 3$	stejnoseměrná složka takového mapování – udává, že senzor v klidu se mapuje na stav $[S_1, S_2]$ s vahou W

Tabulka 3.3: Význam a dimenze matic pro zpracování vjemů při volbě akcí s autostimulací. Tabulka upřesňuje proměnné užívané v rovnici autostimulace (3.2). Jádrem perceptuálních schémat jsou vnímané stavy prostředí. Jejich polohu a váhu ve vnitřní reprezentaci vystihuje matice X_{sens} . Stavy prostředí jsou přirozeně přímo úměrné vstupním sensorickým datům u . Akční schémata se naopak vyvíjí samostatně podle stavové matice **A**. Stavová matice vystihuje interakce na nejnižší úrovni, bez explicitního vlivu chování. Nepřímo se ovšem chování projevuje prostřednictvím aktivovaných elementárních entit; chování výrazně určuje zatlumené či vybuze akce, a tedy i jaká část stavové matice bude použita. Obdobně nepřímo se projevují sensorické hodnoty: pokud je detekován konkrétní vzor, dojde ke změně chování. Taková změna chování se v rovnici (3.2) projevuje novým nastavením stavové matice **A**.

to přímé podněty identifikované ve stavech X_{sens} , ale i očekávané důsledky, posléze ohodnocené ve stavech X_{act} . Jak po identifikaci stavů X_{sens} a X_{act} (v krocích algoritmu 1÷3) proběhne hledání nejlépe hodnoceného kompromisu, už bude tématem následující podkapitoly 3.4.3.

3.4.3 Zapomínání u intencionálního modelu

Navržený způsob iterativní akumulace hodnocení jednotlivých podnětů v sobě implicitně obsahuje též možnost iterativního procesu exponenciálního zapomínání. Obecně se započtením aktuální hodnoty zaváděné do intencionálního modelu zvnějšku platí iterace

$$\begin{aligned} x_0 &= u_0 \\ x_k &= qu_k + (1 - q)x_{k-1} \end{aligned} \quad (3.3)$$

kde q je koeficient exponenciálního zapomínání – časová konstanta dozrívání podnětu ve vnitřní reprezentaci problému v intencionálním modelu, $0 < q < 1$.

Možnost zapomínání v paměti intencionálního modelu využívám zejména v kombinaci se sociálním modelem. Uplatnění najde při jednorázovém zavedení již ohodnoceného faktu jiným robotem. Tak například v okamžiku akceptování pohybového schématu druhého robota vstupuje do intencionálního modelu jednorázový impuls udané velikosti váhy W . Iniciálně zadaná váha W následně dozrívá, s časovou konstantou q .

Procedura zapomínání se projevuje již v kroku 7 Algoritmu 1, který podává základní pravidla pro činnost Hybridního mechanismu výběru akcí. Ve zmíněném 7. kroku se váha přechodných sensorických podnětů sníží zapomínáním.

Nyní uvádím Algoritmus 3. Tento algoritmus v pěti krocích vystihuje operaci autostimulace a následnou volbu optimální akce. v operaci autostimulace je též zahrnuta funkce zapomínání krátkodobých podnětů.

Algoritmus 3 *Volba optimální akce s uvážením autostimulací*

-
- 1: {přímé buzení akce}
 - ke všem senzorům aktualizuj jejich **SensRepr**
 - 2: {autostimulace akcí a zapomínání krátkodobých podnětů}
 - interpretuj seznam vybuzených akcí
 - ke všem akcím a krátkodobým podnětům uži rovnice společenských potřeb (váhová matice **A**)
 - 3: {hledání extrému mezi podněty}
 - použít zobecněnou charakteristiku podnětu $N(\vec{S}, W, \sigma_1, \sigma_2, \varrho)$
 - urči mapování senzoru do pohybového schématu $[\mathbf{x}; \mathbf{y}]$
 - 4: urči extrém $[\mathbf{Xopt}; \mathbf{Yopt}]$
 - 5: inverzní zobrazení k $[\mathbf{Xopt}; \mathbf{Yopt}]$
-

3.4.4 Intencionální funkce popsaná eko-gramatickým systémem

Předchozí oddíl vymezil pravidla, která zamýšlím použít v ASM. V Umělém životě se prakticky tentýž princip pravidlového zpracování užije hned v několika krocích naznačeného algoritmu. Po sémantickém filtru v hybridním ASM obvykle následuje analytické zpřesnění, případně selekce. Mezi sémantickými částmi v kolektivním systému existuje i několik explicitních či implicitních interakcí. Z nich při integraci pravidel do hybridního systému rozeznávám:

- Evoluční zobrazení, kterým se deformuje skutečný stav problému do reaktivního modelu. Tyto deformace nejčastější plynou z charakteristiky použitých senzorů. Závěry reaktivního modelu se ihned přebírají do intencionálního modelu „evolučními“ pravidly.
- Akční zobrazení, které vnitřní reprezentaci interpretuje v Kolektivu.
- Popis prostředí, který se může měnit jak po akcích uvažovaných agentů, tak i dle vlastních pravidel (například zdroje zastarávají i bez zásahu agentů).

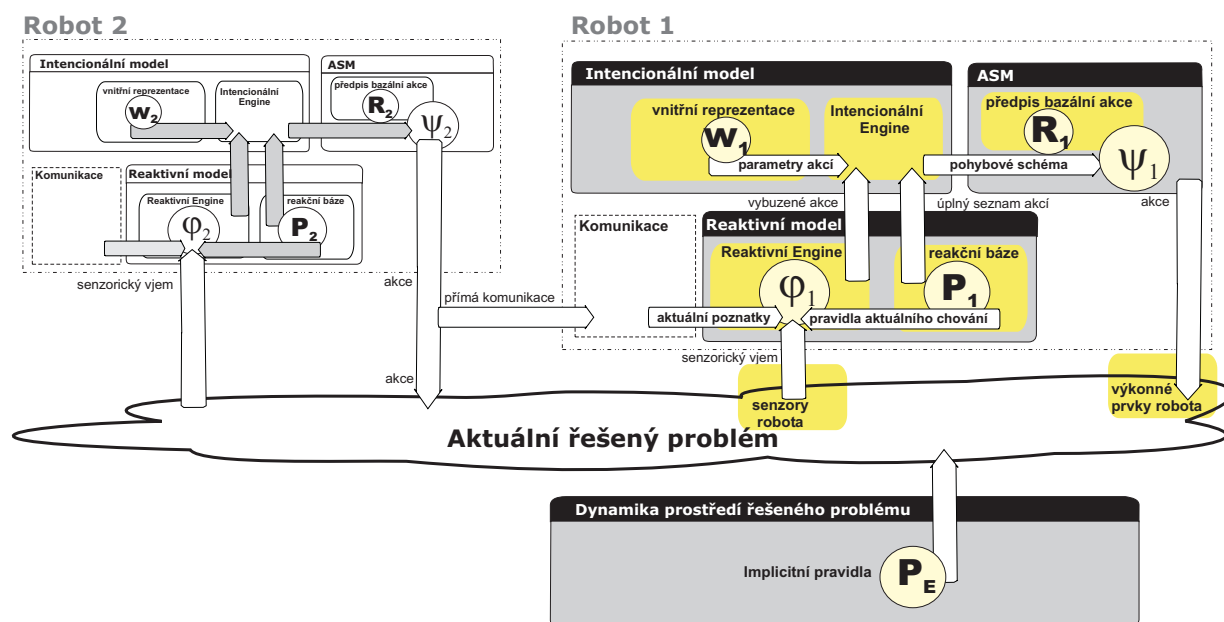
Všechny způsoby interakce kombinuje v jeden celek: eko-gramatický systém – (EG).

Definice 3.4.1 *Eko-gramatický systém*

Mějme modely agentů $A_i = (V_i, P_i, R_i, \varphi_i, \psi_i)$ shrnující
 jejich konečnou abecedu V_i
 konečnou množinu evolučních pravidel (P1L) $P_i = \{x \mapsto y : x \in V_i^+; y \in V_i^*\}$
 množinu jejich akčních pravidel R_i
 evoluční zobrazení $\varphi_i : V_E^* \rightarrow 2^{P_i}$
 akční zobrazení $\psi_i : V_i^+ \rightarrow 2^{R_i}$
 a model prostředí $E = (V_E, P_E)$ slučující
 konečnou abecedu prostředí V_E
 evoluční pravidla (ve tvaru P1L-systémů) $P_E = \{x \mapsto y : x \in V_E^+; y \in V_E^*\}$
Eko-gramatický systém (stupně n) pak uspořádaná $(n + 1)$ -tice
 $\Sigma = (E, A_1 \dots A_n)$

□

Na této definici eko-gramatického systému 3.4.1 dle [Csuhaj-Varjú a kol., 1997] se dobře vymezují hlavně operace nad intencionálním modelem autonomních kontextově orientovaných elementů (agentů). Novým prvkem tohoto formalismu je paměť jednotlivých agentů. Použití paměti je důležité obecně v každém hybridním řešení. Já zde paměť využívám na uložení krátkodobé vnitřní reprezentace problému formou intencí. V číselné škále hodnotím zejména váhy bazálních akcí.



Obrázek 3.4: Použití eko-gramatických opeátorů v řídicím systému. V pravé části se zaměřuji na interakce Robotu 1 s prostředím. Reaktivní Engine už pracuje se symboly. Transformace jsou shrnuty do jedné abstrakce evolučního zobrazení φ_1 . Z reaktivního modelu se paralelním přepisem se do paměti intencionálního modelu dostávají symboly použitelné za daného chování. Mechanismus výběru akce pracuje právě nad vnitřní reprezentací intencionálního modelu w_1 . Jeho počáteční hodnocení bazálních akcí slouží i jako axiom při výběru akce. Z nich se odvozuje sekvence pro přepis řešeného problému. Programové sekvence pro stavový přechod jsou uvedené v akčních pravidlech R_1 . Nad těmito pravidly pracuje ASM, jeho funkci formálně vystihuje akční zobrazení ψ_1 .

V tomto pojetí je model prostředí oddělen od řídicího systému robota. Implicitně předpokládám vývoj prostředí podle jeho vlastních pravidel P_E . Navrhovaná architektura na vývoj prostředí reaguje stejně jako na aktivní změny ve stavu řešeného problému.

Ve stejném duchu je princip implicitní interakce použit při paralelním řešení problému dvěma roboty. Robot 2 v eko-gramatickém systému kromě implicitní interakce může i zasahovat do vnitřní reprezentace přímo symboly společné abecedy. Symbolická reprezentace se převádí v dalším komunikačním modulu – k transformaci slouží například dále uváděný sociální model.

Sada základních (bazálních) chování R pak představuje determinismus pravidlového systému. Při použití spojitě hodnotící funkce bude mít zobrazení ψ za úkol nalézt extrém hodnotící funkce.

3.4.5 Senzuální a akční vazby

A. Akční, percepční a pohybová schémata

Analogicky k fuzzy přístupu lze Akc_i odvozovat od extrému pohybového schématu $MotSch$. Váha zvolené Akc_{opt} je v něm posílena příspěvkem aktuálních významově blízkých závěrů.

Oproti fuzzy logice lze odvozenou akci ještě parametrizovat — specifické parametry zvolené $\mathbf{Akce}_{opt}$ dostaneme inverzním zobrazením

$$\begin{aligned}
SENSSCH_i(x, y) &: \mathbf{Vjem}_i \rightarrow R^2 \\
ACTSCH_j(x, y) &: \mathbf{Akce}_j \rightarrow R^2 \\
MOTSCH(x, y) &: \sum_i SENSCH_i(x, y) + \sum_j ACTSCH_j(x, y) \\
\mathbf{Akce}_{opt} &= ACTSCH_j^{-1}(x_{opt}, y_{opt}) : \quad \begin{aligned} MOTSCH(x_{opt}, y_{opt}) &= \max_{x, y} \\ ACTSCH_j(x_{opt}, y_{opt}) &= \max_j \end{aligned} \quad (3.4)
\end{aligned}$$

V rovnici (3.4) jsem označil zobrazení dílčího perceptuálního schématu $SENSSCH$, zobrazení dílčího akčního schématu $ACTSCH$, výsledné pohybové schéma $MOTSCH$, a $ACTSCH_j^{-1}$ inverzní zobrazení pro určení parametrů zvolené akce. Aktuální vjemy tedy budí jednotný prostor pohybového schématu (R^2) paralelně.

B. Integrace předpokladů blízkého významu v pohybovém schématu

Můj návrh je postaven na hrubém ošetření senzorických vstupů v reaktivním modelu. Inkrementální model každý jeho kusý závěr reprezentuje v samostatném perceptuálním schématu. Princip perceptuálního schématu (3.4) implementuji zcela účelově. Hledám kompromis, a pro něj musí být definičním oborem schémat celá vnitřní reprezentaci stavů z okolí agenta. V něm přímou reprezentaci detekovaného stavu okolí hodnotím ještě plnou vahou W , tak jak ji udává reaktivní model. A protože fakticky se detekovaný stav v prostředí neizoluje, i ve vnitřní reprezentaci předpokládám normální rozdělení jeho váhy. V tomto smyslu jedno perceptuální schéma vystihuje hodnotící plocha

$$SENSSCH_i(x, y) = \frac{W_{(i)}}{2\pi\sigma_{1(i)}\sigma_{2(i)}\sqrt{1-\varrho_{(i)}^2}} e^{\frac{1}{2}Q(x,y;\sigma_{1(i)},\sigma_{2(i)},\varrho_{(i)},S_{1(i)},S_{2(i)})} \quad (3.5)$$

V normálním rozdělení se už objevují statistické parametry, vhodné k adaptivní modifikaci vnitřní reprezentace. Takovými jsou zejména běžně užívané *rozptyly podle báze* x (σ_1), respektive σ_2 v y bázi. Standardně označuji též ϱ jako *korelaci mezi pravděpodobnostmi pohybu v těchto bázích*.

V dále rozebíraných příkladech navigace robotů v aréně v normálním rozdělení uvažuji dva volné parametry, pohybové báze x a y . Toto rozdělení modifikuji pěti vázanými parametry. Kromě zmiňovaných rozptylů σ_1 , σ_2 a korelace ϱ se jedná o souřadnice namapovaného senzoru $\vec{S} = [S_1, S_2]$. Volné a vázané parametry kvadriky Q tak v exponenciále rovnice (3.5) vymezují vliv vybuzeného senzoru na další reprezentované stavy.

$$Q(x, y; \sigma_1, \sigma_2, \varrho, S_1, S_2) = \frac{1}{1-\varrho^2} \left[\frac{(x-S_1)^2}{\sigma_1^2} + \frac{(y-S_2)^2}{\sigma_2^2} - 2\frac{(x-S_1)(y-S_2)}{\sigma_1^2\sigma_2^2} \right] \quad (3.6)$$

V praktickém návrhu rovnice (3.5) a (3.6) determinují váhu při hledání kompromisu mezi významově podobnými závěry právě vázané parametry. Pro jednotou manipulaci s nimi je obdobně jako v praktické realizaci definuji v souhrnné struktuře *zobecněné charakteristika*

podnětu $N(\vec{S}, W, \sigma_1, \sigma_2, \varrho)$. Pro tuto práci tak uvažuji váhu podnětu $\mathbb{S}_{\text{CH}}$ v normálním rozdělení, zcela ve významu rovnic (3.5) a (3.6). Přičemž kalkulovat mohu jak s podněty sensorické povahy, tak i s autostimulacemi. Se zobecněnou charakteristikou se tak setkáme v optimalizaci algoritmů:

- Při *autostimulaci*, kde se uplatní váha zatlumených akcí.
- Při *zpřesnění reakce*, kdy se v pohybovém schématu hledá extrém (ve vnitřní reprezentaci je prohledávaná oblast vymezena nejvhodnějšími kandidáty).

3.4.6 Intencionální ASM

*Mechanismy intencionálního modelu dodávají pro mechanismus výběru akce **pohybová schémata**. Nad váženým nástrojem pohybových schémat pracuje intencionální ASM — tentokrát má kromě výběru nejvhodnější reakce ještě upřesnění, za jakých podmínek by se zvolená bazální akce nejlépe uplatnila a jak. **Intencionální ASM** má oproti reaktivnímu ASM rozšířenou funkci:*

- stejně jako v případě reaktivního výběru se volí nejvýše ohodnocená akce
- navíc dochází k většímu zaostření této zvolené akce. Je parametrizována do stavu, kde je vlastní hodnocení předpokladů ke spuštění takové akce nejvyšší.

A. Vnitřní zpřesnění nejvhodnější reakce

Pokud šlo o obecnější hodnocení rozhodných podnětů, volba pro normálové rozložení hodnotící funkce $\mathbb{S}_{\text{CH}}$ byla ryze účelová. Funkce (3.5) má algoritmicky dobře vyčíslitelné derivace až do druhého řádu. Hodnota první derivace je přímo úměrná hodnotě funkce, a pro celé pohybové schéma platí

$$\begin{aligned} \frac{\partial}{\partial x} \text{MOTSCH}(x, y) &= \frac{1}{2\pi} \sum_{(i)} \frac{W_{(i)}}{\sigma_{1(i)}(1 - \varrho_{(i)}^2)} \left(\mathbb{S}_{\text{CH}(i)}(x, y) \left(\frac{x}{\sigma_{1(i)}} + \frac{y}{\sigma_{2(i)}} \right) + \frac{\overline{S}_{1(i)}}{\sigma_{1(i)}} + \frac{\overline{S}_{2(i)}}{\sigma_{2(i)}} \right) \\ \frac{\partial}{\partial y} \text{MOTSCH}(x, y) &= \frac{1}{2\pi} \sum_{(i)} \frac{W_{(i)}}{\sigma_{2(i)}(1 - \varrho_{(i)}^2)} \left(\mathbb{S}_{\text{CH}(i)}(x, y) \left(\frac{x}{\sigma_{1(i)}} + \frac{y}{\sigma_{2(i)}} \right) + \frac{\overline{S}_{1(i)}}{\sigma_{1(i)}} + \frac{\overline{S}_{2(i)}}{\sigma_{2(i)}} \right) \end{aligned} \quad (3.7)$$

a obdobně lze lineárně odvodit z první derivace i hodnotu derivace druhé:

$$\begin{aligned} \frac{\partial^2}{\partial x^2} \text{MOTSCH}(x, y) &= \sum_{(i)} \left(\frac{W_{(i)} \mathbb{S}_{\text{CH}(i)}(x, y)}{\sigma_{1(i)}^2(1 - \varrho_{(i)}^2)} + \frac{\partial}{\partial x} \mathbb{S}_{\text{CH}}(x, y) \left(\frac{x}{\sigma_{1(i)}} + \frac{y}{\sigma_{2(i)}} \right) \right) \\ \frac{\partial^2}{\partial y^2} \text{MOTSCH}(x, y) &= \sum_{(i)} \left(\frac{W_{(i)} \mathbb{S}_{\text{CH}(i)}(x, y)}{\sigma_{2(i)}^2(1 - \varrho_{(i)}^2)} + \frac{\partial}{\partial y} \mathbb{S}_{\text{CH}}(x, y) \left(\frac{x}{\sigma_{1(i)}} + \frac{y}{\sigma_{2(i)}} \right) \right) \end{aligned} \quad (3.8)$$

Této vlastnosti využiji při iterativním hledání extrému celého pohybového schématu. Numericky se v pohybovém schématu přibližují k hledanému stavu s variabilním krokem ve směru gradientu pohybového schématu. Iterace končí v místech, kde funkce $\mathbb{S}_{\text{CH}}$ nabývá

nulového diferenciálu. Gradientní krok prakticky vychází z Newtonovy metody hledání extrému.

$$\begin{aligned} s_1(n+1) &= s_1(n) - \frac{\frac{\partial}{\partial x} MOTSCH(s_1(n), s_2(n))}{\frac{\partial^2}{\partial x^2} MOTSCH(s_1(n), s_2(n))} \\ s_2(n+1) &= s_2(n) - \frac{\frac{\partial}{\partial y} MOTSCH(s_1(n), s_2(n))}{\frac{\partial^2}{\partial y^2} MOTSCH(s_1(n), s_2(n))} \end{aligned} \quad (3.9)$$

Algoritmus 4 *Určení maxima pohybového schématu*

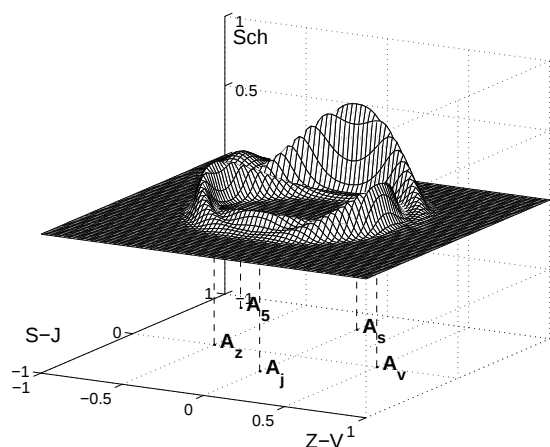
- 1: {inicializace}
 - pro každý podnět napočítat konstanty K1,K2,K3,K4 dle rovnic (3.6) a (3.7)
 - začíná se od nejlepšího podnětu (S1,S2,W); $W=\max \Rightarrow \text{Best}=[S1, S2]$
 - 2: **iterace** {gradientní hledání extrému}
 - 3: inicializace sum: Diff1=[0;0], Diff2=[0;0]
 - 4: kalkulace Diff1.x,Diff2.x přes všechny podněty $dSCH=K1*SCH(\text{Best})+K2$,
Diff1.x=Diff1.x+dSCH, Diff2.y=Diff2.y+K3*dSCH+K4
 - 5: korekce nalezeného stavu: $\text{Best.x} = \text{Best.x} + \text{Diff1.x}/\text{Diff2.x}$
 - 6: body 4 a 5 analogicky pro Diff1.y,Diff2.y,Best.y
 - 7: **dokud** $\|\text{Diff1}\| < \text{tolerance}$
 - 8: nalezený stav: **Best**
-

Algoritmus 4 slouží především k určení maxima pohybového schématu. Algoritmus hledá optimum $s = [s_1, s_2]$ odděleně v obou bázích. Po interpretaci rovnice (3.9) v v bázi x už může být nové s_1 použito při hledání s_2 podle báze y . Vzhledem k tvaru derivací (3.7) a (3.8) roste složitost algoritmu iterativního hledání lineárně.

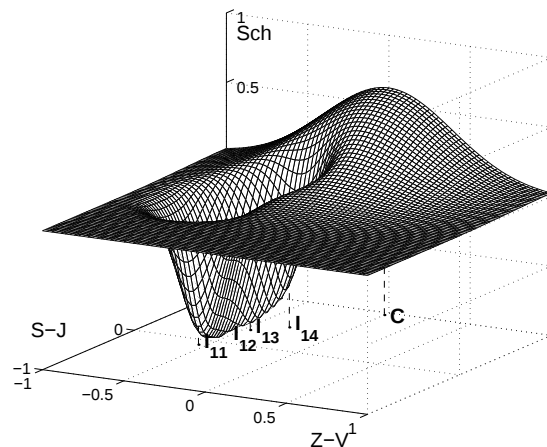
B. Význam pohybového schématu

Na parametrické reprezentaci pohybového schématu jsem ukázal výpočetně jednoduchou superpozici příspěvků do jednoho kompromisního řešení. V jednotném pohybovém schématu se do optimálního stavu započte vjem vnějšího prostředí spolu s vnitřními podněty. Najít takové optimum je v pohybovém schématu polynomiálně náročné. Počítá se pouze příspěvek každého podnětu se do gradientu pohybového schématu. Další příklad ukazuje, jak lze skládat vnitřní hodnocení souslednosti s podnětem k jízdě směrem světelného vjemu.

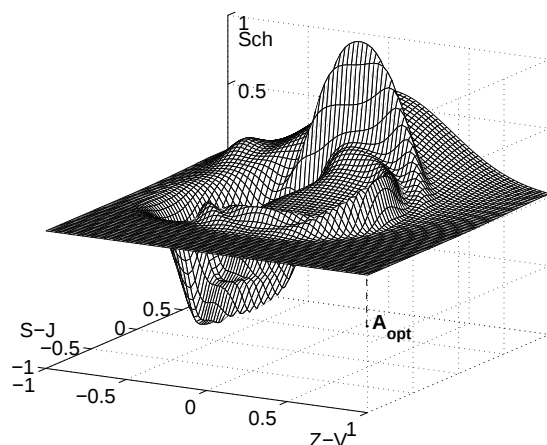
Na sérii obrázků 3.5 se ukazuje, s jakou vahou pohybové schéma zahrnuje jednotlivé podněty do výsledného rozhodnutí. Směr zaneseného podnětu odpovídá směru vedené bazální akce či vjemu vůči čelnímu směru robota. Všechna schémata jsou pak bezrozměrová, dosah vnitřní reprezentace se omezuje pouze na lokální rozhodnutí. Stav v každé ze souřadnic nabývají hodnot $\langle -1; 1 \rangle$. V těchto souřadnicích akce mapují od středu pohybového schématu $\vec{S}_0 = [0; 0]$ ve vzdálenosti, odpovídající délce kroku ℓ . U nepohybových akcí odpovídá délce sekvence, nezbytné k dosažení podobného efektu alternativními akcemi.



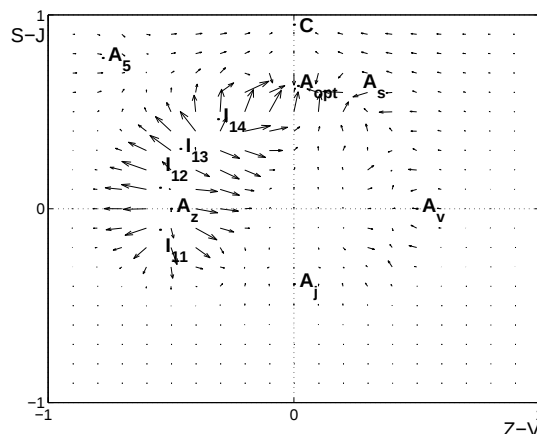
(a) akční schémata



(b) perceptuální schémata



(c) výsledné pohybové schéma



(d) gradient v pohybovém schématu

Obrázek 3.5: Příklad užití pohybových schémat. Obrázek (a) ukazuje mapování použitých bazálních akcí do akčních schémat. Poloha stavů odpovídajících pohybu ve směrech S , J , V , Z je zřejmá z průmětu jader A_S , A_J , A_V , A_Z . V SZ směru robot očekává přítomnost dalšího robota. Ve vnitřní reprezentaci je záměr komunikovat s druhým robotem ohodnocen v akčním schématu s jádrem A_5 . Podle těchto akčních schémat robot v klidu je preferuje bazální pohyb v čelním směru (S). Vnitřní podněty pro otáčení robota jsou s nižší prioritou rovnocenné v obou bočních směrech Z, V .

Perceptuální schéma z obrázku (b) lze vyhodnotit obdobným způsobem. V příkladu uvažujeme bezprostřední detekci infračidly $I_{11} \div I_{14}$ a rostoucí intenzitu světla (jádro tohoto vjemu označují C). Zobrazení podnětů do stavů v perceptuálním schématu respektuje polohu jejich detekce, hodnota přímo vystihuje významu podnětu pro zavedené akce. Aktivní světelný podnět je hodnocen pozitivně. Naopak podněty z infračervených čidel signalizují překážku, jsou proto hodnoceny negativně.

Na obrázku (c) je znázorněna superpozice akčních a perceptuálních schémat s vyznačeným průmětem extrému pohybového schématu A_{opt} . Vzhledem k rozložení stavů je stavů optimální stav blízký jádru akce $krok_+(A_S)$.

Obrázek (d) pak vykresluje tutéž funkci pohybového schématu v gradientech, v souvislosti se všemi dosud reprezentovanými stavy.

V akčním schématu na obrázku 3.5(b) se blíže středu $\vec{S}_0$ objevují stavy odpovídající bazálním pohybovým akcím. Naopak stavy dále od středu $\vec{S}_0$ odpovídají například bazálním komunikačním aktivitám. Stejným způsobem se promítají významy aktuálních vjemů do perceptuálního schématu na obrázku 3.5(a). V něm je hlavní atraktor gradientního navádění odvozen od čelní detekce světla. Na druhou stranu detekce objektů infračervenými čidly představují repulsivní podněty pro pohyb v daném směru.

Z pohybového schématu odvozují jednu optimální akci po superpozici obou schémat dle Algoritmu 1. Mechanismus výběru akce z obrázku 3.5(c) upřednostňuje jízdu v čelním směru $[\ell; \frac{1}{20}\ell]$. Mírné vychýlení pohybu od přímé jízdy ve směru rostoucí intenzity světla způsobila boční detekce překážky.

Shrnu-li:

V této podkapitole 3.4 jsem navrhl a prezentoval *intencionální část hybridní architektury*. Přitom jsem použil autostimulace. Ta v mém pojetí představuje *jeden z popisů stavového přechodu ve vnitřní reprezentaci*. Autostimulací prakticky implementuji pouze příčinnou vazbu, stejně jako Maes.

V další části jsem implementoval důležitou *funkci zapomínání* a uvedl *iterativní postup*, jakým se odvozuje a klasifikuje hodnota (význam) zprostředkované informace. V ovážení takových podnětů se odráží, že ASM se na senzory neověřená data spoléhá pouze krátkodobě. Poté je reakce ASM na takový podnět anulována.

V další podkapitole jsem ukázal, že navrhované pojetí je v souladu s pojetím eko-gramatického systému. *Zvláštností mého návrhu oproti přístupu Prof. Kelemenové je to, že jednotky i nejnižší úrovně jsou kontextové*. Eko-gramatický formalismus mně zde otevírá možnosti i pro *zavedení kriteriálních funkcí pro průběžné hodnocení kvality řídicího systému*.

Věnoval jsem se zde i možnosti *obecnější parametrizace bazálních pohybových akcí v pohybovém schématu*. Nejprve jsem zavedl *formalismus pohybového schématu*. Poté jsem se zaměřil na speciální případ hodnocení významu jednotlivých podnětů v pohybovém schématu v normálním rozložení a ukázal výhody právě použitého normálního rozložení významu podnětu.

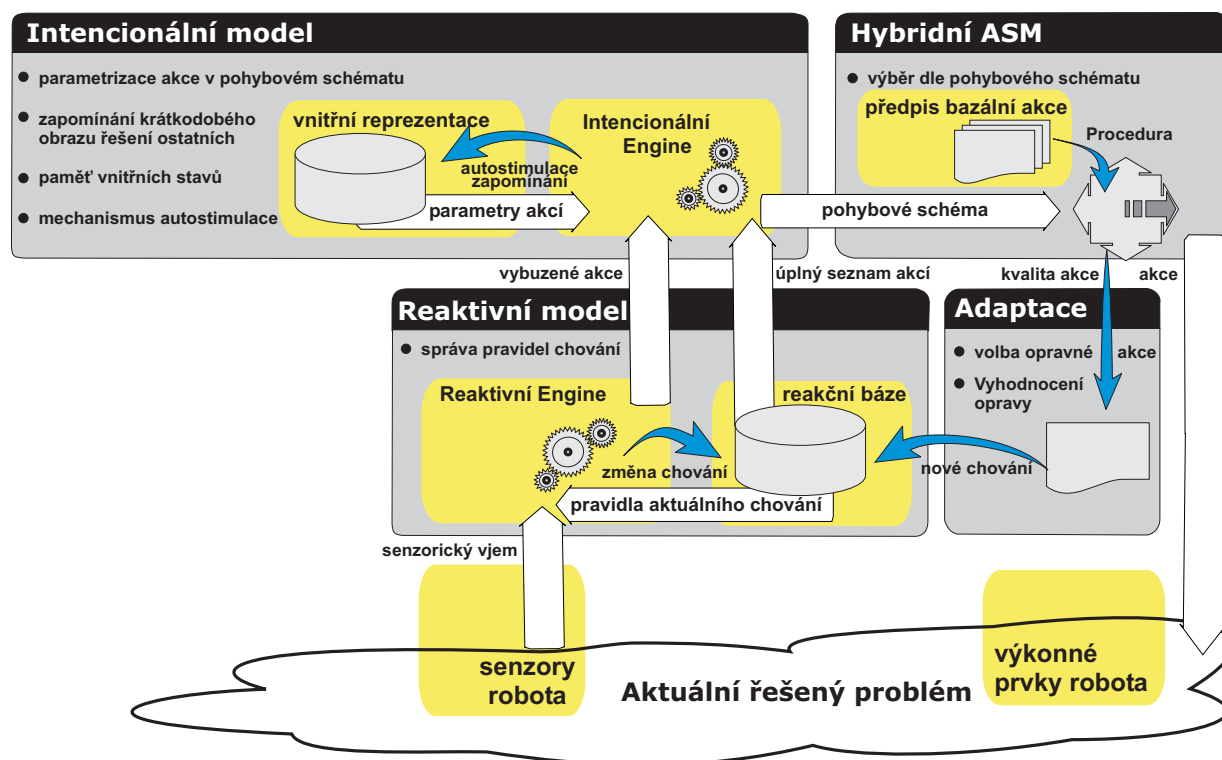
V poslední části podkapitoly 3.4 jsem všechny výše zavedené prvky pospojoval (integral) do jednotného intencionálního mechanismu výběru akce a předvedl, jak se zvolená akce odvozuje z kombinace průběžně hodnocených bazálních akcí a aktuálních vjemů. V nich se zohledňuje rozpracovanost řešení v okolí robota.

Pro realizaci hybridní architektury jsou dále **důležité vztahy**, zavedené v podkapitole 3.4:

- pro gradientní hledání pracovního bodu v rámci ASM (3.9)
- model autostimulace jako jednoho ze stavových přechodů (3.2)
- užití exponenciálního zapomínání jako dalšího filtru impulsů registrovaných v paměti intencionálního modelu (3.3)

3.5 Adaptivní část hybridní architektury ASM

3.5.1 Blokové schéma hybridní architektury s adaptací



Obrázek 3.6: Hybridní architektura řízení s adaptací. Adaptace je zapojena jako paralelní modul k předchozí kombinaci reaktivního a intencionálního modelu v hybridním ASM. Modul adaptace pracuje ve stavu, kdy robot reaguje podle pravidel předdefinovaného chování pokusný-krok. Kontroluje zejména jeho ověření či zamítnutí. Pokud jsou pravidla pro pokusný-krok potvrzena, modul Adaptace v reakční bázi definuje odpovídající nové chování. V opačném případě, při detekci chybného postupu, modul Adaptace zanáší do intencionálního modelu pokutu pro použití chybného směru.

Systém je navržen tak, aby bylo možné jednotlivé případy v evolučním zobrazení φ nastavovat individuálně. V první etapě návrhu hybridního systému jsem mohl nastavit vnitřní reprezentaci podnětu ve stavech $S = [S_1, S_2]$ s vahou W velice zjednodušeně: simulační server RoboCup poskytoval informaci vůči pevným prostorovým značkám. Jejich význam pak byl v určení geometrické odchylky vůči pracovní oblasti. Taková logika nastavení evolučního zobrazení je ovšem nepřenositelná na řízení reálných robotů s těmi nejjednoduššími senzory.

3.5.2 Rozbor kritérií kvality navržené akce ASM

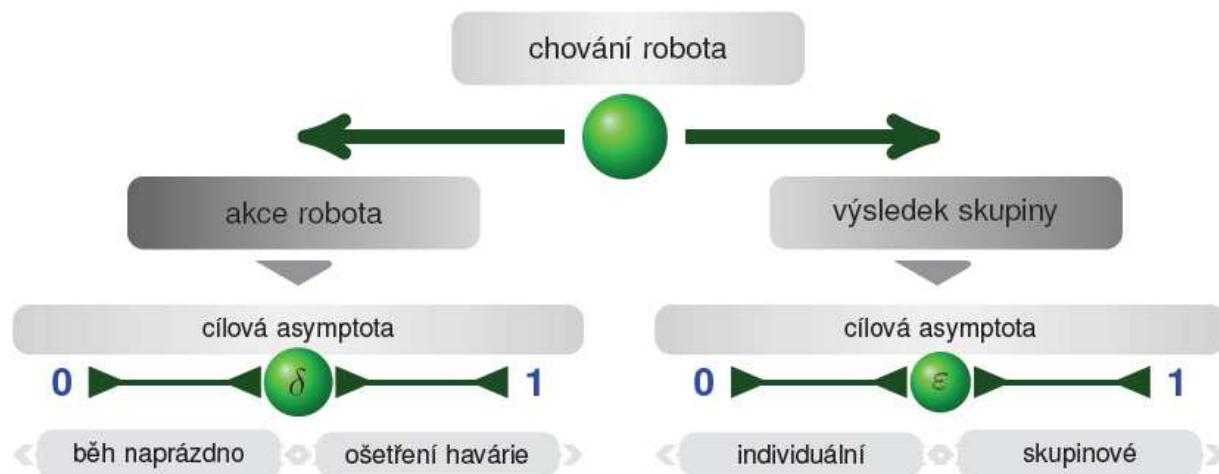
Pro vlastní nastavení se nejprve zaměřím na vlastnosti kritérií, podle nichž bude možné určit, zda aplikace bazálních akcí byla úspěšná či naopak upozornit na nevhodné použití akce. Nevhodná volba akce se může projevit i selháním robota. Pohled na kvalitu bazálních akcí a na kvalitu chování robota je vzájemně provázaný. Jejich současné hodnocení navíc napomáhá odhalit nedostatky, které mohou být v jednom z úhlů pohledu zastřené.

Nastavené měřítko při řešení úlohy přístupem Umělého života významně určuje i chování Kolektivu v cílovém stavu. V naší hybridní architektuře se mechanismus výběru akce ASM opírá o vyhodnocení rozpracovaného řešení ze dvou měřítek — kritérií kvality

- Diferenciální kritérium δ . Pomocí diferenciálního kritéria δ hodnotím, jak bazální akce nízké úrovně dál sleduje trend, určený na nejnižší úrovni robota. Po každém kroku ASM hodnota δ vyjadřuje, nakolik se sám robot svým novým stavem odchýlil od původně vytyčeného směru řešení.
- Integrální kritérium ϵ . Kvalifikuje celkový efekt chování Kolektivu (robotů) jako celku. Dává do poměru dosaženou frekvenci požadovaných vlastností oproti četnosti ostatních, nevyžádaných stavů.

Měřítka kvalit δ a ϵ jsou vhodná i pro práci s reaktivním ASM. V něm diferenciální kritérium δ předurčuje přímo zvolenou aktivitu. Hodnoty integrálního kritéria ϵ reaktivní ASM zpracovává implicitně (využívá je jako omezující podmínky v pravidlech chování).

Kritéria jsou navržena z logiky interpretace rozdílů v reaktivním chování robota coby jednotlivce (využíváme δ) oproti chování Kolektivu jako celku (využíváme ϵ). Schématicky ideu návrhu dobře vystihuje následující obrázek 3.7. Kritéria by měla vymezit práci mechanismu výběru akce v mezích «klidového a kritického», stejně jako «skupinového a individuálního». V těchto mezích budu v dalších částech práce poměřovat i různé mechanismy ASM. Vůči dosahu ASM bude kritérium poměřované a normované. Dále v práci uvažuji rozsah hodnot kritérií od minima (0) po maximum (1). V této části návrhu se nicméně



Obrázek 3.7: Idea použití kritérií výběru akce. Kritéria v práci podchycují hlavní úrovně realizace řešení přístupem Umělého života. Diferenciálním kritériem δ hodnotím zapojení bazálních akcí pro uspokojení cíle agenta. Na druhou stranu integrálním kritériem ϵ hodnotím zapojení agentů pro dosažení cíle kolektivu. Obě kritéria jsou postavena jako poměřová, normovaná do intervalu hodnot $\{0; 1\}$.

zaměřím na vypovídací schopnost zavedených kritérií o krizových situacích. Uvedu použití kritérií pro detekci takových situací a jejich ošetření. Omezení (upřesnění) zde nastíněných kritérií pak dále rozvedu v kapitole 4 ke každému z experimentů.

A. Metrika při distribuci zájmových oblastí

Úloha je shora omezena konečným rozsahem problémové situace S , nikoliv však přesným stavem. Proto se v distribuci zájmových oblastí zaměřuji pouze na přiřazení mezi apriorní konečnou množinou robotů $AGS = \{R_i\}$ a částmi problémové situace. Rozsah pokrytý robotem závisí na jeho aktivitě. Robot R_i k řešení volí vždy jednu akci ze své sady $\{\mathbf{akce}_{ij}\}$. V chování systému se v úloze distribuce zaměřuji na možnou aktivitu robota ze dvou pohledů:

Integrální kritérium ϵ : hodnotí, jak řešení v Kolektivu pokrývá zájmovou oblast. Distribuční úlohu omezuje plocha oblasti S . Cílem integrálního kritéria ϵ je minimalizovat plochu, v které se příspěvky jednotlivých robotů $S(R_i)$ překrývají.

$$\epsilon_{distr} = \bigcap_{AGS} S(R_i) \rightarrow \min \quad \text{pro} \quad \bigcup_{AGS} S(R_i) \rightarrow S; \quad (3.10)$$

Diferenciální kritérium δ : v úloze distribuce minimalizuje náklady robota na vyhodnocení pokryté oblasti. Minimalizace nákladů se obecně charakterizuje jako udržování agenta ve stavu minimální entropie, či alespoň snahou o udržení preferované akce. Při jejich známé preferenci $c_{distr}(\mathbf{akce}_{ij})$ minimalizují

$$\delta_{distr} = \sum_j c_{distr}(\mathbf{akce}_{ij}) \log c_{distr}(\mathbf{akce}_{ij}) \rightarrow \min \quad (3.11)$$

Preference pohybové akce $c_{distr}(\mathbf{krok}_x)$ je dobré posuzovat ve zpětné vazbě — zda se zvolenou akcí dosáhneme očekávaného výsledku. Při praktickém řízení pohybu lze porovnávat počet požadovaných pulzů servopohonu oproti skutečnému posunu, určenému z odometrie, (rozdíl mezi teoretickou hodnotou a skutečností atd. Kde zpětná vazba neposkytuje dostatečně detailní popis problému, nezbyvá než určit kritéria paušálně podle aktivity robota.

B. Metrika při koordinovaném pohybu

Úloha je omezena počátečním stavem S a cílovým stavem T . Koordinace je vázána relativně vůči ostatním indikátorům polohy. Především koordinuji vzájemně roboty $AGS = \{R_i\}$ tak, aby každý z nich byl ve svém relativním rovnovážném stavu užil $\mathbf{akce}_{ij}$, s cenou $c_{coord}(\mathbf{akce}_{ij})$. Podle kolektivního omezení se robot se zvolenou akcí přesouvá od počátečního stavu S do cílového T . Ve formulaci kritérií proto zdůrazňuji následující omezující předpoklady:

Integrální kritérium ϵ : v Kolektivu se volí nejlevnější aktivity pro přesun po cestě z počátku S do cíle T přes všechny roboty zapojené do úlohy vzájemné koordinace AGS

$$\epsilon_{coord} : \sum_{AGS} \int_S^T c_{coord}(\mathbf{akce}_{ij})(l) dl = \min \quad (3.12)$$

Diferenciální kritérium δ : se vztahuje na robota R . Jeho rovnováha se kvalifikuje ve váženém těžišti od stavů generovaných roboty $R_i \in AGS$. S ostatními agenty srovnává svoji polohu, sám u sebe pak srovnává svoji skutečnou polohu R s očekávaným výsledkem $\hat{R}$.

$$\delta_{coord} : \sum_{AGS} \omega_{R_i} \cdot ||R_i - R|| = \min \quad (3.13)$$

Váhou ω_{R_i} zavádím omezení mezi robotem R_i a R : při nulové hodnotě roboti spolu nekoordinují. S větší vahou ω_{R_i} se robot R v koordinaci více zaměřuje právě na robota R_i .

V navrhovaném řešení stále mohu cíl T zadat výčtem ohraničujících podmínek. Robot ovšem musí být z dostupných ohraničujících podmínek stále schopen určit svůj odhad $\hat{R}$ a míru rozdílu mezi odhadem a skutečným stavem R .

3.5.3 Zpřesnění reakce adaptací

Abych zachoval co nejjednodušší vnitřní reprezentaci řešeného problému i v reálném prostředí, přistoupil jsem ve druhé fázi k adaptivnímu určení základních parametrů na základě bezprostřední zkušenosti robota.

Adaptaci navrhuji bez učitele, jen na základě diferenciálního kritéria δ . Rozhodným okamžikem v adaptaci je negativní zkušenost, vyvolaná rapidním poklesem $\dot{\delta} = \delta(t) - \delta(t-1)$, určeným rozdílem diferenciálního kritéria δ ve dvou po sobě následujících okamžicích $t-1$ a t . Pozitivní zkušeností rozumím chování (akceptované robotem i prostředím), blízké předchozímu: $\dot{\delta} \rightarrow 0$.

Cílem adaptace je minimalizovat krátkodobou ztrátu. Pro minimalizaci zahrnuji do rozhodnutí podněty (zkušenosti), plynoucí z aktuálních kritérií.

- Po negativní zkušenosti do pohybového schématu přechodně zapojuji podnět s repulsivním ohodnocením, $W = -1$. Chyba plyne z posledního rozhodnutí za stavu $\text{LastBest} = [\text{LastBest.x}, \text{LastBest.y}]$.
- Po pozitivní zkušenosti trvale dodávám nový atraktivní podnět do stavu, v němž se chyba napravila. Opět předpokládáme, že změna plyne z posledního stavu LastBest . Váhu podnětu W interpoluji z hodnot diferenciálního kritéria před provedením pokusného kroku δ_{Last} , a po jeho dokončení δ_{Actual} : $W = \frac{\delta_{\text{Actual}} + \delta_{\text{Last}}}{2}$.

Repulsivní či atraktivní podnět je reprezentován v rozsahu stavů, zahrnutých do vyhodnocení adaptačního pokusu. Odvozený fakt platí v ℓ -okolí stavu LastBest , přičemž ℓ je normalizována délka jednoho pokusného kroku. V tomto okolí rovnoměrně (Gaussovsky) zobecňuji váhu podnětu v obou bázích ($\sigma_1 = \sigma_2 = \ell$). Není důvod tedy uvažovat kovarianci mezi bázovými složkami, $\varrho = 0$. A to jak v repulsivním ohodnocení chybového směru ERRSCH , tak atraktivním ohodnocení nového, pokusným krokem ověřeného postupu NEWSCH .

$$\begin{aligned} \text{ERRSCH}(x, y) &= -\frac{1}{2\pi\ell^2} e^{\frac{1}{2}Q(x,y;\ell,1,\text{LastBest.x},\text{LastBest.y})} \\ \text{NEWSCH}(x, y) &= \frac{\delta_{\text{Actual}} + \delta_{\text{Last}}}{4\pi\ell^2} e^{\frac{1}{2}Q(x,y;\ell,1,\text{ActualBest.x},\text{ActualBest.y})} \end{aligned} \quad (3.14)$$

Interpretace rovnice 3.14 je více intuitivní z naznačeného Algoritmu 5. V jeho několika bodech je i jádro mého návrhu, jak dosáhnout adaptivního zpřesnění reakce.

Adaptace nové akce **TrialAction** na podmínky **TrialSense** probíhá v závislosti na tom, jestli už nějaká přímá vazba na dané podmínky existuje:

- v případě **nové vazby** se registruje přímo **TrialAction**.

Algoritmus 5 *Adaptace zkušební akcí***Předpoklady:** LastBest, LastDelta, LastSense; určím ActualDelta

- 1: **pokud** LastDelta-ActualDelta > tolerance, {proved adaptaci *kroky A*}
- 2: A.1. - krok zpět
- 3: A.2. - rozšíření repulsivního ohodnocení ErrSch - nový podnět v LastBest
- 4: A.3. ActualBest = calculate(MotSch+ErrSch)
- 5: A.4. - pokusný krok: jízda podle ActualBest
- 6: A.5. - zapominání vah v popisu ErrSch
- 7: A.6. - vyhodnocení: pokračuj od bodu 1
- 8: **jinak** {jsou splněny podmínky pro fixaci}
- 9: adaptuj novou akci TrialAction na podmínky TrialSense

- v případě **existující vazby** dochází ke zpřesnění zobecněné charakteristiky podnětu k takové akci koeficientem q :

V dále popisovaných úlohách snižuji rozptyl poměrem $\frac{1}{\sqrt{2}}$, tedy $q=2$.

3.5.4 Význam adaptivních kroků při změně chování

Použití adaptace vysvětluji na situaci, kdy robot postupuje okolo rohu stěny podle chování sledování-zdi nebo přímá-jízda. Změna chování je v tomto příkladě dána detekcí alespoň jednoho z proximitních čidel. Při přímé-jízdě jakákoliv detekce vede na chování sledování-zdi. Na druhé straně se robot při úplné ztrátě detekce vrací od sledování-zdi zpět k přímé-jízdě. Na omezení takových nevhodné změně chování je určena zde zaváděná adaptace.

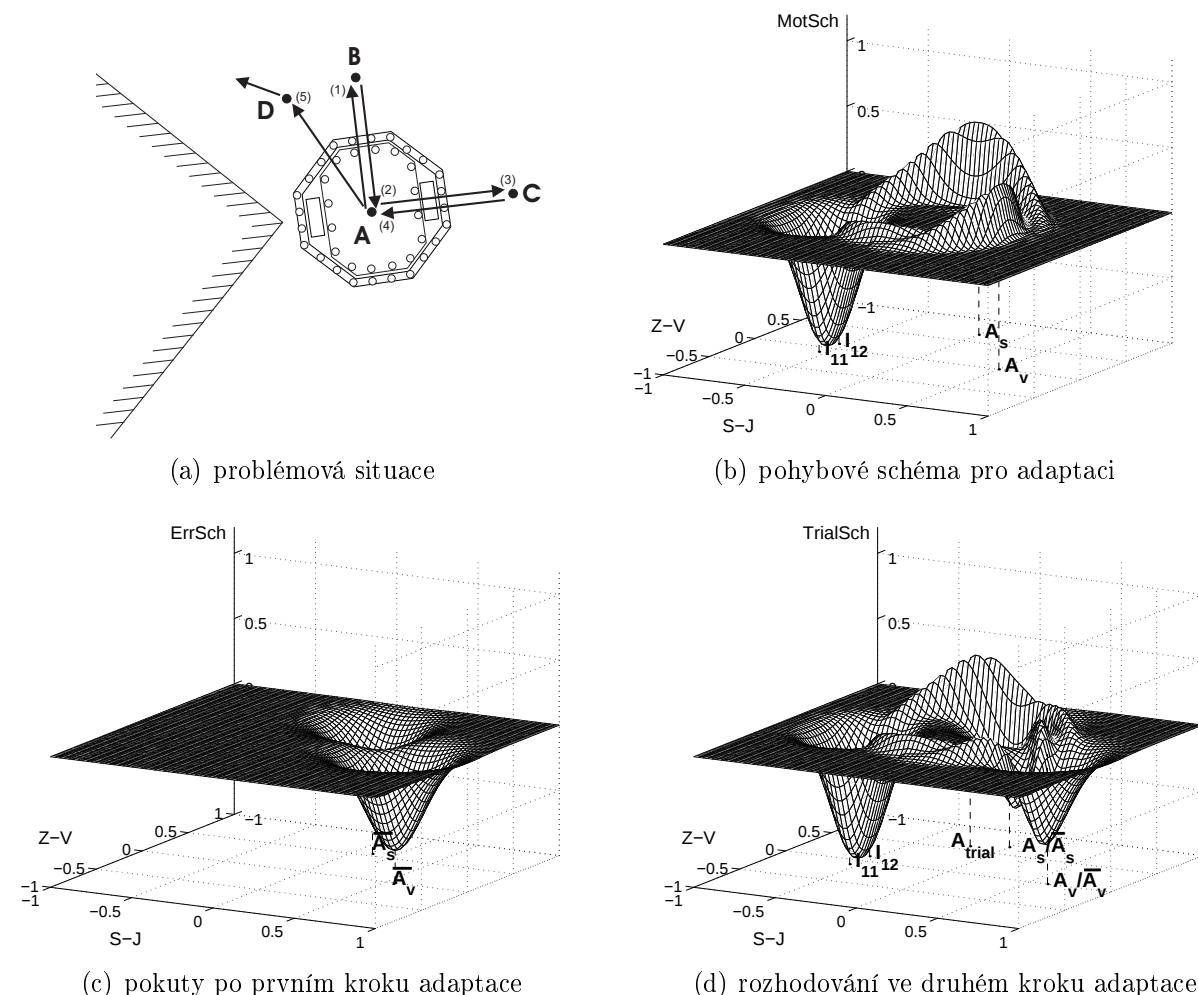
Pro tuto ukázkou jsou na obrázku 3.8(a) nevhodné změny chování situace načrtnuty jako jednotlivé úseky prvního pokusného kroku: **A** představuje výchozí místo, kde byla zvolena nevhodná reakce. V místě **B** robot vyhodnotil situaci jako nevhodný aktuální pohyb a navrácí se zpět. Obdobně (jako nevhodný krok) hodnotí robot první pokusný krok, kterým se dostal do bodu **C** a rovněž se vrací zpět do výchozího bodu **A**. V bodě **D** si však robot proximitními čidly ověřuje správnost prvního kroku a následně je zakreslen i druhý pokusný krok.

Adaptační procedury v tomto příkladu můžeme abstrahovat aktivací pravidel dvou předdefinovaných chování, návrat a pokusný-krok. Chování návrat se aktivuje při výrazné rychlosti poklesu diferenciálními kritéria ($\dot{\delta} < \tau_{adapt}$; $\tau_{adapt} < 0$). Realizace chování pokusný-krok pak probíhá podle Algoritmu 5.

Úloha k obrázku 3.8. Adaptace reakce proximitních čidel při objíždění rohu:

Nechť platí předpoklady, podle nichž se pravidla chování pokusný-krok navenek projeví ve volených bazálních akcích v místech **A**, **B**, **C**, **D**:

- A. Popis ošetření chybného rozhodnutí začínáme v místě **A**, kde byla nevhodná akce vybrána.
- B. v okamžiku detekce nevhodného chování se aktivuje deterministická posloupnost pravidel pro chování návrat. Odvozují se inverzí z předchozího postupu.



Obrázek 3.8: Adaptace reakce proximitních čidel při objíždění rohu. Pro ukázkovou situaci na obrázku (a) jsou načrtnuty jednotlivé úseky pokusného kroku: A představuje místo, kde byla zvolena nevhodná reakce. V místě B robot vyhodnotil jako nevhodný aktuální pohyb. V bodě C pak se stejně hodnotí první pokusný krok. V bodě D si nakonec robot ověřil správnost druhého pokusného kroku. Následující grafy ilustrují podklady pro vytyčení směru pokusných kroků. Pohybové schéma *MOTSCH* na obrázku (b) vystihuje vnitřní reprezentaci problému v bodě A. Přibližuje hodnocení bazálních akcí pro pohyb vpřed A_s a vzad A_j , stejně tak otáčení napravo A_v a nalevo A_z . Pozitivně jsou hodnoceny akce, negativně se do hodnocení promítají signály detekce překážek I_{11} , I_{12} . Tuto informaci poskytují proximitní čidla.

K základnímu hodnocení pohybovým schématem se po detekci nevhodného chování připočítává chybové schéma *ERRSCH*. Na obrázku (c) je například vyznačeno chybové schéma po prvním pokusném kroku, kdy je nejvyšší pokuta kladena na poslední krok – A_v . Pokuta za směr původního nevhodného chování A_s (jízda do B) se projevuje s menší vahou. Tato pokuta už prošla procedurou zapomínání.

Graf na obrázku (d) pak ukazuje souhrnné předpoklady pro odvození vhodné reakce ve druhém pokusném kroku: *MOTSCH* + *ERRSCH*. Podle tohoto vyhodnocení robot volí pohyb do bodu D, kde si následně na hodnotě diferenciálního kritéria δ potvrzuje správnost takového rozhodnutí.

- A. Po dosažení předchozího stavu je základním podkladem pro výběr akce pohybové schéma *MOTSCH*. Jeho extrémy jsou naznačeny na obrázku Obrázek 3.8(b). K tomuto pohybovému schématu se v důsledku detekované chyby připočítává negativní hodnocení pohybu v čelním směru. Na obrázcích je tento podnět ve vnitřní reprezentaci problému označován $\overline{A_s}$.
Tato negativní hodnocení mohou vést na otáčení od stěny. Při časté ztrátě detekce takové chování opět vede na přímou-jízdu do bodu C.
- C. Takováto reakce je ovšem po několika krocích vyhodnocena jako nevhodná a podruhé se aktivuje chování návrat pro storno tohoto pokusného kroku.
- A. Předpokládejme, že zde robot řeší ve stejném lokálním kontextu již jednou vyhodnocovanou situaci. Po zkušenosti z předchozího pokusného kroku se ovšem mění podmínky chybového schématu — pokutě ve směru A_s se dává v důsledku zapomínání dává menší váha. Nově se s maximální vahou zanáší do pohybového schématu pokuta spojená s poslední zkušeností, nevhodným chování ve směru $\overline{A_s}$. Oba tyto extrémy jsou patrné na obrázku 3.8(c). Toto schéma se připočítává k původnímu pohybovému schématu *MOTSCH*. Tato reprezentace je opět zřejmá z obrázku 3.8(b) Jak je dále patrné z obrázku 3.8(c), z jejich součtu se pak odvozuje jízda směrem A_{trial} .
- D. V tomto směru robot v posledních několika krocích detekoval překážku proximitními čidly. Aktivní chování je sledování-zdi. Hodnota dosaženého diferenciálního kritéria δ pak opravňuje k registraci tohoto směru jako nového chování podle Algoritmu 5.

Shrnu-li:

V této podkapitole 3.5 jsem vnitřní reprezentaci problému rozšířil a precizoval *zařazením adaptace*. Poté jsem zde navrhl a popsal párová kritéria – jakási měřítka pro posouzení kvality řešení Úloh na dvou různých úrovních detailu. Zavedená párová kritéria posuzují rozpracovanost řešení (výběru – vybuzení bazální akce). Stejnou událost mohou hodnotit současně (ale z odlišných úhlů pohledu) integrálním a diferenciálním kritériem.

Algoritmus adaptace pracuje ve třech krocích:

- (i) detekce selhání,
- (ii) volba opravného manévru včetně jeho bezprostřední realizace a nakonec
- (iii) vyhodnocení takového manévru. Dále navrhuji způsob zavedení ověřených reakcí na problematiku situace. v závěru oddílu jsem pak na příkladu ukázal, jak taková adaptace probíhá v praxi.

Pro realizaci hybridní architektury jsou pak dále důležité v této části zavedené vztahy:

- Určení krátkodobých pokut v hodnocení podnětů (3.14). Zamezuje se tím postupu ve směrech, kde došlo k selhání.
- Zavedení nového schématu (3.14), případně zpřesnění již existujícího.

3.6 Souhrn podstatných znaků navrženého hybridního řízení

Navržená hybridní architektura vychází z bazálních akcí. Mechanismus jejich syntézy *shrnuje výhody reaktivního a intencionálního modelu chování*:

- dle *reaktivního modelu* je silný podnět ošetřen bezprostředně,
- dle *intencionálního modelu* postupně dochází i k ošetření slabších alternativ v řešení

3.6.1 Význam analytické části mechanismu výběru akcí ASM

V navrženém ASM je rozhodující vjem. Jednoznačná akce z vjemu plyne ovšem zřídka. Častěji je třeba rozhodovat mezi několika přípustnými variantami. Zde o celé sadě variant uvažují jako o vybuzených akcích. Právě volbě mezi vybuzenými akcemi v mém návrhu hybridního agenta napomáhá zapojený *intencionální model*. Poskytuje doplňkové *atributy pro rozhodnutí*. Z nich se v ASM uplatní *dodatečná váha vybuzené akce* – (a) diferenciální odchylka uvažované aktivity od individuálního cíle robota, a (b) nakolik uvažovaná alternativa dlouhodobě napomáhá dosažení cíle celého Kolektivu.

V navrhovaném intencionálním modelu jsou vybuzené akce doplňovány vahou variant dosud nepoužitelných, ale žádaných. Váhový příspěvek od nepoužitelných akcí odráží logickou souslednost jednotlivých akcí. Od podporované vybuzené akce se očekává, že přispěje ke splnění všech podmínek pro použití žádaných variant.

Čisté vnitřní formování záměrů lze také nalézt v dále uváděných experimentech. Jako příklad uvádím strategii sledování pohyblivého objektu z kapitoly 4. Dobře se na ní ukazuje doplňková, zobecňující funkce intencionálního modelu: zatímco při rozhodování o vybuzených akcích se ASM zaměřuje přímo na podněty relevantní za daného chování, autostimulace může probíhat mezi všemi dostupnými akcemi bez omezení.

3.6.2 Významová superpozice vnitřních a externích podnětů

Vnitřní reprezentace problému je volně rozšiřitelná ještě o další alternativy. V intencionálním modelu jsou alternativy opět zastoupené odpovídajícími podněty. V průběhu řešení může být odvozena (a) nová akce jako kombinace dosavadních bazálních pohybových aktivit v těsné vazbě na určitý vzor. Případně se též objevují (b) nové explicitně získané poznatky z rozhodovacího procesu ostatních robotů. Metrikou pro všechna později zaváděná *dodatečná hodnocení* jsou dříve uváděná kritéria. *Pro nové podněty vyčísľují*:

- **Diferenciální kritérium δ .** Hraje roli při zavedení nových pravidel do schématu. Předpoklady pro zavedení pravidla po takzvaných pokusných krocích byly popsány v kapitole 3.5.2.
- **Integrální kritérium ϵ .** Projevuje se při komunikaci s ostatními roboty v Kolektivu. Tak mohl robot explicitně získat již integrálně ohodnocené („předoceněné“) stavy i mimo dosah svých senzorů.

3.7 Shrnutí závěrů a důležitých faktů Třetí kapitoly

V této třetí kapitole nejprve definuji mechanismus výběru akce hybridního agenta. Ve svém návrhu používám *formální odvození bazálních akcí z vnitřních termů robota, a zobrazení reálné situace do takových termů*. Navrhovaný mechanismus výběru akce *ASM odvozuje výsledné rozhodnutí (kterou akci jako další vybrat) syntézou bazálních akcí*. Připomínám zde velmi užitečný *důsledek eko-gramatického systému, parametrizaci pravidel*. Parametrickým vyjádřením problému výrazně redukuji množství pravidel. Připravuji si tak i nástroj pro lepší přizpůsobení reaktivní části ASM na reálný problém.

Pozitivním v návrhu je *nadřazení chování* (okolí, ostatních robotů apod.) nad *bazálními akcemi ve smyslu převrstvování*. Tímto způsobem eliminuji v té době uváděnou nevýhodu superpozice, *totiž posun významu při snaze o kompromis různorodých závěrů*. Tento princip dále rozvíjím podle aktuálních potřeb buď jen jednotlivého robota či celého Kolektivu robotů.

Dále jsem zde navrhnul, jak vhodnost rozhodování ASM při výběru dalšího kroku (akce) dodatečně podložit ještě dalšími argumenty. *K vjemům zpracovávaným reaktivní částí dodávám významově blízké vnitřní podněty*. Nově zde posuzuji i minoritní (vedlejší) tendence pro dosažení cíle, které se vyskytly až bezprostředně při akci. *Váhu takových podnětů započítávám do výsledného rozhodnutí superpozicí v pohybovém schématu*. Pro podchycení minoritních trendů v ASM používám:

- zobecnění podnětů i na významově blízké varianty, v normálním rozložení $N(\vec{S}, \sigma)$
- bazální akce pro vnitřní podněty a autostimulace mezi nimi
- adaptivní doplnění nových bazálních akcí, odvozených v režimu učení bez učitele

Kapitola 4

Experimenty

V této kapitole dokládám mé experimenty, sloužící k ověření uváděných vlastností hybridní architektury. Celá její první část se věnuje **popisu a vyhodnocení** mnou realizovaných **počítačových simulací**, které jsem prováděl na (rovněž mnou navrženém) modelu zde popisovaného hybridního řídicího systému. Pro ověření zavedené teorie v praxi jsem také využil platformu (skupinu čtyř reálných) jednotně navržených mobilních robotů. Takovýto „Kolektiv robotů“, na jehož realizaci jsem se na katedře kdysi aktivně podílel [Svatoš, 1997], mi posloužil pro ověření většiny mých (do té doby jen teoretických) závěrů. Do této práce jsem ale z mnohem širší sady experimentů vybíral jen ty úlohy, které dobře ilustrují splnění mých záměrů v kontextu **Cílů** dizertační práce. Experimenty s roboty vychází z jejich konstrukčních omezení. Konstrukční parametry robotů jsem zde stručně uvedl v příloze B.1. Návrh robotů je zcela původní. Tato má práce, opírající se finálně právě o platformu mobilních robotů, tak významně zhodnocuje i mé předchozí dlouholeté zkušenosti z realizace mobilních robotů a výzkumu jejich řízení.

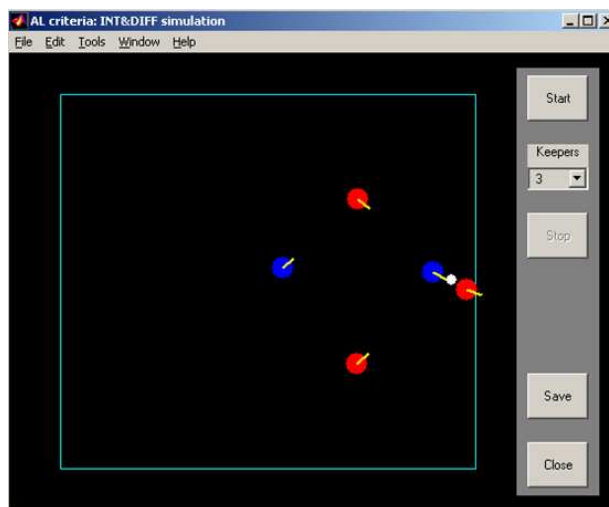
4.1 Experimenty – ověření návrhu simulacemi na modelech

4.1.1 Simulace útočení a obrany ve srovnávací úloze Keepaway

V následující části v experimentech doložím *správnou funkci navrhovaného hybridního mechanismu výběru akcí ASM* a to nejprve **ověřením v simulacích** (při reaktivním ošetření situace v doméně Robotického fotbalu). Výchozí situace představuje hráče ve dvou skupinách s rozdílným záměrem. Budu zde popisovat simulaci soupeření o míč ve formaci dle obrázku 4.1: 3 (modří) na 2 (červené). Na jedné straně se zde skupina tří robotů snaží *maximalizovat dobu, po kterou udrží míč pod svou kontrolou*. Naopak cílem pro skupinu dvou protihráčů je *minimalizace doby, než získají míč zpět na svoji stranu*. Standardní průběhy této Stonem navržené srovnávací úlohy Keepaway [Stone a kol., 2006] jsou k dispozici na stránkách [web Keepaway].

Provedení dále popisované *Keepaway simulace* se od původní verze [Stone a kol., 2006] liší v použitém prostředí **Matlabu**. Příklad řešení *jedné herní situace a její vykreslení* je patrný z aplikačního okna na obrázku 4.1. Grafický výstup simulace je ryze informační; vykresluje pouze aktuální pozice robota v kartézských souřadnicích.

Simulace následovaly s časovým odstupem po jednotlivých experimentech s reálnými roboty (uváděných v kapitolách 4.2 a dalších). Oproti navigaci reálných robotů v uzavřené aréně tentokrát simulace v doméně Keepaway dávají ještě jeden pohled na uplatnění navrhované architektury řídicího systému. Větší váha je v této srovnávací úloze kladena na nezbytnou spolupráci ve skupině. Nicméně pokud šlo o prvotní nastavení ASM jednotlivých robotů pro simulaci, ze zkušeností z dřívějších experimentů jsem rutinně vyšel i při skládání řídicího systému pro Keepaway doménu.



Obrázek 4.1: Simulace použitelnosti kritérií na standardní srovnávací úloze Keepaway. Ve čtvercovém poli se útočníci (modří) snaží získat míč, který je iniciálně v držení obránců (červených). Poloha míče je vyznačena bílou značkou.

A. Obecné předpoklady počítačové simulace

Počítačová simulace má demonstrovat *vzájemné interakce robotů, řízených dvěma typy ASM*. Stejně jako v experimentech s eko-gramatickým systémem stavím proti sobě možnosti *řízení na základě reaktivního a intencionálního modelu*.

Odděleně uvažuji o řízení ve skupině s míčem v držení (dále budu hovořit o skupině **obránců**) a řízení skupiny jejich protihráčů (dále **útočníci**). Zatímco skupina obránců pracuje hlavně se základním (reaktivním) odvozením akce (viz kapitola 3.3), skupinu obránců vybavuji hybridním systémem včetně intencionálního modelu.

Parametry systému pro simulaci jsou blíže rozvedeny v příloze B.2.

B. Cíle simulace

Za teoretických předpokladů, uváděných podrobně v kapitole 3, simulací **ověřuji vlastnosti hybridního mechanismu výběru akce**. Ve svých simulacích na počítači *kombinuji funkce reaktivního modelu, intencionálního modelu a modelu adaptace*.

Reaktivní model popisuje pravidla navádění robota na určenou pozici – momentální cíl.

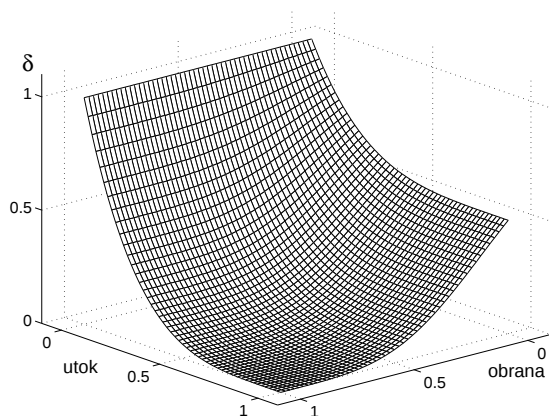
Intencionální model doplňuje do rozhodování i další podněty, které přímo nevyplývají ze simulované senzorky (zde z absolutní polohy hráčů a míče, a odvozených relativních pseudopodmínek). Významnou měrou určuje podobu momentálního cíle také *autostimulace*.

Modul adaptace slouží k dodatečnému rozšíření vlivu bazální akce na ostatní, významově podobné situace. Metrika významu je v **tabulce B.1**.

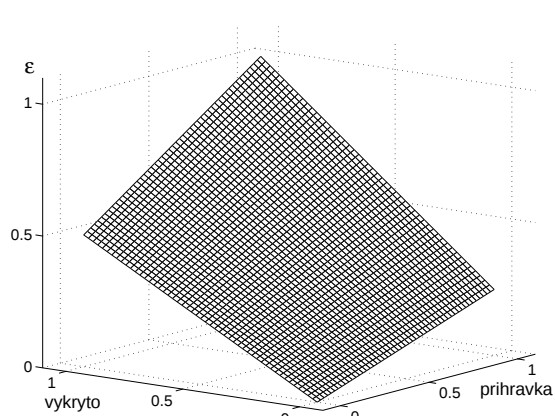
4.1.2 Kritéria hodnocení

V simulaci je hlavní otázkou kvalita chování každé ze skupin soupeřících o míč. Relaci hráč–míč odpovídá i konkrétní interpretace kritérií δ a ϵ , vymezených vztahy (3.10) a (3.11). Metodika kritérií je tedy nezávislá na mechanismu výběru akce; odvozuje se výhradně z pozice robotů.

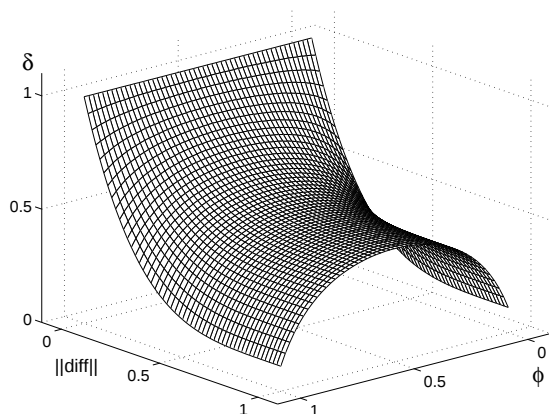
Už při zavedení kritérií v kapitole 3.5.2 byl zmíněn jejich omezený rozsah hodnot. V této simulaci uvažují kritéria v závislosti na dvou normovaných parametrech. Cílová metrika posuzovaných parametrů je patrná z obrázku 4.2. Vzhledem k záběru hráčů v rámci simulovaného herního pole (čtverec 2×2) beru v úvahu podněty maximálně do vzdálenosti 1. Vzdálenější události pak spadají do kategorie „velká“ – viz následující **tabulky 4.1 a 4.2**. Situace mimo rozsah tedy hodnotím stejně jako odpovídající mezní stavy.



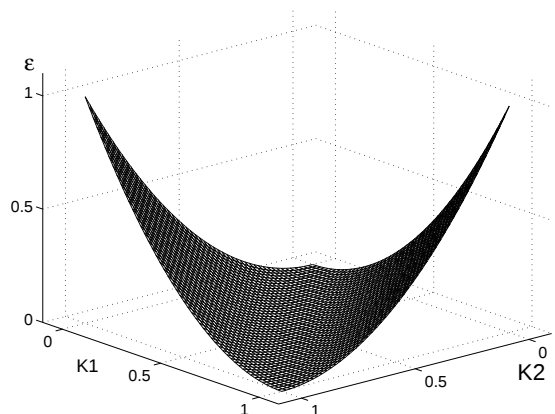
(a) obránce : individuální δ



(b) obránce : skupinové ϵ



(c) útočník : individuální δ



(d) útočník : skupinové ϵ

Obrázek 4.2: Navržené kritéria pro posouzení simulovaného chování. Rozlišují se kritéria chování skupiny obránců (a-b) a útočníků. Individuální kritérium obránců (a) δ závisí na vzdálenosti k nejbližšímu spoluhráči (osa «obrana») a protihráči (osa «utok»). V části (b) se pro skupinové kritérium ϵ obránců hodnotí vykrytá plocha (osa «vykryto») a velikost prostoru neobsazeného protihráči (osa «prihravka»). Při posouzení kvality útočníka se v individuálním pohledu δ (c) porovnává odchylka vzdálenosti (osa $\|diff\|$) a směru (osa ϕ). Nakonec skupinový pohled (d) na útočníky ϵ kombinuje vzdálenost jednotlivých útočníků od míče. Vzdálenosti jsou vztaženy k poloviční straně čtverce vymežujícího simulační oblast, úhly k π .

Návrh a interpretaci kritérií načrtnutých v obrázku 4.2 v následujících odstavcích dále samostatně popisují pro případ každé ze soupeřících skupin. Vychází se ze znalosti mezních stavů. Od nichž pak ostatní stavy volně aproximují. Pouze při volbě takové aproximace je třeba mít na zřeteli, že u útočníka rozhoduje o aplikování pokusných kroků.

A. Diferenciální kritérium δ obránců

U obránce δ kritérium vynucuje zpracovat míč. Na chování obránce poměřuje dva soupeřící trendy. Na jedné straně to jsou možnosti obránce, který kritérium vyhodnocuje. Na druhé straně stojí možnosti protihráčů. Na obrázku 4.2(a) je v tomto smyslu naznačeno hodnocení chování obránce. Tentokrát v závislosti na vzdálenosti obránce a útočníka od míče. Vzhledem k rolím obránce si je třeba také uvědomit, že v konečném důsledku se chování bude vyhodnocovat vůči jedné ze dvou možných poloh. Cílem chování Rozehrávka je míč, cílem chování Záloha je volná oblast.

Hraniční případy očekávané kvality chování jsou diskutovány v tabulce 4.1. Z těchto případů je pak navržen vlastní průběh hodnocení kvality chování ve skupině obránců při simulaci.

B. Integrovaný kritérium ϵ obránců

Pro interpretaci vztahu (3.10) v integrovaném kritériu ϵ posuzují především plochu konvexního obalu polygonu odvozeného z polohy hráčů na ploše. Pokud jde o cenu takového pokrytí, nedělám rozdíl mezi uvažovanými rolemi obránců. Průnikem, o který se hodnocení snižuje, je naopak plocha vykrytá protihráči. Tato úvaha se dále promítla do případového výčtu hodnot ϵ v tabulce 4.1.

vzdálenost od míče		typ reakce	hodnota δ
útočník	obránce		
malá	libovolná	ošetření havárie	1
velká	malá	cílová situace	0.5
velká	velká	běh naprázdno	0

vykrytá plocha	prostor k přihrávce	typ reakce	hodnota ϵ
malá	malý	individuální	0
malá	velký	částečný vliv skupiny	0.25
velká	malý	částečný vliv skupiny	0.5
velká	velký	skupinová	1

Tabulka 4.1: Případové odvození průběhu kritérií Obránce. Tabulka vytyčuje mezní stavy kritérií při kombinaci dvou nezávislých veličin. V horní části se pro odvození diferenciálního kritéria δ uvažuje vzdálenost spoluhráče a protihráče od míče. V dolní části se pak do integrovaného kritéria ϵ uvažuje polygon vymezeného skupinou oproti velikosti plochy, která je mimo kontrolu protihráčů. Mezi vytyčenými stavy se kritérium aproximuje monotónní funkcí těchto dvou veličin.

C. Diferenciální kritérium δ útočníků

Podstatnými složkami pro diferenciální vyhodnocení reaktivity útočníka je odchylka v úhlu a vzdálenosti. Větší důraz je kladen na vzdálenost $\|diff\| \rightarrow 0$ (striktní požadavek na zachy-

cení míče) a úhel $\Phi \rightarrow \pm\pi$ (klidové směřování robota na cíl). V souvislosti s těmito případy v tabulce 4.2 podle typu reakce uvádím i očekávané hodnocení kritériem δ . Z těchto mezních případů aproximuji ideální průběh kritéria δ . Navržené hodnocení diferenciální kvality je pak vykresleno na následujícím obrázcích 4.2(c) a 4.2(d).

D. Integrální kritérium ϵ útočníků

Cílem skupiny útočníků je minimalizovat čas do zachycení míče. Dosažení takového cíle se dá také interpretovat ze vzdálenosti jednotlivých útočníků od skupinového cíle (míče). Mezní vzdálenosti hodnotím v tabulce 4.2. Grafický průběh kritéria rozvedeného do prostoru $\langle 0; 1 \rangle^2$ je naznačen na obrázku 4.2(d).

odchylka od míče		typ reakce	hodnota δ
$\ diff\ $	$\pm\Phi$		
malá	Libovolný	ošetření havárie	1
velká	malý	cílová situace	0.5
velká	velký	běh naprázdno	0

vzdálenost od míče		typ reakce	hodnota ϵ
K_1	K_2		
malá	malá	individuální	0
velká	velká	chybová	0.5
velká	malá	skupinová	1
malá	velká	skupinová	1

Tabulka 4.2: Případové odvození průběhu δ kritéria Útočníka. Tabulka vytyčuje mezní stavy při kombinaci dvou nezávislých veličin – vzdálenosti a úhlové odchylky od cíle (míč nebo místo příhodné pro přihrávku). Mezi vytyčenými stavy se kritérium aproximuje monotónní funkcí těchto dvou veličin. V dolní části se pak do integrálního kritéria ϵ uvažuje vzdálenosti jednotlivých členů skupiny od míče. Mezi vytyčenými stavy se kritérium aproximuje monotónní funkcí těchto dvou veličin.

4.1.3 Zapojení simulačního počítačového modelu

Na navržené skupinové simulaci se ukazuje, že bez kritériálního zaměření je možné formulovat úlohu pro jednotlivce. Kritérium je ekvivalentní jeho pravidlům. Model s adaptací přirozeně rozšiřuje vlastnosti behaviorálního bottom-up přístupu. V obou případech je znalostní část postavena na ohodnocených pravidlech.

Pro simulaci ovšem není nutné hledat ekvivalentní výčet skupinových pravidel. Evidentně bez formulace skupinového kritéria by si řešení vyžádalo mnohem složitější heuristiku. Skupinové kritérium v simulaci působí zprostředkovaně. Požadavek na udržení míče se jím přirozeně dekomponuje podle možností zapojeného hráče.

V simulaci se také ukazuje jeden ze způsobů zapojení kritérií do ASM. Kritérium δ například určuje, kdy dojde k přepnutí bazální akce v rámci ASM. Takový přínos δ kritéria se pak v modelu ukazuje hlavně zprostředkovaně – adaptací. Účelem adaptace je ostatně (stejně jako u δ kritéria) maximalizovat kvalitu rozhodování. V případě této simulace se adaptací posilují alternativní tendence v intencionálním modelu. Aby se předešlo nevhodné reakci, Algoritmus adaptace (5) se používá po překročení $\dot{\delta} > 1$.

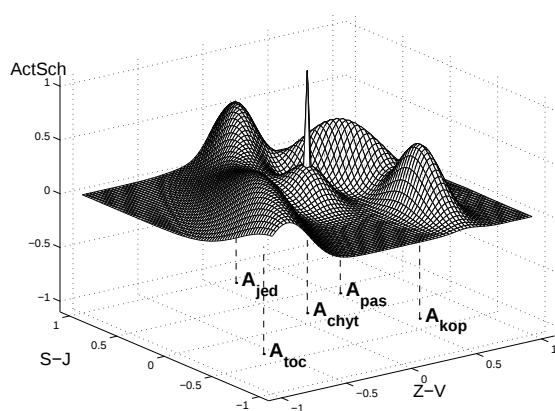
K obdobnému přepínání mezi závěry na úrovni kolektivu dochází zapojením kritéria ϵ . Implicitně je role obránce přímo odvozena od relací v prostředí. Aktivně je může změnit hráč s míčem. Přihrává – a tedy předává roli – spoluhráči v lepším postavení. Tuto kvalitu z pohledu skupiny ohodnocuje právě kritérium ϵ .

4.1.4 Průběh a výsledky simulace

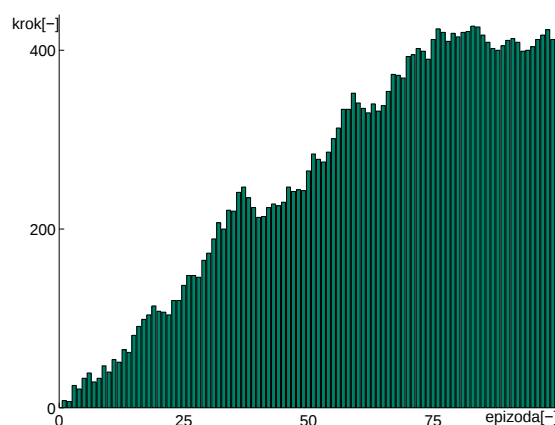
Výše uváděné modely chování obránců a útočníků jsem ověřoval ve 100 epizodách simulace. Každá epizoda začíná vhazováním míče. Ukončena je v okamžiku, kdy míč získává skupina útočníků, nebo když míč končí za hranicí vymezeného pole.

V této pasáži budu hodnotit chování obránců jak dodanými měřítky, tak pragmaticky. To jest dobou, po kterou zůstává míč v držení obránců. Takový markantní výsledek optimalizace je vidět na grafu na 4.3(b), trvání jedné epizody se postupně prodlužuje.

U dodané kvality už není hodnocení určité herní situace tak jednoznačné. V každé epizodě se řeší v podstatě nová situace. K ní jsou ovšem poměrně úzce vztažena obě kritéria. Detailní průběh kritérií tedy nevypovídá jen o dynamice rozhodování, ale také o dynamice řešené situace. Bez znalosti situace (a v této simulaci jednotlivé herní situace ani nemá význam zpětně rekonstruovat) jsou kritéria postavená na poloze jsou mezi jednotlivými simulovanými epizodami vzájemně neporovnatelná. Dále proto na sérii obrázků 4.4 podávám pouze statistické vyhodnocení naměřených kritériálních funkcí.



(a) adaptované schéma



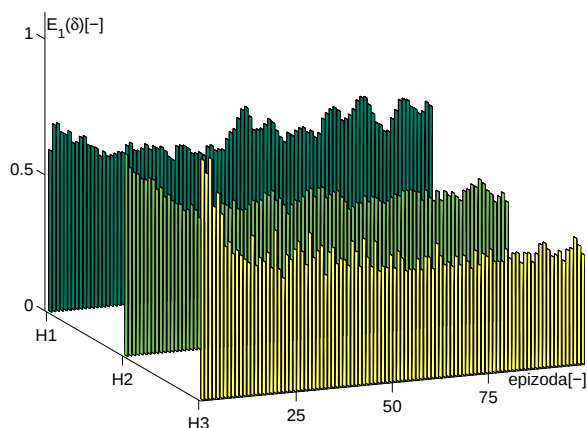
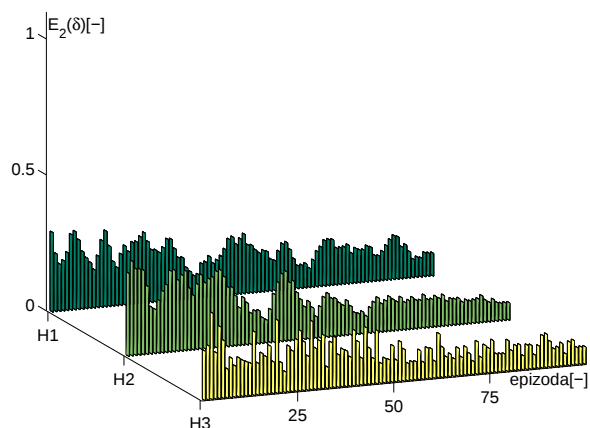
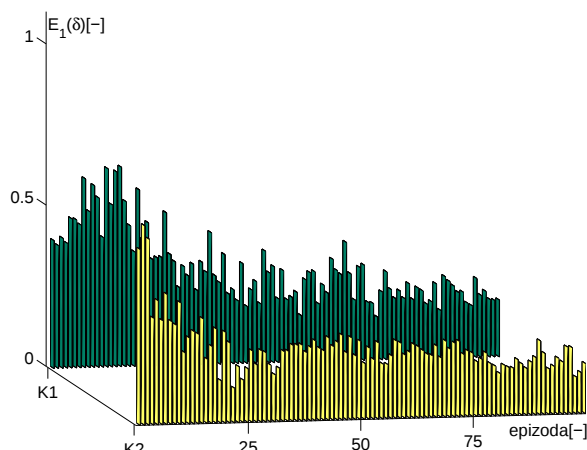
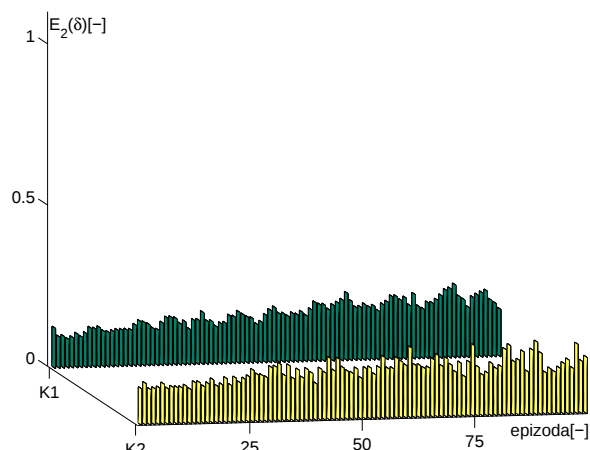
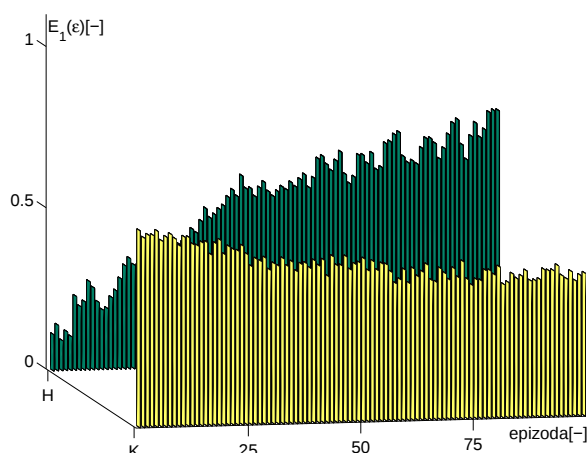
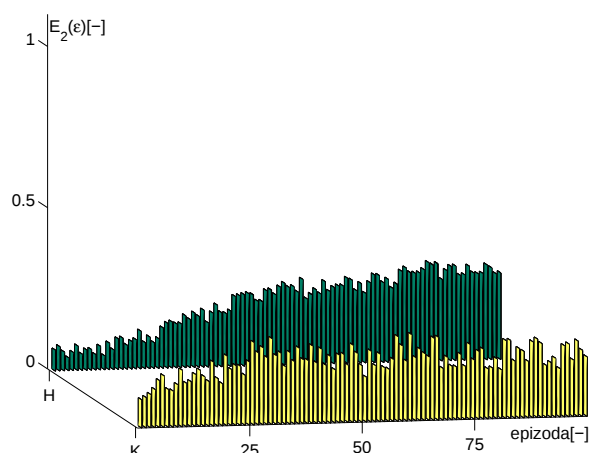
(b) držení míče skupinou obránců

Obrázek 4.3: Výsledky simulace navrženého hybridního ASM. V obrázku (a) je adaptovaný intencionální model obránce (příklad H1) a (b) rostoucí úspěšnost skupiny obránců v průběhu 100 simulačních epizod. Pro srovnání uvádím počet ASM cyklů, než protihráči získají míč na svoji stranu.

A. Chování skupiny útočníků

Pro řízení útočníků byl použit jednodušší systém. Tak i δ kritérium útočníků má menší rozptyl než u obránců (obránci zapojují do chování více variant odpovědi; srovnání je patrné z obrázků 4.4(b) a 4.4(d)). Při větší spolupráci obránců jsou pak i útočníci nuceni častěji přepínat role. Tím lze vysvětlit další mírný nárůst rozptylu δ v pozdějších epizodách.

Obecně hodnoty δ kritéria útočníků jsou poměrně vysoké. Hodnoty ve srovnání v **obrázku 4.4(a) a 4.4(c)** ale spíše spíše poukazují na neadekvátní chování obránců. Útočníci

(a) Obránci : průměrné δ v epizodě(b) Obránci : rozptyl δ v epizodě(c) Útočníci : průměrné δ v epizodě(d) Útočníci : rozptyl δ v epizodě(e) ϵ skupin : průměrné(f) ϵ skupin : rozptyl

Obrázek 4.4: Statistické vyhodnocení kritérií v jednotlivých epizodách simulace. Grafy (a)/(c)/(e) ilustrují průměrné hodnoty, grafy (b)/(d)/(f) rozptyl hodnot v epizodě. První dvě řady shrnují individuální pohled na chování $H_1 \div H_3$ (a-b), respektive útočníků $K_1 \div K_2$ (c-d). Obrázky (e-f) vyznačují integrální kvality útočné (K) a obranné (H) skupiny.

se pak poměrně snadno a rychle dostanou k míči (což je pro ně ošetření kritické situace s vysokým δ). Navíc tím, jak si tento model neudrží žádnou reprezentaci, se v jeho souhrnném hodnocení více odráží průběh kritérií obránce.

Obdobně se jednoduchost řídicího systému útočníků projevuje na rozptylu ϵ kritérií. Srovnání skupinového pohledu na obránce a útočníky je patrné na **obrázku 4.4(f)**.

B. Chování skupiny obránců

Ve strategii obránců je od počátku zdůrazněn kontakt hráče s míčem. Kvalita hodnoceného chování není výrazněji závislá na výsledku útočníků. Více se v chování projevuje vlastní intencionální model obránců. Model obránce zavádí do rozhodování další aspekty, a jinou skladbu podnětů. Velice zjednodušený příklad takových podnětů uvádím dále v popisu experimentálního ověření interakcí eko-gramatického systému v kapitole 3.4.4.

Zpočátku lze chování obránců hodnotit jako značně individuální. S ním ovšem rychle ztrácí míč v souboji jeden-na-jednoho. Později intencionální model obránce více zpřesňuje vnitřní reprezentaci v pohybovém schématu – adaptované rozložení podnětů pak ukazuje **obrázek 4.3(a)**. Vnitřně se eliminuje podíl nerozhodnutelných stavů, což se navenek projevuje v nižším rozptylu δ obránců.

Na obrázku 4.3(a) je také vidět, jak se adaptací v intencionálním modelu vytváří těsnější vazba mezi bazálními akcemi *posun* a *otáčení*. V relaci těchto dvou bazálních akcí je majoritnější *posun* vpřed. Na hodnotách ϵ adaptovaného řídicího systému se dále ukazuje víceméně očekávaná vlastnost této simulace: rostoucí ϵ naznačuje (stále ještě není $\epsilon \rightarrow 1$), že z podstaty zvolené úlohy lze v úloze vyzpozorovat (a využít) koordinaci mezi jednotlivými členy skupiny.

4.1.5 Zhodnocení simulací chování na počítači

Na mnou simulovaném chování *Kolektivu* se těžko vymezuje hranice, zda se ve výsledku projevuje více individuální δ nebo skupinové ϵ . V dokumentované simulaci jsem obě kritéria použil jako doplňkové signály pro změnu chování a jako takové se bezprostředně neprojevují. O oprávněnosti takových zásahů můžeme usuzovat až kumulovaně s určitým časovým odstupem.

V behaviorálním přístupu (bottom-up) se nicméně v tomto případě ukazuje alternativa, která stejně jako cíle mé práce vychází z nejjednodušších možností řízených robotů. Za cenu heuristiky umožňuje nalézt východiska z udané situace bez nutnosti detailně analyzovat hlavní stavy, přechody mezi stavy a váhu takových přechodů. Simulace ověřila případ, jak lze diferenciálními kritérii doplnit (v ASM) gradientní heuristiku. Připomínám, že i v této praktické ukázce ASM přepíná v okamžiku poklesu gradientu funkce vztahované na několik vybraných klíčových parametrů. Jejich růst je patrný například z grafů na **obrázcích 4.2(c)** a **4.2(c)**. V takto pojaté simulaci se ukazuje možný převod procedurálního řešení do bottom-up výběru okrajově vymezeného kritérii, které v podstatě kopírují podmínky reálné cílové situace.

4.2 Experimenty – ověření návrhu na reálných robotech

Na rozdíl od *Experimentů a ověřování navržených postupů v Simulacích na modelech*, v této části se jedná o experimenty prováděné v reálných situacích (úlohách) s reálnými roboty.

Příprava platformy mobilních robotů pro ověření mnou navrhované hybridní architektury řízení je důležitou částí ověření mého návrhu. Ta postupně zahrnuje:

- volbu měřítka pro hodnocení experimentů
- přípravu senzorické části robota
- volbu bazálních akcí

V tomto pohledu nejprve zrekapituluji podstatu všech dále uváděných experimentů. V následujících podkapitolách pak shrnuji detailní průběh experimentů, kde rozvedu zde nyní zdůrazněné principy:

4.2.1 Obecné aspekty realizace návrhu – kalibrace měřítka

A. Apriorní váhy rozhodnutí

Dále uváděné experimenty jsou hodnoceny s ohledem na rozsah vah W používaných pravidel reakční báze. Váhy jsou normalizované v intervalu $\langle 0, 1 \rangle$. Hodnotou W oceňuji

- pravidla pro řešení kritických situací $W \rightarrow 1$
- pravidla užitá v cílovém stavu $W \rightarrow 0,5 = W_{mod}$
- pravidla užitá v klidu $W \rightarrow 0$

V reálném experimentu aplikace Umělého života tak váhy pravidel odpovídají četnosti použití entit nízké úrovně. Snažím se o návrh vyvážených pravidel, aby nejčastěji užívaná pravidla byla charakterizována vahou $W = 0,5$.

B. Kritéria pro aposteriorní hodnocení

Chování Kolektivu hodnotím integrálním kritériem ϵ a chování jednotlivých robotů diferenciálním kritériem δ . U každého z experimentů uvádím iterativní postup výpočtu měřítka, splňujícího předpoklady kapitoly 3.5.2 Rozbor kritérií kvality navržené akce ASM. Odhaduji v něm náklady na použití akce opět z četnosti entit Umělého života; v cílovém stavu oceňuji aktivity s vahou $W \rightarrow W_{mod}$. Obě kritéria proto normalizuji v intervalu $\langle 0, 1 \rangle$, s možností jejich zpětného zapojení do aktuálního běhu algoritmu:

- Posun skutečné hodnoty *diferenciálního kritéria* δ v požadovaném stavu oproti W_{mod} lze brát i jako doporučení pro korekci váhy procedur, jenž jsou pro požadovaný stav charakteristické.
- *Integrální kritérium* ϵ je současně podmínkou pro zastavení iterativního algoritmu.

Těmito kritérii v následujících podkapitolách vyhodnocuji trajektorii, kterou jsem rekonstruoval z časových záznamů změn chování.

Diferenciální kritérium δ v této časové ose například signalizuje okamžiky vhodné pro změnu chování.

Na integrálním kritériu ϵ se pak dobře ukazuje konvergence řešení pro $\epsilon \rightarrow 0$.

4.2.2 Obecné aspekty realizace návrhu – Senzory robota

Konstrukce robota byla (při jeho návrhu, na kterém jsem se podílel) podřízena požadavkům instinktivních rozhodnutí v reálných situacích. Konstrukce tak např. výrazně omezuje vstupní informace, vyhodnocované při instinktivní volbě bazální akce.

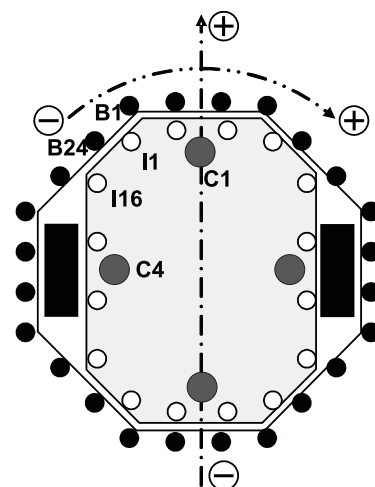
Hlavní charakteristiky robota nyní stručně popíši: **Senzorická výbava robota** a její rozmístění je schématicky vidět na obrázku 4.5. Robot se čtyřmi koly (pro větší pohyblivost) má v půdorysu *tvar oktagonu*. Akční a senzorické prvky jsou na podvozku rozmístěny v osové symetrii. Na čerchované ose je vyznačena orientace pohybu $krok_+$ a $krok_-$. Je patrné souměrné umístění dvou bočních hnaných kol. Dvojitým čerchováním je vyznačena orientace rotačních pohybů $toč_+$ a $toč_-$.

Ověření vlastností zde zaváděné hybridní architektury řízení mobilního robota jsem postavil na senzorech, uváděných v následující tabulce. Po obvodu robota je symetricky umístěno **16 infračervených senzorů I_n** , **24 mechanických nárazníků B_n** a **čtyři citlivá směrová čidla světla C_x** .

Podstatné **informace o bezprostředním okolí** robota tak dostávám od infračervených detektorů přiblížení. Pro tyto **proximitní čidla** používám někdy v dalším textu i název *infračervené nárazníky* nebo jen IR. V experimentech robot reaguje (s předstihem) hlavně na stav těchto dvanácti úzce (cca 35°) směrových proximitních čidel.

Z informace ze *čtveřice senzorů lokální intenzity světla* odvozuji přímo směr, odkud k robotu přichází ze vzdáleného silného zdroje nejvíce světla.

Jako doplňkovou informaci pro výběr instinktivního rozhodnutí zahrnuji i tu, získanou ze sady *mechanických spínačů*. Těchto 24 mechanických nárazníků slouží hlavně *pro detekci kritického přiblížení*.



Obrázek 4.5: Konstrukce a základní charakteristiky mobilního robota..

Robot v půdorysu má tvar oktagonu. Akční a senzorické prvky jsou na podvozku rozmístěny v osové symetrii. Na čerchované ose je vyznačena orientace pohybu $krok_+$ a $krok_-$. Podle této osy je patrné souměrné umístění hnaných kol. Dvojitým čerchováním je vyznačena orientace rotačních pohybů $toč_+$ a $toč_-$. Po obvodu jsou umístěny infračervené senzory $\circ I_n$, mechanické nárazníky $\bullet B_n$ a světelná čidla $\bullet C_x$.

senzor	symbol pro podnět		mohutnost n
	bezprostřední	akumulovaný	
aktivní infračervený nárazník	I_n	-	1..12
aktivní mechanický nárazník	B_n	—	1..24
senzory ve stavu klidu	F	—	1
světelný senzor	C_x	$\mathbb{C}_x$	4 složky: s,j,v,z

Tabulka 4.3: Redukce vjemu do vnitřní reprezentace eko-gramatického systému. Pod uvedenou symbolikou se na senzory odkazují v dále uvedených pravidlech. V experimentech robot reaguje hlavně na stav dvanácti proximitních čidel. Pro detekci kritického přiblížení pak slouží 24 mechanických nárazníků. Informace o směru vzdálených zdrojů světla se určuje z měření lokální intenzity čtyřmi světelnými čidly. Ve většině pravidel používám bezprostřední hodnotu senzoru. V úloze gradientní navigace stavím pravidla na integrální hodnotě, intenzitě naakumulované po dobu jednoho kroku vyhodnocovacího algoritmu.

V tabulce 4.3 uvádím zavedené značení senzorických vstupů. Senzorické vjemy takto shodně označuji i při popisu užitých *odvozovacích pravidel* (Příloha B) Na stejných literálech jsem předvedl i mapování senzorických vstupů do pohybového schématu na obrázku 3.5. Každý z aktivních infračervených nárazníků je ve schématu interpretován v souřadnicích, odpovídajících poloze senzoru na podvozku robota. Ze světelných senzorů se ve schématu reprezentuje pouze směr nejintenzivnějšího světla. Ve většině pravidel používám bezprostřední hodnotu na senzoru v okamžiku jeho vyhodnocování. V úloze gradientní navigace stavím pravidla na intenzitě naakumulované po dobu jednoho kroku vyhodnocovacího algoritmu.

Omezené možnosti senzorů omezují volbu experimentálních úloh i po technické stránce. *Navigace robota je založena na interakci s detekovanými objekty.* Možnosti hybridní architektury ověřím ve dvou druzích navigace – při:

- **Kontaktní navigaci** – v úloze *detekce pasivních překážek*,
- **Gradientní navigaci** – v případě úlohy *detekce aktivního světelného zdroje*.

4.2.3 Obecné aspekty realizace návrhu – Volba bazálních akcí

Důležitým aspektem realizace návrhu v praxi je „Volba bazálních akcí a jejich kombinace při navigaci“. Při kontaktní a gradientní navigaci demonstruji možnosti zavedené hybridní architektury na kombinaci dvou typů bazálních akcí:

- **Jízda v přímém směru.** Dále budu pracovat s bazálními akcemi pro jízdu vpřed $krok_+$ a jízdu vzad $krok_-$.
- **Otáčení robota na místě.** Otáčení je realizované bazálními akcemi $toč_+$ a $toč_-$.

Pokud dojde ke změně podmínek, rozpracovaná bazální akce je ukončena. Bezprostředně poté se spouští nově odvozená bazální akce. Zásadní změna podmínek se pak odrazí ve změně chování, a tedy i použitím pouze určité části pravidel.

Chování jsem experimentálně skládal z entit uvedených v tabulce 4.4. Svým chováním robot reaguje na senzorický vjem pohybem určeným směrem. Pohyb v oblasti bez překážek lze řešit chováním **průzkum-naslepo**. Při detekci překážek pak zavedený mechanismus výběru akce přepíná na specifitější chování. V tabulce přímo udávám scénáře, v kterých se cíleně zaměřená chování uplatní. Pravidla pro skládání daného chování v sobě udržují i podmínky platnosti takového chování. Jakmile dojde ke změně podmínek, mechanismus ASM podle určeného pravidla mění chování.

Všechna chování jsou postavena na bazálních akcích; vesměs se robot pohybuje přímo anebo se otáčí na místě. *Symboly pro akční entity a senzorické literály tvoří abecedu použitelnou pro odvození akce (jako reakce na dané chování).* Na neuvedené podněty je chování necitlivé. Primárně (přednostně) se reaguje na události detekované proximitními IR čidly. Podle povahy chování se však také vyhodnocuje stav (informace z) čidel světla či nárazníků.

chování	abeceda	účel
START	S	startovní symbol gramatiky; <i>počáteční stav všech scénářů</i>
sledování-zdi	I_n, B_n, F $krok_+, krok_-, toč_+, toč_-$	sledování obvodu překážek; <i>více scénáře na obrázku 4.6, 4.9, 4.12</i>
průzkum-naslepo	I_n, B_n, F $krok_+$	vyhledávání překážek; <i>užito ve scénářích na obrázku 4.6, 4.9</i>
pronásledování	I_n, B_n, F $krok_+, krok_-, toč_+, toč_-$	sledování pohyblivých objektů; <i>scénář užití naznačen na obrázku 4.9</i>
zacílení	$I_n, B_n, F, C_x, \mathbb{C}_x$ $krok_+, toč(\alpha)$	gradientní navigace na vzdálený cíl; <i>užívám ve scénáři dle obrázku 4.12</i>
průzkum-plochy	I_n $krok_+, krok_-, toč_+, toč_-$	bezkolizní pohyb; <i>užívám ve scénáři dle obrázku 4.12</i>
STOP	T	indikace cílového stavu; <i>koncový stav všech scénářů</i>

Tabulka 4.4: Testovaná chování systému. Všechna chování jsou postavena na bazálních akcích; vesměs se robot pohybuje přímo anebo se otáčí na místě. Symboly pro akční entity a senzorické literály tvoří abecedu použitelnou pro odvození akce za daného chování. Na neuvedené podněty je pak chování necitlivé. Primárně se reaguje na události detekované proximitními čidly. Podle povahy chování se také vyhodnocuje stav nárazníků či světelných čidel.

4.2.4 Společné znaky experimentů na reálných robotech

V předpokladech všech experimentů jsou konstrukční omezení laboratorního robota použitelná ve všech experimentech. Robot se pohybuje pouze omezenou rychlostí, obdobně je omezena i rychlost vyhodnocení reálné situace senzory. S ohledem na připomenutá konstrukční omezení robotů jsem pak účelově volil následující tři experimenty s reálnými roboty (dál popsané). Účelem bylo **ověřit předpokládaný přínos**

- pohybového schématu,
- jeho adaptace na základě pokusu a omylu,
- vzájemné autostimulace chování,
- a implicitní interakce mezi roboty ve smyslu eko-gramatického systému.

Společným znakem takových experimentů je *reaktivní vyhodnocení proximitní znalosti o okolí*. Proximitními senzory jsem se rozhodl popisovat jak statické překážky, tak později i objekty s vlastní dynamikou.

Shrnu-li:

V prvním oddíle jsem zdůraznil předchozí teoretické závěry návrhu. Zopakované závěry jsou předmětem zde uváděných experimentů.

Ve druhém oddíle jsem upřesnil, jak úlohy omezuje senzorická výbava robota. Proximitní čidla jsou použitelná pro kontaktní navigaci robota. Z detektoru intenzity světla lze odvodit vektor pro gradientní navigaci robota.

Ve třetím oddíle jsem načrtl chování, která lze v kontaktní a gradientní navigaci použít.

4.3 Experimenty srovnávající různé realizace volby bazálních akcí

Experimenty patří do úloh kontaktní navigace. Smyslem experimentů je prokázat přínos pohybových schémat při syntéze jednoduchých bazálních akcí do chování.

Srovnám proto tři rozdílné **varianty realizace volby bazální akce v Algoritmu 1:**

- **Reaktivní realizace:** samostatná volba aktuálně nejlépe hodnoceného pravidla,
- **Intencionální realizace:** zahrnutí významově podobných závěrů v pohybovém schématu (Algoritmus 4) bez možnosti autostimulace,
- **Adaptivní realizace:** závěry pohybového schématu modifikovaného v pokusných krocích (Algoritmus 5).

	varianta 1	varianta 2	varianta 3
chování	sledování-zdi, průzkum-naslepo	sledování-zdi, průzkum-naslepo	sledování-zdi, pokusný-krok průzkum-naslepo
abeceda	$I_n, B_n, F,$ $krok_+, krok_-, toč_+, toč_-$	$I_n, B_n, F,$ $krok_+, krok_-, toč_+, toč_-$	$I_n, B_n, F,$ $krok_+, krok_-, toč_+, toč_-$
ASM	čistá reakční báze	pohybové schéma	pohybové schéma, adaptované v pokusných krocích
kriteria	distribuce zájmových oblastí	distribuce zájmových oblastí	distribuce zájmových oblastí

Tabulka 4.5: Experimentální ASM při kontaktní navigaci – různé varianty volby akcí. Tabulka shrnuje hlavní znaky jednotlivých variant tohoto experimentu. Chování jsou ve všech modifikacích experimentu stejná, mění se pouze způsob volby akcí v mechanismu ASM. Všechny varianty experimentální úlohy sledování stěn také hodnotím jednotným kritériem, kvalitou pokrytí celé oblasti podél stěn arény.

Experiment prokazuje kvality jednotlivých algoritmů při sledování stěn. Z jednotné abecedy se skládají dvě principiálně rozdílná chování. Změna chování je podle scénáře na obrázku 4.6 určena hlavně událostmi v prostředí robota. Ve volném prostoru se robot pohybuje bez omezení podle pravidel chování průzkum-naslepo. Při detekci překážek mechanismus pravidel používá z reakční báze pravidla chování sledování-zdi.

Cílový stav scénáře na obrázku je přímo určen kvalitou integrálního kritéria ϵ . Po dosažení požadované hodnoty integrálního kritéria ϵ chování sledování-zdi přechází v terminální chování STOP (pravidlo 0.2). Význam jednotlivých pravidel obrázku je uveden v příloze práce.

Cílový stav scénáře na obrázku 4.6 je přímo určen kvalitou integrálního kritéria ϵ . Po dosažení požadované hodnoty integrálního kritéria ϵ chování sledování-zdi přechází v terminální chování STOP (pravidlo 0.2). Význam jednotlivých pravidel obrázku 4.6 je uveden v příloze B.3.

4.3.1 Kriteria pro vyhodnocení experimentu

A. Diferenciální měřítko

Vlastnosti diferenciálního kritéria δ (3.13) splňuje už **normalizovaná váha pravidla**, naznačená v úvodu této kapitoly 4.2.1. Pro potřeby adaptace váhu W_{ASM} do diferenciálního

kritéria δ *filtruji mechanismem exponenciálního zapomínání*:

$$\delta(t) = q \cdot W_{ASM} + (1 - q) \cdot \delta(t - 1), \quad \delta(0) = 0, \quad (4.1)$$

kde W_{ASM} je váha, na základě které ASM určil aktuální bazální akci. V exponenciálním zapomínání zohledňuji posledních pět kroků, $q = \frac{1}{3}$. V diferenciálním kritériu δ se tak odrazí události maximálně z předchozích dvou sekund.

B. Integrální měřítko

V této úloze je cílovým chováním sledování-zdi. Při integrálním *porovnání nejlevnější cest* ve smyslu definičního vymezení (3.12) proto rozhodují *o hodnotě kritéria ϵ pouze náklady na průzkum*. V kvalitě tak rozlišuji pouze jízdu ve volném prostoru, s kvalitou $c_{distr}(t_{pruzkum}) = 1$. Za žádaná považuji ostatní chování (a to včetně později naučených) s cenou $c_{distr}(t_x) = 0$. Vyhodnocuji proti sobě dvě součtové složky: 1) **dosavadní dobu experimentu t** a 2) **cenu času $T_{pruzkum}$** , kdy se robot namísto sledování věnuje jinému chování:

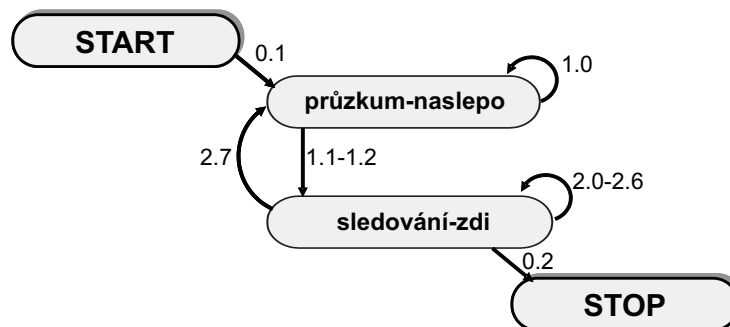
$$\epsilon(t) = \frac{1}{t} \cdot T_{pruzkum}. \quad (4.2)$$

4.3.2 Průběh experimentu

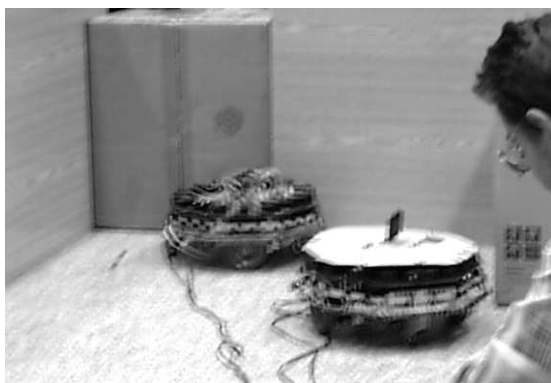
Pro srovnání kvalit těchto algoritmů robot vyřázel ve všech třech případech z téže startovní pozice. Dále pak už je trajektorie robota určena především rozmístěním překážek. Představu o jejich uspořádání dává **obrázek 4.7(b)**. Zajímavá místa trajektorie označuji písmeny. Stejnými symboly jako v tomto schématu pak značím časové okamžiky, kdy robot odpovídajícím místem projíždí. Získané průběhy kritérií δ a ϵ uvádím v grafech na obrázku 4.8.

A. Reakční báze

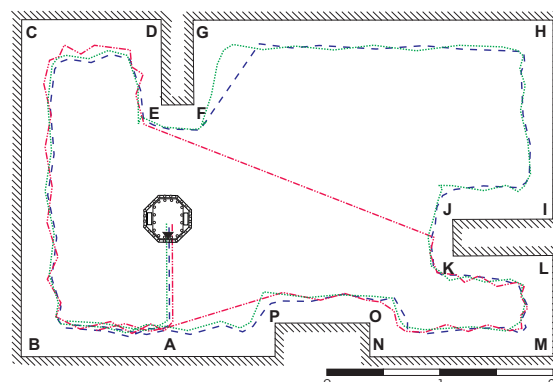
Reakční báze byla navržena tak, aby robot mohl předvést chování podél zdi i bez další optimalizace. Mechanismus výběru pak pracuje na principu majority – vybírá se závěr nejvýše hodnoceného pravidla bez ohledu na další závěry.



Obrázek 4.6: Schéma chování robota v úloze sledování stěn. Graf schematickeje scénář přepínání mezi dvěma chováními robota. Vrcholy grafu jsou použítá chování průzkum-naslepo a sledování-zdi. Na hranách grafu jsou naznačena pravidla stavového přechodu mezi chováními.



(a) Situace



(b) Individuální trajektorie robota

Obrázek 4.7: Experimentální scéna v úloze sledování stěn. V obrázku (a) je ilustrativní záběr na jízdu robotů podél stěny v rohu skutečné arény. Půdorys arény je vyznačen na obrázku (b), fotografie odpovídá levému hornímu rohu. Představu o velikosti arény dává přiložené měřítko — rozměry scény byly $5,5 \times 3,5$ m. Po obvodu scény jsou šrafovány stěny, uvnitř scény jsou šrafovány volně stojící překážky. Robot je zakreslen ve startovní poloze — uprostřed arény. Z rozložení překážek vychází trajektorie při použitých modifikacích experimentu — červenou čerchovanou čarou jízda pomocí reakční báze, modrou čárkovanou podle pohybového schématu a zelenou podle adaptovaného pohybového schématu. U robota je zvýrazněn jeho počáteční čelní směr.

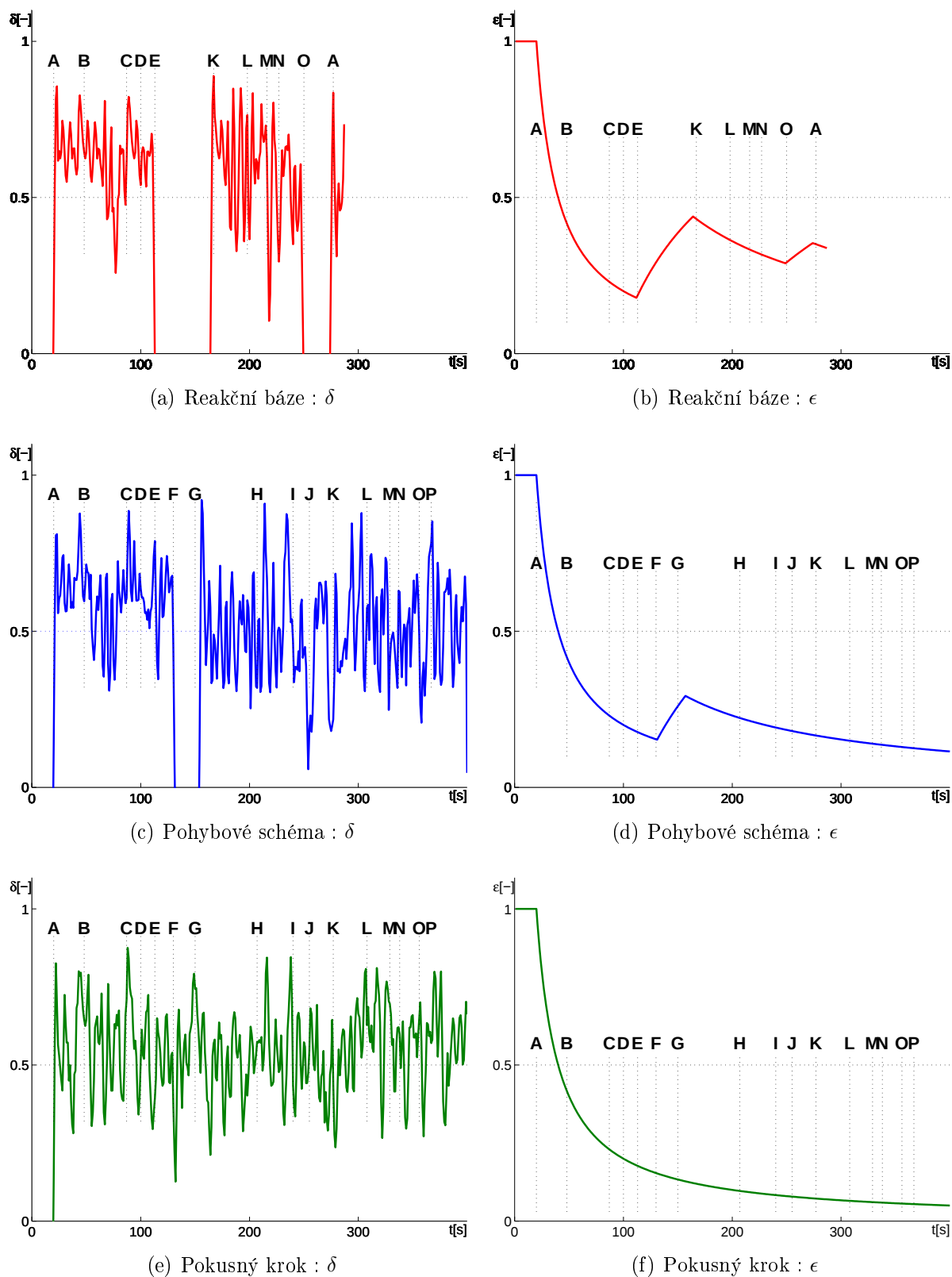
V **diferenciálním měřítku** je už z průběhu δ na obrázku 4.8(a) patrné, jak dobře se robot s reakční bází vypořádá se situací popsanou rozprostřenými příznaky. Pokud robot detekuje v prostředí dostatek příznaků, má ještě možnost na situaci reagovat postupně. Základní pravidla už ale neumožňují ošetření singularit. Jakmile se náhle změní popis situace, robot už nemá možnost se vrátit k předchozím výhodnějším podmínkám. Problém často vzniká z nedostatečně vyhodnocené situace v bodech E, F, J, K, O, P. Situace v okolí ostrých rohů překážek je dobře popsána v příkladu kapitoly 3.5.4. Důsledky takového výpadku pro kritérium δ hodně záleží na podmínkách v prostředí. S pravidly založenými na proximitních čidlech robota od průzkumu odvede až další překážka. A do té doby je δ nulové.

Integrální měřítko ϵ zachycené v grafu 4.8(b) ukazuje opět na problémy v místě E a P. Tentokrát odráží chybné chování robota. Požadovanému chování sledování-zdi odpovídá klesající průběh funkce ϵ . Za každým lokálním minimem kritéria ϵ je vždy změna chování při výpadku detekce proximitními čidly. A stejně jako diferenciální kritérium, i hodnota ϵ odvozená z reakční báze závisí na podmínkách v prostředí. Jak je vidět z grafu 4.8(b), v případě experimentální scény kritérium ϵ roste asymptoticky $\epsilon \rightarrow 1$. Robot se při takto postavené úloze nezastaví.

B. pohybové schéma

Na pohybovém schématu jsem se omezil na způsob, kterým se ze stejných pravidel odvodí vhodná akce syntézou akcí. Použil jsem k tomu přímou transformaci vybuzených akcí do pohybového schématu. Podmínky transformace jsou popsány v příloze, tabulka B.3.

V **diferenciálním měřítku** graf 4.8(c) vykazuje menší rozkmit kritéria δ . Ve důsledku odráží nižší různorodost volených akcí. Potvrzuje předpoklad, že v kompromisu ně-



Obrázek 4.8: Srovnání různých odvození navigace v úloze sledování stěny. Grafy (a)/(c)/(e) ilustrují (δ) odchylku bazální akce v kontextu žádané aktivity, dílčího kroku podél stěny. K nim je na grafech (b)/(d)/(f) vyjádřen (ϵ) podíl požadovaného chování sledování-zdi. Grafy (a)/(b) jsou vztaheny k místům, která projíždí robot řízený reakční bází. Takto vybavený robot mívá celý úsek D÷J. Grafy (c)/(d) se týkají řízení pohybovým schématem. Kvality pohybového schématu s adaptivně doplněným chováním jsou nakonec naznačeny v grafech (e)/(f). Grafy interpretuji ve smyslu tabulky 4.6; prakticky robot s pohybovým schématem více přimyká k sledované stěně.

kolika vhodných bazálních akcí se více přiblíží žádané aktivitě. Stejně se v tomto experimentu eliminují zřejmé excesy chování. Těsnější sledování zdi například snižuje riziko nevhodné reakce na problematické rohy překážek F, G a J. V odpovídajícím časovém průběhu δ se stále ukazují situace vhodné pro adaptaci.

V integrálním pohledu ϵ se také cení eliminace problematických výpadků percepce při objíždění rohu. Průběh na obrázku 4.8(d) už popisuje specifický případ: robot správně vyřešil situaci v bodě E, a jeho selhání v bodě F není fatální. Proto lze jeho jízdu hodnotit úspěšně i v integrálním měřítku. Kritérium ϵ asymptoticky klesá $\epsilon \rightarrow 0$, a tak se po robot zastaví po více než jednom kole.

C. Mechanismus Pokusného kroku

Pro zvýšení citlivosti jsem použil *mechanismus pokusného kroku* – Algoritmus 5.

Diferenciální měřítko δ ve srovnání grafu 4.8(e) s předchozími případy nabývá kvalit požadovaných akcí. Projevuje se v něm především *chování doplněné adaptací*. Adaptivně doplněná sada chování už pokrývá standardní situace řešeného problému. Kritérium δ se významně neodchyluje od střední (a požadované) kvality $\delta = 0,5$.

V integrálním pohledu na obrázku 4.8(e) vystihuje chování monotónní funkce ϵ . Odpovídá použití dvou chování: stejně jako původní sledování zdi se kvalifikuje i nové chování. Takové chování je zavedené adaptací. Obě chování jsou hodnocena jako žádaná. Tato dvě chování v experimentu pokryla všechny běžné situace. Z pohledu robota také dojde ke splnění cílové podmínky ještě před obkroužením celé experimentální arény.

4.3.3 Výsledek experimentálního ověření na robotech

Dosažený průběh interpretuji s ohledem na mezní hodnoty vah W , kalibrovaných v oddíle 4.2.1. Z nich prakticky plyne význam kritérií δ a ϵ , jak jej shrnuji v **tabulce 4.6**.

měřítka	0 ←	→ 0,5 ←	→ 1
ϵ	chování požadované	vhodné a nevhodné chování je v rovnováze	nežádoucí chování
δ	robot překážku neregistruje	krok podél překážky	náraz do překážky

Tabulka 4.6: Význam kritérií ve srovnávací úloze. Tabulka ukazuje interpretaci ve srovnávacím experimentu zavedených kritérií ϵ a δ v oboru hodnot $(0; 1)$. Kritérium ϵ odpovídá entropii požadovaného chování. V měřítku předdefinovaných priorit jednotlivých pravidel pak δ hodnotí, nakolik je aktuální bazální akce vhodná pro řešení úlohy.

V tomto významu kritéria zdůrazňují zejména přínos jemnější kalkulace podnětů v pohybovém schématu. S ním ASM rychleji konverguje k preferenci cílové akce. V navrhované architektuře se s rozšířením konceptu rozšiřují i možnosti robota pro dosažení požadovaného chování.

Rostoucí kvality v řízení, postaveném na *čisté reakční bázi*, na *pohybovém schématu* a na závěrečné *adaptaci* (metodou pokusu a omylu), se projevují v obou dříve zahrnutých pohledech na problém. Z pohledu jednotlivých bazálních akcí užívám jako měřítko kvality

(vhodnosti z nich vybrané aktuální akce) **diferenciální kritérium** δ . U složitějších modelů (*pohybová schémata, doplněná adaptací*) se častěji užívají akce, požadované v cílovém stavu.

Z pohledu chování je asi nejdůležitější **otázka zastavení robota**, hodnocená **integrálním kritériem** ϵ . Stejně jako u výběru bazálních akcí, tak chování robota závisí na těsné interakci s řešeným problémem. Pak je mnohem pravděpodobnější, že nepodchycené stavy nejsou v protikladu se zvoleným chováním.

Těsné vazbě mezi kvalitou dílčích akcí δ a integrálními kvalitami ϵ odpovídá i mé subjektivní **srovnání výsledků tří výše uváděných modelů vnitřní reprezentace řešení**:

- s čistě reakční bází nedošlo ke splnění podmínky pro zastavení robota.
- k zastavení jízdy dle pohybového schématu musel robot projet více než jedno kolo (v závislosti na počtu chybně interpretovaných rohů).
- s pohybovým schématem s adaptací stačilo robotu objet testovanou arénu pouze jednou.

Shrnu-li:

V této podkapitole jsem *na úloze sledování zdi* demonstroval uplatnění *pohybového schématu (Motor Schema)*. Ve stejném prostředí jsem poté ukázal, jak při řízení mobilního robota behaviorálním přístupem lze výchozí pohybové schéma robota rychle **zadaptovat na reálnou (novou) situaci**.

V prvním oddíle jsem shrnul a vysvětlil algoritmické předpoklady mechanismu výběru akce – užití chování, scénáře pro změnu chování a pravidla zvoleného chování.

Ve druhém oddíle jsem zavedl vhodnou *integrální a diferenciální interpretaci kritéria Distribuce zájmových oblastí*. Podle zavedeného kritéria se robot v úloze sledování stěn zaměřuje na akce, kdy může průběžně detekovat blízkou překážku na levé straně.

Ve třetím oddíle jsem *popsal pohyb robota*, pozorovaný během experimentu, a jemu odpovídající průběh kritérií. Kvality různých úrovní mechanismu výběru akce srovnávám na situaci, kdy robot svými senzory nezíská úplný popis ani svého bezprostředního okolí. Reaktivní ASM nízké úrovně od neúplného popisu odvozuje nežádoucí chování. *ASM založená na hybridní syntéze akcí tento problém postupně eliminuje. Ověřil jsem, že navržený hybridní Mechanismus výběru akce s přidáním blokem se i s takovou situací vyrovná a posléze vybírá než žádoucí chování*.

Ve čtvrtém oddíle jsem na párových *kritériích kvality* δ , ϵ ukázal, jak se odrazí kvalitní výběr bazálních akcí i na celkovém chování robota. Díky syntéze obecněji pojatých podnětů robot žádaným způsobem opravdu sleduje terén těsněji – „kopíruje“ přesně obrys (půdorysné: x-y) překážky v aréně. *Eliminuje tak chyby způsobené neúplným popisem lokální situace*. Dále se v navržené architektuře osvědčila možnost adaptivně dodefinovávat zcela nová chování. Rozšířená sada chování pokryla při experimentech s reálnými roboty většinu standardních situací řešené úlohy *sledování zdi*.

V dalším oddíle jsem zdůraznil významné znaky pozorovaného chování reálného robota s ohledem na původní teoretické předpoklady. Praktickým ukazatelem je zde bezproblémové *dosažení cílového stavu*, a tím i zastavení celého Algoritmu. Pomocí ASM nižší úrovně nelze dosažení cílového stavu vždy zajistit. K tomu se nesporně lépe hodí hybridní mechanismy ASM, zahrnující i intencionální přístup, příp. i blok realizující adaptaci.

4.4 Experimentální ověření autostimulace na robotech

Na novém chování robota nyní ukáži, jak lze rozšířit chování robota pomocí intencionálního modelu. Budu vyhodnocovat jízdu robota za *naváděcím robotem*. Chování sledujícího robota dále dokládám v příloze. Na stejném principu (pokusy uskutečněny v roce 2001) robot sledoval i figuranta (doloženo videonahrávkou C). **Intencionální model sledujícího robota** vychází ze *zjednodušeného pohybového schématu*:

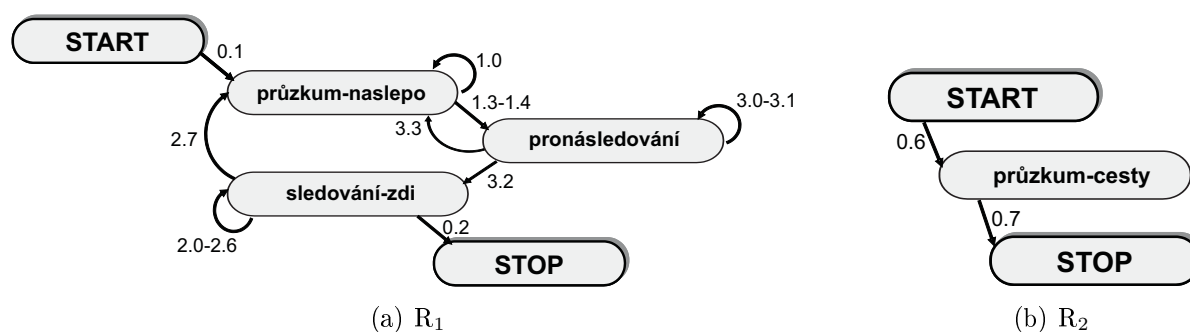
- Rozšiřuji sadu reakcí použitých v předchozím *srovnávacím experimentu*, přičemž situace v okolí rohů překážek ponechávám nepokryté.
- Tyto termy doplňuji na na úrovni intencionálního modelu *autostimulací*.
- *Intencionální model slouží i jako paměť vnitřní kritériální funkce* robota (s funkcí zapomínání).

Vlastnosti a hlavní znaky chování robotů, použitých v tomto experimentu jsou v této tabulce 4.7.

	R ₁	R ₂
chování	sledování-zdi, průzkum-naslepo, pronásledování	průzkum-cesty (apriorně daná posloupnost)
abeceda	$I_n, B_n, F, krok_+, krok_-, toč_+, toč_-, toč, čekej$	—
ASM	pohybové schéma s autostimulací	ovládání
kritéria	distribuce zájmových oblastí	δ není vyhodnocováno, ϵ skoková funkce

Tabulka 4.7: Experimentální ASM. Tabulka shrnuje hlavní znaky robotů použitých v experimentu. Standardní funkce chování se uplatňuje pouze u robota R₁. Robot R₂ se pohybuje podle vlastního fixovaného schématu. Přesto oba dva roboty v úloze hodnotím jednotným kritériem, kvalitou pokrytí oblasti vymezené stěnami a trajektorií robota R₂.

Chování popisuje stav robota. Například bazální akce *čekej* v porovnání s abecedou předchozího chování představuje jen mezikrok mezi různými dalšími aktivitami robota. Bazální akce *čekej* v tomto experimentu nahrazuje klasifikaci «statický/dynamický objekt». Pokud je chování vybuzené (vybrané a spuštěné), průběžně stimuluje chování *sledování-zdi*. Stimulace probíhá maximálně do okamžiku, kdy *sledování-zdi* převáží nad do té doby aktivním chováním *pronásledování*.



Obrázek 4.9: Schéma chování pro ověření autostimulace. Vrcholy na grafu jsou použita chování. Hrany vyznačují stavové přechody mezi chováním. Ke každé hraně jsou připojeny identifikátory pravidel, která přechod iniciují.

4.4.1 Kritéria pro vyhodnocení experimentu

A. Diferenciální měřítko

Diferenciální kritérium má význam u sledujícího robota. V rozšíření úlohy *sledování* plně využívám i diferenciálního kritéria δ pro předchozí úlohu (4.1). Vynucená funkce v experimentu je sledování robota R_2 robotem R_1 . Není ani cílem R_1 robota R_2 chytit, ale plně akcemi robotů pokrýt oblast, v níž se pohybují.

B. Integrální měřítko

Cílový stav je opět popsán integrálním kritériem ϵ . Na celkovém chování robotů se snažím minimalizovat jízdu robota bez znalosti kontextu, pouze «naslepo». Jednotliví roboti k tomuto cílovému stavu přispívají následovně:

Robot R_1 svým chováním rozšiřuje předchozí srovnávací experiment. Do ceny (čas strávený na cestě) se mimo pronásledování (3.12) se stejnou vahou započítává i doba strávená čekáním; $c_{distr}(t_{pruzkum}) = c_{distr}(t_{pruzkum}) = 1$. Jinak $c_{distr}(t_x) = 0$. Vztah (4.2) se tím rozšiřuje o intenciobnální hodnocení takového chování T_{cekej} .

$$\epsilon_1(t) = \frac{1}{t}(T_{pruzkum} + T_{cekej}) \quad (4.3)$$

ASM v T_{cekej} načítá čas, po který volí bazální akci *čekej*.

Robot R_2 po celou dobu provádí průzkum. Po nárazu končí svoji aktivitu. Vztah (4.2) prakticky nabývá dvou hodnot: do nárazu $\epsilon_2 = 1$, po nárazu $\epsilon_2 = 0$.

V úloze není implicitní jev, jenž by centrálně vystihl přínos kolektivu. Dosud každý z robotů vyhodnocoval svým způsobem. Obě integrální hodnocení jsou na sobě závislá pouze lineárně, neboť vazba mezi roboty R_2 a R_1 je pouze jednostranná. Usuzovat o chování dvojice (a poté ho ukončit) lze na součtu kolektivního pohle. Ve smyslu obecné práce s kritérii tento součet normalizujeme na $0 \leq \epsilon \leq 1$. Tedy $\epsilon = \frac{1}{2}(\epsilon_1 + \epsilon_2)$.

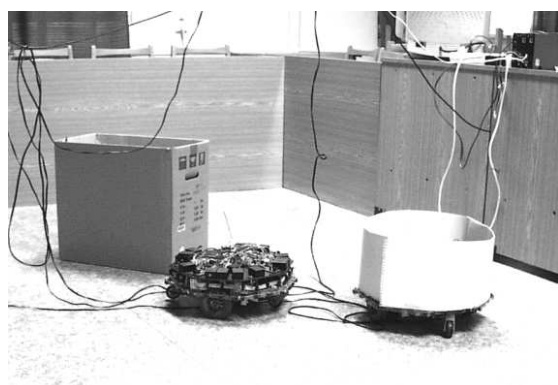
4.4.2 Průběh experimentu

V tomto experimentu jsem nevyhodnocoval pohyb naváděcího robota. Naváděcí robot fakticky jen realizoval předem určenou trajektorii. Jeho omezení na „pokrytí dostupného prostoru“ nedává robotu příliš možnost na změnu strategie podle stavu prostředí. Překážku uprostřed arény plánovitě objíždí do okamžiku, kdy detekuje náraz. Jedinou reakcí je tedy jeho chování STOP.

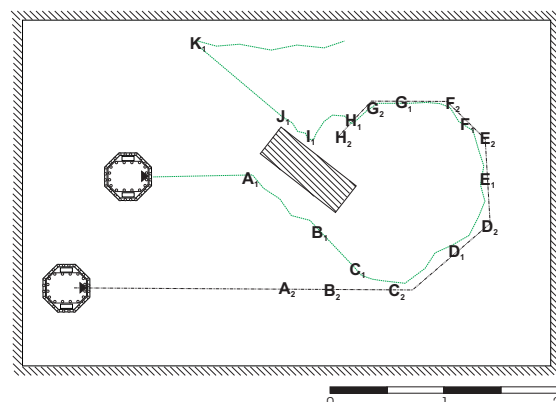
A. Robot R_1

Sledující robot může od zvenku vynucené strategie odvodit poměrně rozsáhlé chování i na základě mnohem jednodušších pravidel.

Na diferenciálním kritériu δ v obrázku 4.11(a) kromě chování pozorovaných v předchozím *srovnávacím experimentu* – průzkum-naslepo do místa A_1 , následovaný do bodu B_1 sledování-zdi. Úvodní chování již byla analyzována. Zde je zajímavé, že ačkoliv



(a) Situace



(b) Trajektorie naváděcího a sledujícího robota

Obrázek 4.10: Pohyb robotů při ověření autostimulace. V obrázku (a) je ilustrativní záběr na jízdu robotů ve skutečné aréně. Půdorys arény je vyznačen na obrázku (b), fotografie odpovídá okolí samostatně stojící překážky. Představu o velikosti arény dává přiložené měřítko — rozměry scény byly $5,5 \times 3,5$ m. Scéna je ohraničena stěnami, uvnitř scény je šrafovaná volně stojící překážka. Roboti jsou zakresleni ve startovní poloze — v pravé části arény. Černou čerchovanou čarou je vyznačena jízda *naváděcího robota* (místa s indexem $_2$), pro dobrou detekci infračidly navíc obaleného tubusem. Plnou zelenou čarou je vyznačena jízda *sledujícího robota* (místa s indexem $_1$). Pro každého z robotů je zvýrazněn jeho počáteční čelní směr.

na vyhodnocení kritických rohů překážek není robot vybaven, stále může být jeho celkové chování v kontextu dalších událostí uspokojivé.

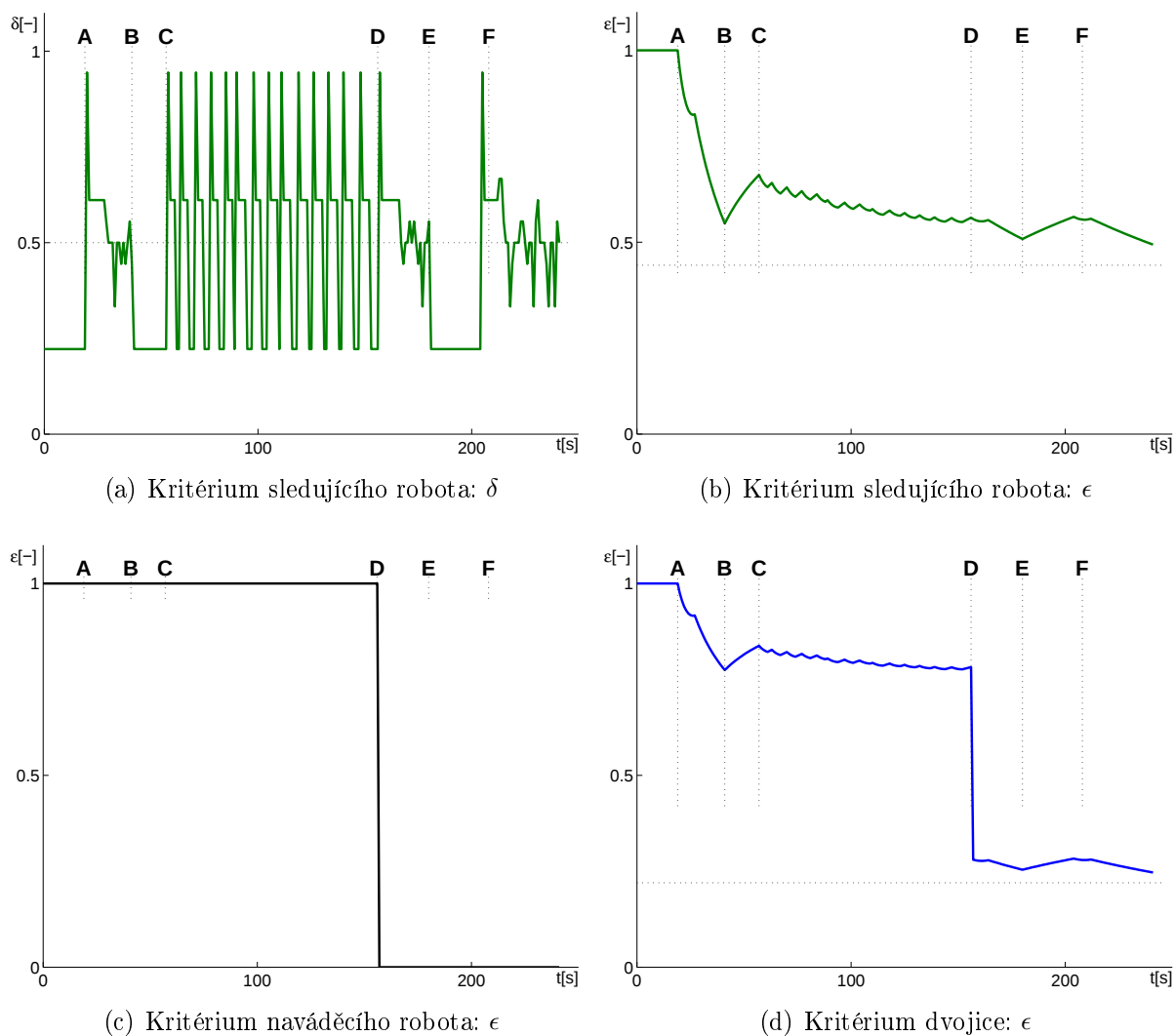
Takovou zajímavou alternativu na průběhu diferenciálního kritéria ukazuje autostimulace v úseku $C_1 \div H_1$. Zde δ hodnotí chování pronásledování. Při něm se plně projevuje *autostimulace*; hodnota diferenciálního kritéria δ periodicky narůstá. V obrázku 4.11 se právě z důvodu *autostimulace* toto chování vyznačuje vyššími výkyvy v prioritě právě použitého pravidla. Rozkmit odpovídá rozdílné rychlosti robotů. Pokud by byla odchylka minimální, pohyb *sledujícího robota* by se zcela synchronizoval s pohybem vůdčího (naváděcího) robota.

Integrálním kritériem ϵ . Sledujícího robota má smysl poměřovat i integrálním kritériem ϵ (využívaným zde převážně k hodnocení akcí Kolektivu robotů). Sledující robot může totiž v závislosti na kvalitě tohoto kritéria (definovaného vztahem 4.2) *přizpůsobit váhu svých produkčních pravidel dynamice chování vůdčího (naváděcího) robota*. To jsem ověřil i praktickými experimenty.

B. chování dvojice robotů

Na úrovni Kolektivu se pak mnohem častěji v předvedeném chování objevuje dosud apriorně neuváděný atraktor. Je dán čistě vjemem z vnějšku. V obrázku 4.10(b) vidíme, jak sledující robot opisuje trajektorii vůdčího (naváděcího) robota, ačkoliv jeho prvotní schopností a cílem bylo pouze sledování obvodu překážek. K tomuto obrázku jsou vztaženy následující průběhy.

Integrálním měřítko ϵ . Dosažené kvality celé dvojice v integrálním kritériu ϵ vystihuje obrázek 4.11(b). Hodnota souhrnného integrálního kritéria ϵ je vytažena plnou čarou.



Obrázek 4.11: Kvality řídicího systému s autostimulací. Měřítkem všech grafů je doba běhu algoritmu sledujícího robota. Graf (a) pro tento čas v diferenciálním kritériu δ hodnotí reakci na překážky registrované během jízdy obou robotů. Vynášená hodnota hovoří o tom, jak často roboti používají chování sledování-zdi nebo STOP. Při hodnotě $\epsilon = 1$ je roboti nikdy nepoužili, při nulové je oba používají výlučně. V průbězích je plnou čarou ohodnocena jízda pomocí obou robotů. Protože průběh integrálního kritéria ϵ vůči robota představuje triviální jednotkový pokles, graf doplňují pouze o ohodnocení sledujícího robota — vyznačeno čerchovanou čarou.

Vynášená hodnota hovoří o tom, jak dalece se aktuální chování vzdaluje jízdě v bezprostřední vzdálenosti překážky. Při vysokých záporných hodnotách robot překážku neregistruje, při vysokých kladných do ní naráží. Pilovitý průběh přibližuje nárůst tohoto hodnocení v okamžicích, kdy se robot snaží rozlišit, zda registruje statickou překážku či pohybující se objekt.

Pro ilustraci je ve stejném kontextu čerchovanou čarou zakreslena hodnota dílčího příspěvku ϵ_0 vůdčího (naváděcího) robota.

Na delším časovém horizontu se projevuje i autostimulace. Sledující robot díky autostimulaci ve svém rozhodování začíná rozlišovat mezi statickými o dynamickými předměty. Dosud robot detekoval pouze objekt bez informace o jeho dynamice. S použitím autostimulace už rozlišuje mezi vnímáním vůdčího (*naváděcího*) robota anebo o statické překážce. Autostimulace probíhá výlučně ve stavu *čekej*.

4.4.3 Výsledek experimentálního ověření

Hodnocení se týká převážně *sledujícího robota* (R_1). Jeho algoritmus se principiálně mění, začíná používat intencionální model. *S ním rozšiřuje velmi jednoduché chování pro sledování stěn dále do prostoru*. Velikost projeté oblasti už závisí na okolních vlivech, zde simulovaných naváděcím robotem (R_1). Ve spojení více robotů se také jinak uplatňuje integrální kritérium ϵ . Neurčuje pouze okamžik zastavení. Stejný způsob akumulace také vynucuje u robota (R_1) synchronizaci s naváděcím robotem (R_2).

Shrnu-li:

V této **podkapitole 4.4** jsem experimentálně doložil *doplňkové možnosti autostimulace*. Zde se jedním mechanismem vyhodnocují akce spolu s podněty. *Průběh experimentu dokumentuje důsledek této minimální reprezentace*: reprezentace částečně zastoupí chybějící senzorické podněty. Tak zatímco robot (jeho ASM) může ze senzorů nanejvýše přepínání mezi chováním v přítomnosti překážky či ve volném prostoru, s autostimulací lze dále rozlišit jeho chování v okolí dynamické překážky (tehdy je robot v očekávání změny).

V prvním oddíle jsou vyjmenované základní nastavení dvou robotů použitých v tomto experimentu. *Autostimulaci ověřuji pouze u sledujícího robota*. Ten autostimulací klasifikuje, zda detekoval statický nebo dynamický objekt. Vůdčí (*naváděcí*) robot je řízený deterministicky.

Ve druhém oddíle jsem interpretoval dříve uváděná kritéria za podmínek sledování pohyblivého předmětu. Diferenciální kritérium sledujícího robota je opět odvozeno od priority akce. Odráží tak postupný nárůst podnětů, aby se robot k překážce choval stejně jako při sledování zdi. Zatlumený nárůst naopak signalizuje chování, kdy se sledující robot synchronizuje s pohybem vůdčího robota. Integrální kritérium pak hodnotí poměr obou typů sledování a ostatních nevyžádaných reakcí.

Ve třetím oddíle se ukazuje přínos autostimulace. Zvláště v integrálním měřítku se s apetencí dociluje vyváženější užívání kombinace prioritních chování pro stabilní sledování. Eliminuje se podíl náhodného vyhledávání cílových oblastí.

V tomto souhrnném čtvrtém oddíle zdůrazňuji nový časový kontext uvažovaných podnětů jako *hlavní výhodu autostimulace*. Robot tak může z dvoustavových detektorů odvodit nejen existenci objektu, ale také i pohyb objektu. Takto lze rozlišit objekty, které se pohybují výrazně pomaleji oproti pohybu robota. Dále je v rozhodování robota implicitně zahrnut přibližný směr pohybu takového objektu; robot se za další detekci vydává ve směru, kde byla předchozí detekce. Zde uvedený experiment a jeho závěry jsem publikoval [Svatoš, 2000a] a diskutoval i na vyzvané přednášce [Svatoš, 2001] v Opavě na katedře Prof. Kelemena, jehož ideu založenou eko-gramatických systémech zde rozvíjím.

4.5 Experimentální ověření interakcí EG systému na robotech

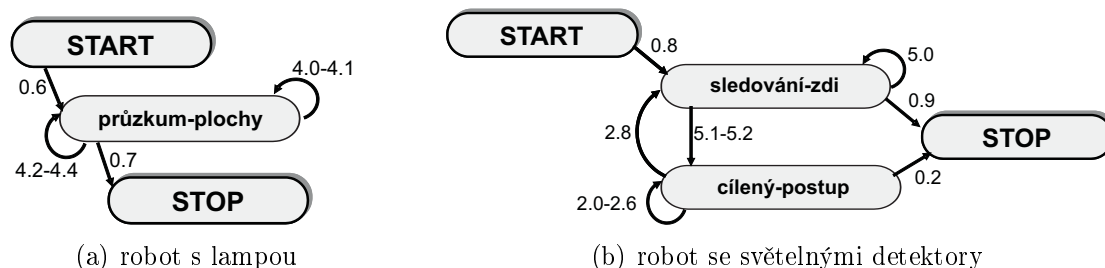
Důsledky interakcí v eko-gramatickém systému jsem nakonec ověřoval na příkladu shlukování dvou robotů s rozdílnými možnostmi percepce okolního prostředí. Na této experimentální úloze mohu ukázat:

- Funkci reaktivního agenta v eko-gramatickém systému; níže robot R_1 .
- Funkci intencionálního agenta v eko-gramatickém systému; níže robot R_2 .
- Nepřímou koordinaci mezi roboty z definice 3.4.1, realizovanou implicitní komunikací.
- Autostimulaci pro rovnoměrné směřování pohybu intencionálního agenta ve vymezené aréně.

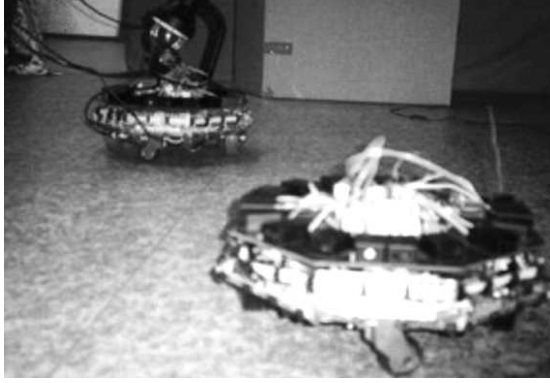
		R_1 :	R_2 :
chování		průzkum-plochy	cílený-postup, sledování-zdi, průzkum-naslepo
abeceda		$I_n, B_n, F, krok_+, krok_-, toč_+, toč_-$	$I_n, B_n, F, \mathbb{C}, krok_+, krok_-, toč_+, toč_-, krok$
ASM		pohybové schéma s autostimulací	parametrizovaná reakční báze
aritm.výrazy $E(\Sigma)$		—	$\alpha = \arctan \frac{C_s - C_j}{C_v - C_z}$
percepce	φ	dvoustavové I_n, B_n	HW senzor světla s integrací; $\mathbb{C}_x = \Sigma C_x$
aktivita	ψ	změna polohy světla	změna polohy robota
kriteria		koordinovaný pohyb; $\omega_{R_1}=1, \omega_{R_2}=0$	koordinovaný pohyb; $\omega_{R_2}=1, \omega_{R_1}=0$

Tabulka 4.8: Experimentální části eko-gramatického systému. V experimentu uvažuji kromě chování robotů také praktickou interpretaci definice eko-gramatického systému. Proto u robota R_1 ukazuji, jak se projeví jeho aktivita (zobrazení ψ). U robota R_2 poukazuji na jeho specifické nakládání se senzory (zobrazení φ). Větší význam než proximální detekci robot dává intenzitě světla. Vyhodnocuje ji předepsaným aritmetickým výrazem $E(\Sigma)$ ve čtyřech bázových směrech $s-j-v-z$.

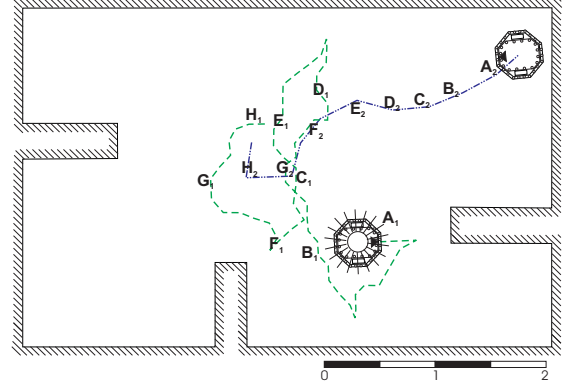
V rámci možností jsem do úlohy zapojil dva laboratorní roboty. Jejich různorodé strategie a funkce jsou uvedeny v tabulce 4.8. Robot R_1 se v reakci na detekci překážky pohybuje opačným směrem. Robot R_1 je v experimentální aréně snadno identifikovatelný díky lampě, kterou veze. Robot R_2 s robotem R_1 koordinuje nepřímo. Registruje jeho polohu světelnými senzory. Do určeného směru gradientu pak zaměřuje své chování cílený-postup. Způsob odvození jsem podrobně přiblížil na příkladu kapitoly 3.3.3.



Obrázek 4.12: Schéma chování v úloze světelné navigace. Vrcholy na grafu jsou použitá chování, hranami jsou stavové přechody mezi nimi. Ke každé hraně jsou připojeny identifikátory pravidel stavového přechodu. Okamžik START a STOP je pro oba roboty identický.



(a) Situace



(b) Trajektorie přiblížení

Obrázek 4.13: Interakce experimentálního eko-gramatického systému. V čele záběru na obrázku (a) je robot R_2 , který sleduje světlo emitované robotem R_1 . Půdorys arény je vyznačen na obrázku (b), fotografie odpovídá okamžiku C. Představu o velikosti arény dává přiložené měřítko — rozměry scény byly $5,5 \times 3,5$ m. Scéna je ohraničena stěnami. Uvnitř scény jsou šrafovány volně stojící překážky. Roboti jsou zakresleni ve startovní poloze – robot R_1 v okolí volně stojících překážek, robot R_2 ve volné pravé části arény. Od těchto poloh jsou zakresleny trajektorie pohybu robota R_1 (zelenou čárkovanou čarou) a robota R_2 (modrou čerchovanou čarou). U robotů je zvýrazněna jejich počáteční orientace.

V algoritmu robota R_2 na obrázku 4.12(b) je ještě naznačená pomocná strategie. V případě nejednoznačného směru gradientu robot může použít své původní chování, sledování-zdi. Při nasazení obou robotů je ovšem po celou dobu experimentu sledování-zdi převrstveno cíleným-postupem. Detailní popis užitých pravidel chování obou robotů udávám v příloze B.5. Dosaženou trajektorii přiblížení ilustruje obrázek 4.13(b).

4.5.1 Kritéria pro vyhodnocení experimentu

Svou podstatou se experiment blíží k simulacím s pachovou stopou. Zde při párování robotů ve vymezené oblasti ošetřují jejich jednostranný vztah. Smysl předvedené kolektivní aktivity mohou vidět v navigaci robota R_2 do oblasti, vycházející z intencionálních rozhodnutí robota R_1 . Oproti distribuci zájmových oblastí se tak zaměřují pouze na současný průnik záběru dvou robotů.

A. Diferenciální měřítko

Robot R_1 V chování tohoto robota nehraje vazba na druhého robota žádnou váhu, ve smyslu (3.13) $\omega_{R_2} = 0$. Robot v diferenciálním kritériu δ pouze porovnává skutečný stav oproti svému odhadu nejlepšího mezikroku. Robot hodnotí souslednost kroků v relacích *konfliktního postupu akce* $\nRightarrow$ *kontra-akce* (zde zřejmě $krok_+ \nRightarrow krok_-$, $toč_+ \nRightarrow toč_-$ a naopak) a *změny pohybové báze* (přechod mezi chováním $krok_{xx} \curvearrowright toč_{yy}$ a naopak). Soulad bazální akce s předchozí aktivitou určují iterativně pro každé nové rozhodnutí $F_{baz-akce}$

$$F_{baz-akce}(t+1) = \begin{cases} 0 & \text{pokud } akce-ASM \nRightarrow baz-akce, \\ F_{baz-akce}(t) + 1 & \text{jinak.} \end{cases} \quad (4.4)$$

Soulad $F_{baz-akce}$ prakticky hodnotí, jak se bazální akce hodnotí po započtení autostimulaci. Proto tento soulad užívám i jako praktickou interpretaci míry mezi odhadovaným nejlepším dosažitelným stavem a skutečností $\|\hat{C}_{R_1} - C_{R_1}\|$ z (3.13):

$$\Gamma(t) = \begin{cases} 0 & \text{pro } akce-ASM = minula-akce, \\ 1 & \text{pro } minula-akce \not\Rightarrow akce-ASM. \\ \frac{-F_{lepsi-zmena}(t)}{F_{lepsi-zmena}(t) + F_{horsi-zmena}(t) + \sum F(t)} & \text{pro } akce-ASM \curvearrowright \{lepsi-zmena, horsi-zmena\}, \\ \frac{F_{horsi-zmena}(t)}{F_{lepsi-zmena}(t) + F_{horsi-zmena}(t) + \sum F(t)} & \text{přičemž } F_{lepsi-zmena}(t) > F_{horsi-zmena}(t). \end{cases} \quad (4.5)$$

Robot R_2 Hodnotu diferenciálního kritéria δ určuje naakumulovaná intenzita světla. Při navigaci chování cílený-pohyb minimalizuje odchylku od největšího gradientu. Pro normalizaci zde užívám funkci $\sin$. S touto odchylkou v diferenciálním kritériu δ pracuji stejně jako s odhadem

$$\Gamma(t) = \frac{\alpha_z(t) - \bar{\alpha}(t)}{\sqrt{(\alpha_s(t) - \bar{\alpha}(t))^2 + (\alpha_z(t) - \bar{\alpha}(t))^2}}, \quad \bar{\alpha} = \frac{1}{4}(\alpha_s + \alpha_j + \alpha_v + \alpha_z) \quad (4.6)$$

kde $\bar{\alpha}$ udává intenzitu světla v těžišti robota, odvozenou z měřených hodnot.

Za předpokladů (4.5), respektive (4.6), pak oba dva roboti filtrují diferenciální kritérium δ analogicky k výchozímu předpisu (4.1). Ve významu tabulky 4.6 je δ vztaženo vůči $W_{mod} = 0,563$:

$$\delta(t) = q \left(0,5 + \frac{\Gamma(t)}{2} \right) + (1 - q) \cdot \delta(t - 1), \quad \delta(0) = 0, \quad (4.7)$$

B. Integrální měřítko

Integrální předvedeného chování je podle (3.12) postaveno na nákladové stránce přesunu robotů.

Robot R_1 Náklady na pohyb intencionálního robota narůstají při opakování protichůdných akcí. Jejich cenu lze vyčíst s ohledem na předchozí rozhodnutí. Ve smyslu (4.4) o nákladech na aktuální rozhodnutí hovoří souslednost protichůdných kroků $F_{kontrakce}$

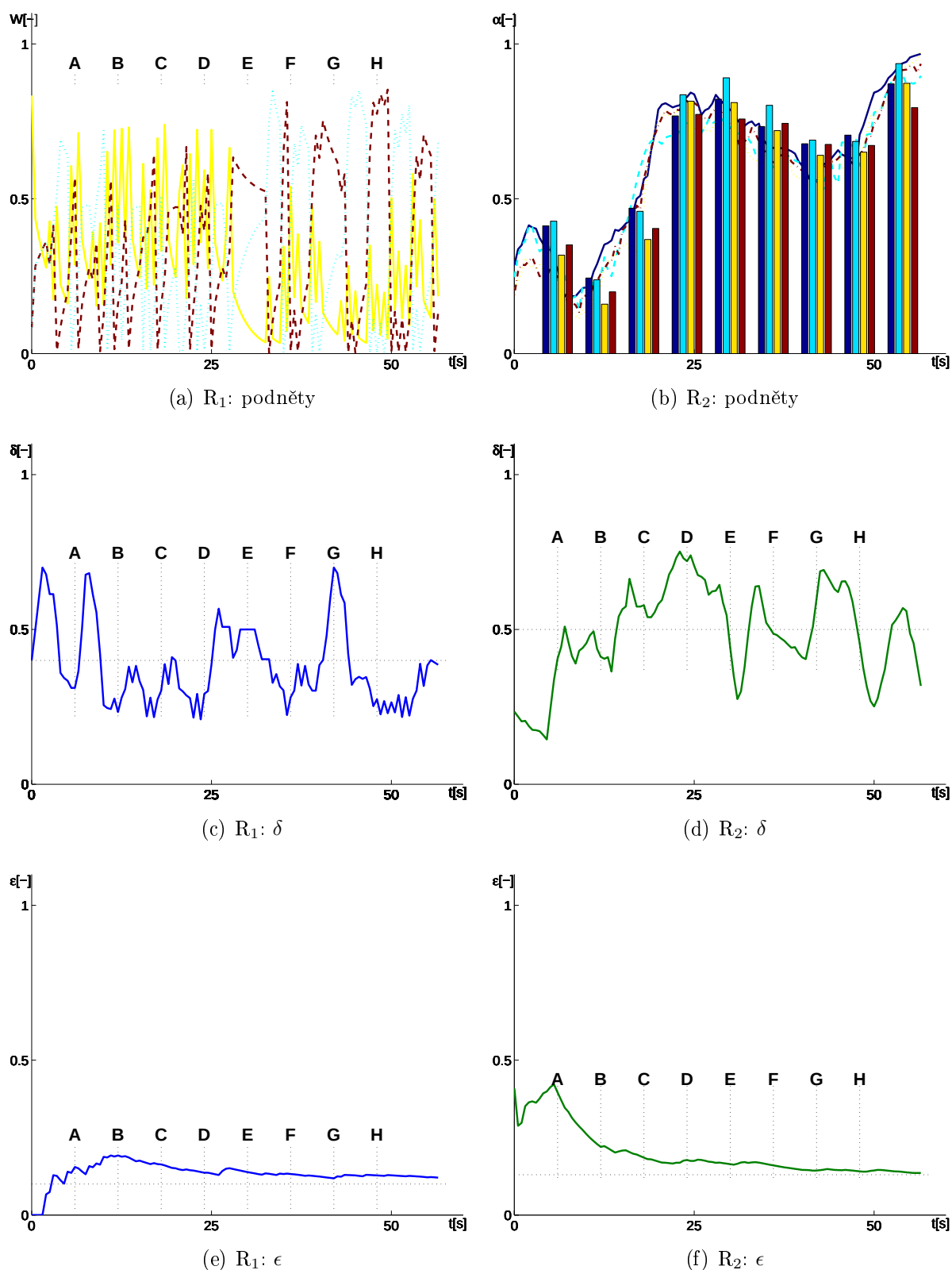
$$c_{coord}(t) = \frac{F_{kontra-akce}(t)}{\sum_a F_a(t)}, \quad kontrakce \not\Rightarrow akce-ASM. \quad (4.8)$$

Robot R_2 Také v ceně jeho reakcí se odráží práce nad rámec optimální cesty. Prakticky ji vyjadřují délkou kroků, které probíhají mimo požadovaný směr. Po každém kroku oceňuji a dále minimalizuji detekovanou kvadratickou odchylku úhlu (4.6).

$$c_{coord}(t) = \Gamma^2(t) \quad (4.9)$$

S předpoklady (4.8), respektive (4.9) si každý z robotů může napočítat své integrální kritérium ϵ ze dvou čítačů, ceny c_{coord} a počtu rozhodnutí:

$$\epsilon(t) = \frac{1}{t} \sum_{\tau=0}^t c_{coord}(\tau) \quad (4.10)$$



Obrázek 4.14: Kvalita chování v úloze světelné navigace. Měřítkem všech grafů jednotný čas robotů zapojených do kolektivu. Písmeny jsou značené okamžiky rozhodování robota R_2 . Grafy (a)/(c)/(e) hodnotí ASM robota R_2 . Jeho podnětem ke změně je váha bazálních akcí. Na grafu (a) je vynesena váha pohybu ve směru s žlutou plnou čarou, směr v červenou čárkovanou čarou a směr z modrou čerchovanou. Grafy (b)/(d)/(f) obdobně hodnotí reakci robota R_2 na intenzitu světla. Křivky na obrázku (b) zobrazují aktuální měřené hodnoty světla ve směrech s (tmavě modrá), j (světle modrá), v (žlutá) a z (červená). Sloupcový graf na témže obrázku ukazuje akumulované hodnoty světla. Z nich robot R_2 určuje další směr postupu.

Integrální kritérium ϵ v této úloze ovšem nedefinuje cílový stav chování. Roboti se k sobě přibližují, a končí v okamžiku, kdy robot R_2 detekuje světlo o intenzitě vyšší než $\alpha_{min} = 95\%$. Tato prahová vychází z doporučení [Starý, 2000].

4.5.2 Průběh experimentu

A. Diferenciální měřítko

Robot R_1 Průběh diferenciálního kritéria δ na obrázku 4.14(c) ukazuje na výběrový mechanismus. Autostimulace plně odstranila průzkum-naslepo. Nezabýváme se ovšem ošetření kritických nárazů. Takové okamžiky jsou patrné jako výrazné změny kritéria (kritérium samo o sobě je totiž zaměřeno na pohyb ve volném prostoru). Diferenciální kritérium $\delta = W_{mod} = 0,5$ naopak hovoří o pokračování akce v nastoleném směru. Z akcí pro změnu pohybu se volí akce s lepší autostimulací, a tedy i sousledností; $\delta < W_{mod}$. V krajním případě je pak volena zcela protichůdná akce, $\delta = 1$.

Robot R_2 Reálná odchylka v diferenciálním kritériu δ na obrázku 4.14(d) charakterizuje soulad s předchozím pohybem robota R_1 . S menší odchylkou robot R_2 vhodně kopíruje střední polohu robota R_1 .

B. Integrální měřítko

Robot R_1 Na průběhu integrálního kritéria ϵ na obrázku 4.14(e) se ukazuje opačný přístup v hodnocení oproti srovnávacímu experimentu. Protože iniciálně jsou překážky mimo dosah robota R_1 , jeho integrální hodnocení je zpočátku velmi příznivé. Aby robot dosáhl cílového stavu, může se v integrálním hodnocení pouze zhoršit (nežádoucími kontakty s překážkou).

Robot R_2 Z integrálního hodnocení postupu reaktivního robota na obrázku 4.14(f) je patrný časový odstup odpovědi od vlastního podnětu. ASM robota R_2 nemá prediktivní algoritmus, který by eliminoval odchylku mezi reálným rozložením světla a průběžně akumulovanou hodnotou. V integrální kvalitě se chování robota R_2 ustálí na poměrně vysoké hodnotě ϵ . Hodnota odpovídá zpoždění robota R_1 oproti robotu R_2 .

4.5.3 Výsledek experimentálního ověření

Experiment ukázal jednu z možností, ***jak implicitně definovat cíl***. V cílovém stavu Kolektivu se stochasticky projevují různé vlastnosti jednotlivých elementů celého ekologického systému. Úloha je založena na implicitní komunikaci. Využívá se zde různé interpretace jedinečné události jednotlivými roboty.

Oproti srovnávacímu experimentu zde stojí za povšimnutí

- **interní paměť robota R_1** . paměťová schopnost světelných senzorů ve výsledku dává stejné možnosti, jako by integrální funkce byla přímo vlastností prostředí robota.
- Stav problému se hodně odráží na pohybu robotů. K tomuto dynamickému cíli roboti vzájemně konvergují mnohem rychleji než ve srovnávacím experimentu. Aktivnější z robotů se navíc může přesněji přiblížit k poloze integrálněji vnímaného cíle.

- **Atraktivní a repulsivní reakce na rozprostřené vlastnosti prostředí.** Díky implicitní komunikaci mezi roboty se tyto reakce zavedené do originálního pohybového schématu ještě zesilují. Změny „intencionálního“ robota nejsou dané pouze stavem prostředí, ale i předchozím postupem robota. „Reaktivní robot“ každou takovou odchylku přenáší i do svého chování.

Jestliže se roboti i v takto omezeném kolektivu zaměřují na řešení podobného problému, jejich řešení zahrnuje více znaků řešeného problému. Zajímavé chování například vznikne po několika nárazech na stěny a odpovídajících reakcích robota R_2 . S užitými pravidly se ovšem celý kolektiv přesouvá do dosud nevymezeného středu plochy (ve výrazně početnějším společenství bych mohl hovořit o *samovolně vznikající emergenci*). U použitých dvou robotů navíc nové poznatky vedly na úplné převrstvení *individuálního průzkumu*, chování sledování-zdi.

Shrnu-li:

V této podkapitole 4.5 jsem ukázal *na dvou robotech možnost reakce na vnější podmínky*. Každý z robotů v experimentu stylizuje odlišný zdroj změn v eko-gramatickém systému. Intencionální robot v něm zdůrazňuje podstatné znaky okolí. Zdůrazněné znaky prostředí jsou předmět nepřímé komunikace mezi roboty. Zdůrazněné vlastnosti dále zpracovává reaktivní robot.

V prvním oddíle jsem vymezil role eko-gramatického systému, K nim jsem uvedl metodu, jak bude každý z robotů demonstrovat přidělenou roli. Algoritmus reaktivního robota jsem v této úloze postavil na přesném odvození akce z detekovaného světelného gradientu. Pohyb intencionálního robota je vymezen obvodem arény. Při nárazu na obvod vůdčí (*naváděcí*) robot reaguje deterministicky. Ve volném prostoru podněty k jízdě a otáčení vychází z autostimulace.

Ve druhém oddíle uvádím interpretaci dříve zaváděného integrálního a diferenciálního kritéria v úloze párování. Reaktivní robot obě kritéria odvozuje ze světelného vjemu. Intencionální robot měří obdobná kritéria z vnitřních podnětů. V obou případech diferenciální kvalita zdůrazňuje tendence v pohybu světelného zdroje. Ohodnocuje změnu směru v po sobě následujících iteracích. Integrální kritérium pak započítává náklady na pohyb robotů. Nákladem reaktivního robota je každá odchylka od odhadovaného cílového směru, náklady na pohyb intencionálního robota přímo plynou z nevhodných protichůdných rozhodnutí.

Ve třetím oddíle těmito kritérii posuzuji dosažený pohyb sledujícího robota. Pohyb sledujícího robota emerguje z předpokladů, obsažených v pravidlech prostředí. V experimentu volím dynamiku řádově rychlejší oproti pohybu sledujícího robota. Do výsledného chování sledujícího robota se tak integruje mnohem více dílčích změn prostředí. Robot se tímto mechanismem i bez znalosti obvodu arény asymptoticky přibližuje do středu experimentální arény.

V tomto souhrnném čtvrtém oddíle se pak věnuji dalším rozšířením užitého principu implicitní komunikace v kontextu dalších vyhraněných relací. V experimentu se robot atraktivně směřuje ke zdroji světla, přičemž repulsivní trendy vymezují atraktivní oblast prostředí mimo obvod arény. Při použití stejného principu v mnohem širším Kolektivu robotů se dá odvodit stochastické zapojení obou podmínek. Stejným vjemem společných příznaků prostředí tak lze docílit přirozeného shlukování robotů ve volném prostranství experimentální scény.

4.6 Shrnutí a důsledky realizovaných experimentů s roboty

Dokumentované experimenty prokazují ***uplatnění navrhované architektury v behaviorálním řízení mobilního robota***. Část experimentů přímo ukazuje, jaké místo v navržené architektuře zaujímá koordinace s dalším robotem v Kolektivu. Všechny experimenty také dokládají, jak lze vzájemně kombinovat různé úrovně předzpracování podnětů před vlastní volbou akce (kapitola 4.3), a jak se zvolené akce mohou vzájemně doplňovat v kolektivním řešení (kapitoly 4.4 a 4.5). Tématicky proto v řešení zdůrazňuji:

- A. *Chování robota jednotlivce ve svém prostředí,*
- B. Příspěvek robota – jednotlivce ke *Kolektivnímu řešení,*
- C. Porovnání dosaženého výsledku s *jinými pracemi v oblasti Umělého života.*

Syntéza funkcí navrženého systému se projevila už na průběhu experimentů. Výsledná funkce systému je pak dána úrovněmi předzpracování. Rostoucí možnosti jsou patrné i ve sledu, jakým jsem postupně naplňoval definici eko-gramatického systému. Přímo v experimentech se například s každou dodanou úrovní rozšiřuje i počet (příjemně) zpracovaných typových situací.

Možnosti postupného skládání a doplňování funkcí jsou patrné už v etapách, jak se do experimentů zapojovaly jednotlivé modely navrhované architektury. Jimi jsem postupně naplňoval definici eko-gramatického systému. Přímo v experimentech se pak s každou dodanou úrovní rozšiřuje i počet (příjemně) zpracovaných typových situací.

A. Individuální pohled – *chování robota jednotlivce ve svém prostředí*

V celé práci sleduji odhad, jak zvolená akce přiblíží stav systému k udanému cíli. V práci tento odhad odráží pravidla reaktivního modelu, a později též ekvivalentní gramatická pravidla pro intencionální model. Krom reakcí na aktuální podmínky se tak do mechanismu výběru akce dostává odvození následků v závislosti na předchozích stavech. K odhadu akcí proto kombinuji (a) z *reaktivního modelu* aktuální intenzitu veličin, (b) z *intencionálního modelu* vnitřní stimuly ke spuštění bazální akce.

B. Pohled z hlediska Kolektivu – *příspěvek robota jednotlivce k výběru akce celku*

V kolektivním pohledu na problém se dále omezují na *implicitní komunikaci*. Vzájemné ovlivňování robotů přes prostředí dokládám ma analogii mezi odvozením aktivity robotů a eko-gramatickým systémem. Projev eko-gramatického systému v navržené hybridní architektuře hodnotím jako jeden z příkladů propojení lingvistických předpokladů s praxí.

C. Porovnání s jinými pracemi v oblasti Umělého života

Výsledky (dosažené mnou zejména na začátku 21. století) tuto práci opravňují zařadit do poměrně úzké skupiny ***realizací lingvistických principů na reálných robotech***. V teoretické rovině proto nejprve připomenu některé tehdy současně s mým výzkumem realizované *práce odborníků z matematické lingvistiky*.

V praktické rovině srovnání neomezují na řešení dílčích problémových situací. Pro srovnání uvádím současné postupy Behaviorální robotiky, analogické mému návrhu. Na nich dokládám přijatelnost zvoleného přístupu. Ke srovnání se například nabízí behaviorálně pojaté architektury řídicího systému v soutěži Robotického fotbalu. Simulace a praktické herní situace mohou ukázat na řešení německého národního týmu [Röfer a kol., 2007]. Srovnávám je v historickém nadhledu. Do kontextu současných trendů dávám své práce z doby, kdy **Problémy** z úvodu této práce byly více než inspirativní i pro mě.

4.6.1 Individuální hodnocení experimentů

Celé experimentální ověření je záměrně postaveno na velmi jednoduchých robotech. ASM i jejich omezenou sadu podnětů může zpracovat s dostatečnou generativní silou. Z hrubých vjemů zde ASM vždy nabízí přijatelné řešení pro jednotlivce. Robot při individuální volbě do základního ASM výhodně kombinuje

- inkrementálně dodávané další úrovně předzpracování
- a vhodné interpretace podnětů.

Celý návrh řídicího systému jsem pojal inkrementálně. Každá nová úroveň v řídicím systému generuje události, které řídí další chod systému. První filtrace probíhá v reaktivním modelu, dodává vnímaným znakům hrubou sémantiku situace. V intencionálním modelu pak následuje syntéza předzpracovaných vnějších událostí spolu s vnitřními podněty.

A. Inkrementální vytváření řídicího systému

Strukturu řídicího systému definuji postupně. Po způsobu Umělého života zapojuji nové moduly pomocí implicitních vazeb. S nově aktivovanými vazbami přecházím od kvantitativního řešení ke kvalitativnímu. Kvalitativní vývoj návrhu je dobře vidět na popsaném sledu experimentů. Každý nový modul rozšiřuje reprezentaci téhož vymezeného vjemu do nové úrovně:

1. Základní *čisté reaktivní ASM* parametrizují. Formálně doplňují pravidla pro odvození gradientu.
2. Nad reaktivní ASM nadřazují *hybridní zpracování*. Na úrovni intencionálního modelu nejprve reprezentují závěry pravidlového modulu. Tuto znalost reprezentují v parametrech pohybového schématu [Nahodil, Svatoš, 2001]. V jednotné metrice pohybového schématu poté odvozují nejvýhodnější postup.
3. Nad hybridním mechanismem výběru akce dále funguje modul *adaptace*. V něm metodou pokusu a omylu odvozují nová pravidla pro ošetření nepokrytých situací. Nová pravidla po ověření zpětně formují reprezentaci problémové situace na nižších úrovních.

Tento princip vrstvení chování byl původně publikovaný v [Petrus, Svatoš, 2000]. Stejná myšlenka je dnes použita v [Sudolský, 2006]. Vyšší kvalitu po zapojení těchto modulů mohu stejně doložit na výsledku mých dřívějších výsledku experimentů [Svatoš, 2000c].

B. Interpretace podnětů

Řídicí systém výhodně kombinuje evidentní sensorické vjemy spolu s doplňkovými vnitřními stavy. Rozšiřuje tak mechanismus výběru akce o odhady dalšího postupu, souslednosti uvažovaných akcí či aktuální hodnoty kritérií. Ve výběru se dávají do souvislosti vnitřní stavy a cíle. V architektuře tomu předchází předzpracování sensorických vjemů a doplňkových vnitřních stavů:

1. V případě reaktivního výběru se mi osvědčila *akumulace podnětů*. V experimentech kapitoly 4.5 jsem vyzkoušel nejčastější způsob, jakým může jedinec dosáhnout maxima. Ekvivalentem k hledání nejintenzivnější pachové stopy bylo hledání maxima světla.
2. Autostimulace pro mě představuje dodatečný podnět k *vyváznutí z klidového stavu*. Užívám ji jako ekvivalent žíhání v klasických algoritmech hledání maxima. Autostimulací se do mechanismu výběru zařazují proaktivní chování typu «poznávání».
3. Autostimulace dává *nový rozměr senzorickému vnímání*. Představuje ekvivalent k pohybu očí či dotýkání. Na experimentech kapitoly 4.4 ukazují, jak autostimulace napomáhá klasifikaci.
4. Stejně jako u vnějších podnětů dává autostimulace *časový rozměr vnitřním podnětům*. Specifickým případem je v práci často používaná negativní autostimulace. Tímto mechanismem realizují zapomínání jednorázových podnětů, kritériálních hodnot. Taktéž dodávám impulsy ke změně chování při detekci statických překážek.
5. Samostatný podnět také vychází z kritických hodnot použitého *kritéria*. Prahovými hodnotami kritérií lze aktuální chování zastavit. Buď tak končí celá úloha, nebo se přechází na chování typu «adaptace».
6. Systém je dále otevřen dalším podnětům z explicitní komunikace.

Všechny tyto podněty formálně shrnuji v rámci jednoho eko-gramatického systému. Z něj v případě nejasností odhaduji vnitřní stavy a chování jednotlivce [Nahodil, Svatoš, 2002]. Z jednotlivých vnitřních stavů si také v eko-gramatickém systému mohu udělat představu o závěru společného kolektivního řešení.

4.6.2 Hodnocení experimentů na úrovni Kolektivu

Použití řídicího systému v Kolektivu hodnotím na Typových Úlohách, kde lze cílový stav vystihnout na poloze robota vůči ostatním. Zajímám se proto o funkci řídicího systému při (a) distribuci zájmových oblastí a v (b) úlohách koordinovaného pohybu. Zaměřuji se u nich na:

- interpretaci událostí, generovaných jiným robotem
- a doplňkové role robotů při řešení společného problému.

Uplatnění navrženého řídicího systému v Typových Úlohách dokládám na průběhu experimentů v kapitolách (a) 4.4, respektive (b) 4.5. K experimentům se krátce vracím v následujících dvou odstavcích.

A. Interpretace událostí jiného robota

V experimentech se omezují na implicitní komunikaci významných znaků problémové situace. Na takových znacích zakládám koordinaci robotů. Oproti výběru akcí jednotlivce nepoužívám převrstvování aktivit robotů. Nechávám průchod přirozené syntéze chování zapojených robotů:

V experimentech se omezují na implicitní komunikaci významných znaků problémové situace. Na takových znacích zakládám koordinaci robotů. Oproti výběru akcí jednotlivce nepoužívám převrstvování aktivit robotů. Nechávám průchod přirozené syntéze chování zapojených robotů:

1. Implicitní komunikace bez dodatečných úprav dodává do mechanismu výběru další signály. Dodané signály buď potvrzují již detekovanou situaci, nebo také signalizují nové okolnosti, dosud nepodchycené senzory robota.
2. Koordinaci ukazují na sledování mezi dvěma roboty. Analogicky k biologickým zřetěžením v hejnech lze i na trajektorii robotů v experimentu kapitoly 4.4 pozorovat kombinaci znalostí obou robotů. Sledující robot zde používá (a) relaci na předchůdce ve stopě, který má znalost o (b) dalším postupu ve stopě.
3. Na úrovni jednotlivých robotů preferují zdůraznění hlavních znaků cílové situace. Ve výsledku se roboti chovají jako přírodního hejno, koncentrované do míst s vyšší hustotou aktivně dodávaných značek – často se jako takové médium například užívá pachová stopa.
4. Událost vyhodnocuji vždy relativně – vazbou mezi lokálním vjemem a bazální akcí robota. Impulsy, které tak budí bazální akce robota, jsou analogické funkci místních buněk v mozku hlodavců [Telenský a kol., 2006].

Princip sledování společného světelného zdroje (ale i tepelného, aj.) se stále používá pro navigaci roje robotů. *Společnou myšlenku* vidím například mezi mými pokusy se světelnou navigací jednoho robota [Svatoš a kol., 2001] a vytvářením obdobné trajektorie při postupu robotů v jednoduchém bludišti [Christensen, Dorigo, 2006]. Nicméně Christensen už při navigaci strukturuje chování do neuronové sítě. I on, stejně jako kdysi já, zavádí do svého návrhu prvky adaptivity.

Zmíněná *fototaxe* se dále dnes využívá i v realizačních úlohách, vzdálených mému původnímu záměru – robotickému fotbalu. Fototaxe, jak ji používá Ieropoulos je tak v současnosti důležitým chováním *při simulaci a ověřování energeticky autonomních robotů* [Ieropoulos a kol., 2005].

B. Role robotů při řešení společného problému

Pro porovnání přínosu jednotlivých robotů jsem zvolil úlohy přirozeně dekomponovatelné na jednotlivé roboty. Výhoda zvoleného oportunistického zapojení robotů je následující:

1. Implicitní komunikace se v experimentech projevuje jednostranně.
2. Od dané role přirozeně odvozují měřítko integrálního kritéria.
3. Mechanismus komunikace není adresný. Roboti přijímají implicitní signály pouze v tom případě, že znají způsob jejich vyhodnocení. Řešení úlohy tak lze snadno rozšířit na více robotů.

Ve své práci [Svatoš, 2001], kterou jsem dále rozšířil [Svatoš, Nahodil, 2008a], dále uvažuji o kombinaci robotů s různými rolami do pohybové formace. Z nedávných řešení [Vail, Veloso, 2003] obdobně jako já svými vymezeným **Úkoly** přistupuje k robotickému

fotbalu i Veil. V obou pracech je to *globální potenciálové pole*, které poskytuje autorům signály pro další lokální koordinaci chování robotů vůči absolutně lokalizovaným událostem. Je tak možné zhodnotit lokální odhady dalšího postupu v Kolektivu (například zavedením distribuce zájmových oblastí).

Zajímavá dizertační práce, která je postavena (stejně jako u mne) rovněž jen na bazální komunikaci a to na *komunikaci mezi robotem a hlodavcem* [Telenský a kol., 2006] ukazuje příklad jednostranného vztahu (vjemy jdou jen směrem od robota k živočichovi). Telenský v této své podnětné práci, nazvané „Biologická evoluce versus evoluční systémy robota“ porovnává učení potkana a robota v dynamických prostorových úlohách srovnatelného typu. Vjemy, které tak budí bazální akce robota, jsou analogické k podnětům, přicházejícími do neuronových buněk mozku (amigdal) hlodavce.

4.6.3 Zhodnocení simulací

Na rozdíl od reálných experimentů jsem k simulaci přistupoval s výrazně menšími omezeními. Bazální akce při simulaci lze definovat mnohem volněji. Obdobně není důvod pro simulaci uměle zavádět sofistikované interpretace měřených veličin a hledat aspoň částečnou spojitost mezi senzorickými hodnotami a cílem.

Proto jsem přistoupil k definici kritérií z opačné strany. Kritéria jsem zaváděl po diskusi nad uvažovanými mezními případy. Mezi uvažovanými případy pak už jen předpokládám monotónní funkci. Kritérium dále podporuje rozhodování; pro takové užití už podstatný tvar funkce. Simulace ukazuje, že v iterativním mechanismu výběru i zjednodušená formule vede na preferované chování. Výběr se omezuje na možnosti vymezené těmito okrajovými podmínkami. Systém v rámci nich konverguje do požadovaného stavu. Mezivýsledky konvergence ve stylu Umělého života pak více závisí na dalších podmínkách vyhodnocovaných mechanismem výběru akce.

V tomto kontextu mohu vzpomenout aktuální koncepci kooperace, kterou jsem mohl přehledově studovat po skončení simulačních experimentů [Espinoza, 2008]. Podle práce Ros je vidět, že případový popis řešení je krok správným směrem. Já jsem ovšem návrh případového ohodnocení chování dále nerozvíjel. Ros ukazuje, že vlastní výběr signifikantních případů je v podstatě samostatný proces, alternativa k expertní analýze vlivu jednotlivých senzorických veličin na cílové situace v kolektivu. Analytické vyšetření jsem mimo jiné striktně dodržoval také v experimentech s reálnými roboty.

4.6.4 Porovnání návrhu s jinými úspěšnými přístupy v oboru

Zvolená architektura řeší třídu problémů, které lze v okolí robota popsat rozprostřenými parametry. Měřítko pro vyhodnocení mnou zvolených experimentů jsou pro tuto práci nová (viz dále bod D.). Srovnání mých výsledků se světem se proto může týkat pouze metodologie. V kontextu s dříve připomenutým stavem problematiky (ve druhé kapitole) se ukazuje *několik zásadních odchylek předkládané koncepce řídicího systému mobilního robota*:

- v autostimulaci hodnotím *důsledky použitých chování mnou navrženým postupem*,
- typové reakce shrnuji do *pohybových schémat*,
- při volbě bazálních akcí se *inspiruji závěry eko-gramatického popisu*,

- a jako *doplňkové kritérium chování* udávám jeho úspěšnost (ověřitelné i ex-post).

A. Specifika behaviorálních vazeb s autostimulací, využívanou v mém návrhu

Autostimulace byla původně prezentována [Maes, 1991] na úrovni chování. V omezení na chování jako mezikroku postupu umožňuje výrazně omezit prudký nárůst uvažovaných alternativ ad-hoc vytvářeného plánu.

1. Oproti profesorce Maesové [Maes, 1991] já zařazuji autostimulaci při aktuálním ohodnocení podnětů do pohybového schématu.
2. Vzájemné vazby vytvářím (definuji) nejen *pro ohodnocení přechodu mezi chováními. Autostimulací udržuji nastolený směr při klidovém buzení bazálních akcí*. V kapitole 4.5 dokládám autostimulaci bazálních akcí v rámci jednoho chování.
3. Autostimulace jako ohodnocení přechodu mezi chováními v mém případě nově napomohla klasifikaci v kapitole 4.4).
4. *Autostimulaci nově podchycuji i pravidly eko-gramatického systému.*

Od původního „proaktivního udržitelného vývoje robota“, zavedeného Maes [Maes, 1991] se liší proto i mé zařazení autostimulace. Mé experimenty prováděné na reálných robotech prokázaly, že *autostimulací minimalizují náklady na odůvodněné aktivity robota* [Nahodil, Svatoš, 2002]. Tyto ***závěry a výsledky simulací i experimentů v reálu jsem nově prezentoval*** [Nahodil, Svatoš, 2007], [Svatoš, Nahodil, 2008a]. Obdobné závěry (ze svého pozdějšího výzkumu na obdobné platformě reálných robotů) publikoval i Peijnenburgův holandský tým [Peijnenburg a kol., 2002].

Autostimulace má dnes evidentně své místo v průzkumných aplikacích. V soutěžích robotického fotbalu je v tomto ohledu řešení holandského Philips Teamu překvapivé. Na rozdíl od českých týmů nedává Philips Team více prostoru na jemnější variace bazálních akcí v rámci jednoho chování. Autostimulace zcela vynechává. Kompletní chování jednoho robota Holanďané předepisují stavovým automatem.

B. Typové reakce v pohybovém schématu – účelné modifikace v mém návrhu

Ve své práci se přísně zaměřuji na lokální hodnocení výhod aktuální situace v Arkinem zavedených a Balchem zpřesněných pohybových schématech [Balch, Arkin, 1995a]. Právě oni používali ***pohybová schémata i na řízení robotů v Kolektivu. Jejich základní návrh jsem já – v mnou navržené architektuře řídicího systému mobilního robota výrazně modifikoval:***

1. **Normální rozložení** [Nahodil, Svatoš, 2001] využívám k parametrizaci vjemů a vnitřních podnětů v pohybovém schématu.
2. V normálním rozložení také zavádím a nově odvozené chování při adaptaci [Nahodil, Svatoš, 2007], [Svatoš, Nahodil, 2008b].
3. Upřesněním **parametrů normálního rozložení** *mohu popsat i externě dodávanou jednorázovou znalost při komunikaci* [Nahodil, Svatoš, 2002].

4. **Pro senzor** definuji zobrazení snímané veličiny do parametrů normálního rozložení (v perceptuálním schématu) [Nahodil, Svatoš, 2007].
5. **Zobrazení akcí**, odvozených pravidlovým systémem v sobě zahrnuje paměťovou složku. Do parametrů bazální akce reprezentované akčním schématem započítávám i autostimulace v kontextu dalších požadovaných aktivit.

Toto své původní řešení konfrontuji nyní se současnými trendy robotického fotbalu: Pro objektivní srovnání by měl být jasný rozsah rozhodování a k němu odpovídající hodnocení přípustného stavu. To autoři (jako své know-how) vesměs nepodávají. Nejvýše pouze *principy*:

Smyslem mých modifikací bylo nahrazení chybějící percepce robotů alternativními veličinami a dodatečnou autostimulací. Naopak německý Laueho tým [Laue, Röfer, 2005] skládá dnes z *percepce kompletní mapy hrací plochy*. Přirozeně, že v takovém záběru dostávají z potenciálového pole mnohem přesnější obraz. V Laueho realizacích každý z robotů svým vlastním potenciálovým polem globálně hodnotí významná místa.

Obdobné globální značení herního prostoru jsem ale již výrazně dříve volil ve svých simulacích i já [Svatoš, 2001]. Velmi podobné pozdější práce Vailovy a Telenského jsem vzpomněl již v předchozím oddíle 4.6.2 [Vail, Veloso, 2003, Telenský a kol., 2006].

C. Eko-gramatické závěry použité v mém návrhu

Odvození výsledků chování robotů dle apriorní gramatiky zůstává bohužel na okraji zájmu. Nižší zájem mimo jiné plyne i z povahy úzce zaměřených experimentů s reálnými roboty.

Zatímco simulace lze velice rychle přizpůsobit na více úloh, v reálných experimentech je vícenásobné použití podstatně omezené. U speciálních funkcí, na nichž jsou praktické aplikace postaveny, se nakonec posuzuje pouze syntéza vybraných chování.

Ve srovnání se *simulacemi eko-gramatického systému* v prostředí mobilních robotů kolegů [Sebestyén, Sosík, 2004, Takadama a kol., 1999] i já hodlám dodržet obdobný přístup. Rozdílně však nakládám s výkonnými komponentami. ***U bazálních akcí realizovaných podle mého původního návrhu mám já tato rozšíření:***

1. Užívám **složitější, kontextové odvození**. V definici eko-gramatického systému uvažuji odpovídající parametrické P1L-systémy.
2. Oproti tranzitivnímu uzávěru celého eko-gramatického systému **posuzují konečné řešení (vybranou akci) mnou nově zaváděným integrálním kritériem**. I na něm je vidět o řád složitější syntéza aktivit, pokud se zapojí více robotů.

Gramatika velmi napomáhá odhadu, kam až se může limitně dostat robot, respektive kolektiv robotů [Csuhaj-Varjú, Jiménes-López, 1998]. V tomto směru se i můj postup odlišuje od pouze účelového pokrytí vymezeného problému [Nahodil, Svatoš, 2002, Nahodil, Svatoš, 2007, Svatoš, Nahodil, 2008a, Svatoš, Nahodil, 2008b].

D. Kritéria v mém návrhu – celkové shrnutí specifík mého původního návrhu

Také v otázce kritérií je zřetelné, že *smyslem mnou navržené architektury není pokrytí předem vymezeného problému*. V hodnocení předvedených úloh se proto výrazně odlišují:

1. *Diferenciálním kritériem* hodnotím trend odvozený z několika bezprostředně předcházejících aktivit robota.
2. *Integrálním kritériem* hodnotím nákladovou stránku celé započaté realizace, případně podíl přípustných mezikroků na pokrytí rozsáhlého problému.
3. V definici kritérií označuji stav, který požaduji nejčastěji; a nakonec i v asymptotě. Dílčí kvality ostatních stavů vyhodnocuji relativně k označenému asymptotickému stavu.
4. Systém v tomto měřítku vynucuje vlastní asymptoty pro situace, kdy se roboti shodnou na interpretaci jedné události.
5. Z kritérií odvozují podněty, které ve zpětné vazbě ovlivní chování robota při adaptaci či ukončení aktuálního chování.

Jistý *ekvivalent mnou zavedeného diferenciálního kritéria* mohu vidět ve způsobu, jak [Bowling a kol., 2004] *indikuje v chování robotů nevhodné reakce*. Právě na úlohách robotického fotbalu Bowling ukazuje, že četnost akcí v požadovaném stavu je jedním z prvních kritérií pro hodnocení každého chování. Já ve své práci obdobně hodnotím i dílčí kroky, které takové chování provázejí. Každá úroveň mého pohledu na detail má opodstatnění – na úrovni chování lépe odvozují přechody, nezbytné pro adaptaci. Na úrovni bazálních akcí se zase lépe detekují výrazné odchylky od očekávané reakce; a to už v rámci předem daného chování.

Integrální vyhodnocení nákladové stránky se dnes už hojně užívá i jako měřítko kvality algoritmů, definovaných rozproštěnými parametry. **Obdobu mnou zaváděných ohodnocení kvality** návrhu nalézám například v Dorigově *hodnocení vhodné kombinace několika chování při navigaci robotů ve světelném bludišti* [Christensen, Dorigo, 2006]. V předložené práci ovšem mé integrální kritérium spolu s diferenciálním kritériem navíc poskytuje dva současné pohledy z různých úhlů na kvalitu jednoho a téhož rozhodovacího mechanismu.

E. Souhrnné srovnání – publikační aktivita

Odborné veřejnosti jsem své výsledky prezentoval jako autor či spoluautor v celkem deseti zadokumentovaných publikacích, z toho sedmi na prestižních mezinárodních konferencích: [Svatoš a kol., 2000, Svatoš a kol., 2001, Petrus a kol., 2001a, Petrus a kol., 2001b], [Svatoš a kol., 2001, Nahodil, Svatoš, 2001, Svatoš, 2001, Nahodil, Svatoš, 2007], [Svatoš, Nahodil, 2008a] a [Svatoš, Nahodil, 2008b].

4.6.5 Další možné rozšíření navrhovaného mechanismu

Inspirací a výchozím bodem pro zájemce o další rozšíření předkládaného tématu mohou být, stejně jako v mém případě, současné přístupy *robotického fotbalu*. Já však *Robotický fotbal* nepovažuji v této práci za svůj cíl, ale za zajímavou formou prezentovaný společný jmenovatel několika principiálních otázek, spojených s návrhem použitelných mobilních robotů.

V analýze možností (a nejen Robotického fotbalu) pak zájemcům doporučuji zaměřit se na nově otevřené body této práce:

1. **Komunikační protokol.** V mém pojetí protokol představuje transformaci dílčích poznatků o současném stavu řešení do podnětů akceptovatelných v ASM druhého robota. Je zřejmé, že v daném protokolu robot nezíská identický obraz situace, jako kdyby řešil zprostředkovanou situaci přímo na místě. Na druhou stranu se přirozeně rozšiřuje abeceda robota. A od rozšíření abecedy robota plyne i jeho generativní síla. Tématem pro další práci může být vyšetření těchto dvou protipólů použití protokolu. Vzhledem k architektuře blízké stavovému automatu se jeví jako perspektivní protokol XABSL (*Extensible Behavior Specification Language*), XABSL (Extensible Behavior Specification Language), používaný nově i vítězným německým Lötzsche-lovým týmem, známým z robotického fotbalu [Löttsch a kol., 2006].
2. **Rozšíření kolektivu.** I přes omezené možnosti laboratorních robotů se v práci ukázaly nové znaky chování dvou komunikujících robotů. Zajímavým námětem, a v dnešní době i experimentálně dostupným, je vyšetření hranice, za níž už podobné individuální znaky přímo indikují emergentní chování.
3. **Citlivost kritérií.** Otázku kvalitativních kritérií algoritmů dnes považuji za přirozený důsledek i současných snah o analýzu a ohodnocování chování přírodou inspirovaných řídicích systémů. Patrně nejpálčivějším problémem bude i nadále asi snaha o sjednocení všech přístupů k hodnocení chování robotů do univerzálnějších kritérií. Platnost takových kritérií by měla být invariantní vůči mohutnosti Kolektivu (s počtem robotů až o několik řádů vyšším). Jsem si vědom, že na základě jen této práce nemohu o vztahu mezi chováním a navrženými kritérii v tak rozsáhlém Kolektivu dělat z tohoto pohledu žádné seriózní závěry.

Shrnu-li:

Tato podkapitola byla vedena v rovině myšlenkových úvah zaměřených na zajímavé vlastnosti navržené hybridní architektury řídicího systému. Dále se zamyslím nad otevřenými možnostmi v práci uváděných rozšíření intencionálního a sociálního modelu chování.

V prvním oddíle připomínám jednotný způsob vyhodnocení podnětů v mechanismu výběru akce. V navrženém formátu lze hybridně kombinovat reaktivní závěry spolu s intencionálními odhady, jak se nejlépe přiblížit cílovému stavu.

Ve druhém oddíle uvádím hlavní znaky navrženého řídicího systému z pohledu Kolektivu. Praktické aspekty posuzuji u implicitní komunikace. Na zavedených rolích pak popisují způsob odvození chování robotů dvou odlišných rolí.

Dále navrhuji možné využití parametrizované vnitřní reprezentace při explicitní komunikaci. Po ní má k dispozici přímo popis, jak stejný problém ve své vnitřní reprezentaci

charakterizuje druhý robot. Ze senzorických vjemů druhého robota získává přehled o dosud nedetekovaných objektech. Z vnitřních stavů jiného robota může předjímat jeho další chování.

Ve třetím oddíle srovnávám oblast uplatnění navržené architektury s ostatními pracemi. Ve srovnání mi vychází výhodná parametrizace různých typových vjemů a jejich následná syntéza na úrovni bazálních akcí robota. Jednotné procedurální řešení na úrovni jednoho modelu pak udává i šablonu, která se bude dávat konkrétní náplň v případě adaptace.

Z pohledu adaptace je také důležitá volba kritérií, kterými indikují okamžik adaptace a poté i parametry nově doplněného chování. Jako samostatný bod uvádím své příspěvky, v kterých jsem uváděné principy prezentoval na vědeckých konferencích. Ve čtvrtém oddíle uvažuji otevřené možnosti ve větším důrazu na explicitní komunikaci, emergentní znaky chování, a korespondenci mezi uvažovanými kvalitativními kritérii a dosaženým chováním Kolektivu.

4.7 Shrnutí závěrů a důležitých faktů čtvrté kapitoly

Ve čtvrté kapitole, nazvané Experimenty prezentuji jak počítačové simulace, prováděné na modelu navrženého hybridního systému (který jsem rovněž vytvořil), tak ověření mých předchozích teoretických závěrů implementací tohoto hybridního řízení do skutečných robotů. Praktické možnosti mnou navržené architektury hybridního agenta ukazují v různých úlohách navigace mobilních robotů. Právě vhodně zvolenými úlohami postupně dokládám splnění jednotlivých cílů mé práce.

V první podkapitole se věnuji ověření návrhu hybridního řízení simulacemi na jeho modelech. Výsledky jsem zpracoval jednak tabulkově, jednak formou četných grafů a doplnil komentářem.

Ve druhé podkapitole popisují průběhy experimentu kontaktní navigace. Ve zvolených experimentech je základní řízení odvozeno z reaktivního modelu. Vycházím v něm z několika zvolených chování. V podkapitole uvedený postup řešení úlohy dvěma mobilními roboty беру za výchozí stav. Tento základní výsledek porovnávám s adaptivní modifikací takové úlohy (v souladu s mými předchozími teoretickými závěry). Model chování tedy doplňuji o adaptaci.

V třetí podkapitole demonstruji možnosti implicitní interakce mezi agenty. V experimentu předvádím reakce dvou heterogenních agentů v rámci společného, předem nedefinovaného řešení společného cíle.

Ve čtvrté podkapitole uvádím několik podrobností k programové realizaci experimentu.

V páté podkapitole srovnávám dosažené výsledky a zamýšlím se nad možnými analogiemi s výsledky jiných pracovišť.

Poslední podkapitola slouží ke shrnutí experimentálních výsledků z hlediska jak nejlepších v praxi dosažených postupů (výběru dalších akcí), tak jejich omezení.

Kapitola 5

Zhodnocení a závěr

Předložená disertační práce «Behaviorální řízení robotů s hodnocením přínosu jednotlivce pro kolektivní cíl» shrnula výsledky mé vědecko-výzkumné činnosti v oblasti mobilní robotiky od roku 1995 v podstatě do současnosti. Při své práci jsem se výrazně inspiroval pracemi v rámci skupiny mobilní robotiky MRG, vedené mým školitelem docentem Pavlem Nahodilem. Ten výzkum zde realizovaných diplomových prací a disertací nasměroval již v roce 1992 k behaviorální robotice.

Podstatnou část své energie jsem v této práci zaměřil na reaktivní řízení robotů. Velmi zajímavých výsledků jsem dosáhl při analýze koordinace chování robotů. Neméně zajímavých výsledků jsem dosáhl při aplikaci této analýzy v mechanismu výběru akce (ASM) především reaktivních robotů. V průběhu práce jsem dále modifikoval algoritmus koordinace chování pro robustnější řešení úlohy. Odvozený algoritmus se ukázal jako silný nástroj pro přirozenou distribuci úkolu uvnitř společenství. Také umožnil účelově ovlivnit chování celku, a to externě dodanými atraktivními vzory. *Tuto část považuji za hlavní příspěvek předkládané práce.*

5.1 Přehled výsledků práce včetně původního přínosu doktoranda

Práce, tak jak ji zde uvádím, především zahrnuje výsledky mého výzkumu v oblasti aplikované robotiky a umělého života coby interního doktoranda na katedře kybernetiky počátkem desetiletí. Doplnkově pak obsahuje i mnou nově navržené a alespoň v simulacích úspěšně odzkoušené a opublikované metody z posledních dvou let [Nahodil, Svatoš, 2007, Svatoš, Nahodil, 2008a].

Z nich si nejvíce vážím svého příspěvku, který jsem mohl úspěšně přednést na 8. konferenci *Kognice a umělý život* (KUZ VIII) v květnu 2008 [Svatoš, Nahodil, 2008b]. Ten navázal zejména na můj někdejší výzkum *Behavioural Formation Management in Robotic Soccer* [Svatoš, 2001], který jsem presentoval v hlavním plénu na renomované mezinárodní konferenci o umělém životě *Advances in Artificial Life* (ECAL 2001) – [In: LNAI 1600, Springer-Verlag, Berlin] a na rovněž oceněný příspěvek *Imitation Principle to Navigation in Robot Group* [Svatoš a kol., 2001] na prestižní akci *IFAC Workshop on Mobile Robot Technology v Soulu*.

V průběhu své práce jsem modifikoval **algoritmus koordinace chování pro robustnější řešení úlohy**. Odvozený algoritmus se ukázal jako silný nástroj pro přirozenou dis-

tribuci úkolu uvnitř společenství. Také umožnil účelově ovlivnit chování celku, a to externě určenými „atraktivními vzory“. Tuto část považuji za hlavní příspěvek předkládané práce. Další, snad neméně významné příspěvky mé práce uvádím bodově:

1. **Parametrizace vnitřní reprezentace.** Pro analytické vyhodnocení kombinace mechanismu ASM používám spojitě ohodnocení založené na superpozici dílčích příspěvků. Původní je v něm způsob mapování akcí. Umožňuje velmi variabilně parametrizovat a provázat pravidla vzájemné excitace dílčích akcí. Můj návrh byl jako původní prezentován na fóru IASTED [Nahodil, Svatoš, 2001].
2. **Adaptace parametrizované reprezentace podnětů.** V užité adaptaci na základě pokusu a omylu původně užívám *repulsivní ohodnocení pro vyjádření chybné reakce na vzniklou situaci*. Repulsivně do parametrizované vnitřní reprezentace registruji zlomové lokální podmínky, za kterých dochází k chybě. Původní návrh jsem s ohlasem prezentoval na fóru IFAC v Soulu [Svatoš a kol., 2001].
3. **Předvedení principů Umělého života na reálných robotech.** Ve vytvoření ucelené základní platformy mobilních robotů vidím praktický přínos. Na mezinárodní studentské konferenci POSTER'99 jsem se s projektem «Autonomous Mobile Robot Platform» umístil na 1.–3. místě. Původně vytvořená *čtveřice mobilních robotů se stala základním experimentálním nástrojem pro ověřování různorodých interakcí uvnitř společenství robotů*, modelovaných ve výzkumných simulačních návrzích členů Skupiny mobotiků na katedře kybernetiky ČVUT FEL.
 Jako aktivní spoluřešitel projektu FRVŠ 564/2000 «Inovace výuky mobilní robotiky využitím společenstva reálných robotů» svůj aktivní přínos i v zavedení platformy mobilních robotů a zde navržených metod do výuky doktorandů. To bylo při závěrečné oponentuře s uznáním vyzdviženo. O aplikacích mých výchozích, zcela původních postupů a algoritmů (na bázi Umělého života) přímo na reálných robotech jsem měl čest referovat i v komunitě IFAC EPAC [Svatoš a kol., 2000].
4. **Publikační aktivita dizertanta.** Odborné veřejnosti jsem své výsledky prezentoval jako autor či spoluautor v celkem deseti významnějších (v IS zadokumentovaných) publikacích. z toho bylo sedm prezentováno na prestižních mezinárodních konferencích: [Svatoš a kol., 2000, Svatoš a kol., 2001, Petrus a kol., 2001b, Petrus a kol., 2001a], [Nahodil, Svatoš, 2001, Svatoš a kol., 2001, Svatoš, 2001, Nahodil, Svatoš, 2007], [Svatoš, Nahodil, 2008a, Svatoš, Nahodil, 2008b].

Zejména z výsledků experimentů vyplývá kvalita mého **řešení**, postaveného na **principech Umělého života a opírajícího se o eko-gramatickou teorii**. Experimenty potvrdily, že užití spojitě hodnotící funkce v hybridním řízení robotů byl krok správným směrem. *Vzájemné kvalitativní ohodnocení elementárních možností aktuálního chování považuji za nedílnou součást pravidel hybridního řízení mobilního robota.* V praxi jsem platnost tohoto mého závěru ověřil na (v disertaci prezentovaných) Úlohách koordinace.

5.2 Splnění cílů disertační práce

Subjektivně hodnotím splnění jednotlivých cílů mé disertační práce, definovaných v jejím úvodu, následovně:

Cíl 1: Celá práce je věnována návrhu, realizaci, a ověření řídicího systému robota. Už od úvodu analýzy uvažuji o kombinaci (a) reaktivních podnětů a (b) dlouhodobě definovaných cílů robota. Architekturu, ve které zpracovávám tyto tři typy znalostí definuji jako otevřenou, a volně rozšířitelnou o znalostí přejaté od dalších robotů.

V této architektuře jsem později spojil události, generované modulem pro reaktivní předzpracování senzorů a modulem intencionální hodnocení aktivity robota. Pro tyto různorodé události jsem navrhl a implementoval algoritmus hybridního mechanismu výběru akce (v práci uváděný Algoritmus 1). Algoritmus jedním mechanismem shrnuje jak původní vjemy robota, tak i doplňkové hodnocení dosavadních reakcí v intencionální rovině.

Cíl 2: Chování odvozené z konečného seznamu bazálních akcí robota a konečného seznamu fyzických senzorů charakterizují eko-gramatickým systémem. Korespondenci jednotlivých entit řídicího systému je věnována kapitola 3.3.1. Do tohoto konceptu zavádím:

- Normální ohodnocení vnitřních podnětů a senzorických vjemů. Princip jsem prezentoval v [Nahodil, Svatoš, 2001, Svatoš, Nahodil, 2008b].
- Lineární ovlivňování vnitřních podnětů autostimulací.
- Jednorázové podněty s doznívajícím účinkem.

Vyhodnocení podnětů v tomto formátu shrnuji v Algoritmu 4, který přímo interpretuje navržené iterativní zpracování normálních podnětů dle vztahu (3.9).

Uvedený formalismus je přímo základem adaptačního Algoritmu 5. V něm adaptuji vnitřní soubor chování robota. Tento můj původní příspěvek jsem prezentoval [Svatoš a kol., 2001].

Cíl 3: Navrženou architekturu řídicího systému jsem koncepčně zacílil na naše laboratorní roboty. Experimentálně jsem demonstroval chování robotů v následujících situacích:

- sledování překážek
- vzájemné sledování robotů; každý z nich vnímá situaci rozdílně
- párování dvou robotů zcela rozdílného chování.

Skupina robotů řešila navigační úlohy v laboratorním prostředí. Rád bych zdůraznil neúměrnou náročnost reálných experimentů – jak časově, tak programátorsky.

Druhá fáze ověření byla v simulacích na standardní situaci Robotického fotbalu. V ní už jsem plně ukázal, jak spojit dva pohledy na chování týmu v intuitivní definici kritérií. Současně jsem v simulaci ukázal funkci postupné adaptace intencionálního modelu simulovaných hráčů na základě zvolených kritérií. Kritéria jsou zavedena v oddíle 3.5.2:

- v integrálním kritériu jsem hodnotil cenu dílčích kroků robota. Prakticky jsem pak tuto cenu přímo zahrnul do experimentů uváděných v oddílech 4.4 a 4.3.

- v diferenciálním kritériu jsem hodnotil dílčí reakci robota vzhledem k žádanému směru.

Pojetí dvou kritériálních hodnocení jedné události je v této práci původní. Díky kritériím u každého z doložených experimentů určuji použitelnost řídicího systému v individuálních problémových situacích. Z hodnot kritérií vyplývají také požadavky na adaptaci nevhodných reakcí.

Všechny algoritmy uváděné v této práci současně popisují funkci programů, které jsem vytvářel během realizace celého projektu. Výsledné programové knihovny poskytují

- rozhraní pro monitoring percepce laboratorního robota, protokolární přizpůsobení bazálních akcí do jazyka nízké úrovně robota
- rozhraní pro komunikaci se standardním serverem RoboCup
- schéma externího datového úložiště pro reakční a intencionální model chování
- komponentu reaktivního ASM, pracující nad externím datovým úložištěm,
- a komponentu hybridního ASM, včetně implementace funkcí autostimulace, zapomínání a adaptace.

V práci zpracovávám vytyčené téma na 112 stranách, členěných do pěti kapitol. Ty jsem na dalších 20 stranách doplnil seznamem obrázků, tabulek, algoritmů, přílohami dokumentujícími provedené experimenty, rejstříkem a úplným seznamem mnou použité literatury. Na příloženém CD je tato práce v pdf formátu. Dále jsou zde kratší pohybové sekvence (*.avi, *.mlv a obrázky *.jpg), vztahující se ke zde prezentovaným experimentům.

5.3 Závěry pro další rozvoj vědy nebo pro realizaci v praxi

Domnívám se, že některé dílčí výsledky mnou předkládané disertační práce představují původní vědecký přínos. Uvědomuji si přitom, že výsledky představují jen jeden z mnoha dalších kroků výzkumu hybridního řízení inteligentních mobilních robotů. Za přínos pro další rozvoj vědní disciplíny lze jistě považovat i publikování dílčích výsledků mé práce na šesti mezinárodních odborných konferencích s velmi příznivým posudky recenzentů i širší odborné veřejnosti při jejich vlastní prezentaci. V rámci mé interní přípravy na katedře kybernetiky jsem se jako autor, či spoluautor podílel na sepsání mnoha dalších (cca 8) dílčích výzkumných zpráv (lokálního významu), které zde již jmenovitě neuvádím.

Konkrétní přínos vidím ve svém návrhu (dalšími bloky) *volně rozšiřitelné architektury hybridního řídicího systému robota*. Rozsah možného rozšíření je vymezen formálním popisem systému. Formálně – v rámci eko-gramatického systému popisují chování Kolektivu, složeného z několika robotů. Tématem práce je stále aktuální integrace (včlenění) různorodých akčních prvků do jednotně pracujícího celku. Od tohoto přístupu se odvíjí i návazná témata, vhodná pro další samostatný výzkum. Zajímavým tématem pro pokračování je zejména aplikace zavedeného modelu, zejména pak ohodnocování kvality (do mnohem rozsáhlejšího kolektivního systému. Dalším stále otevřeným tématem je hlubší analýza vlivu zavedených kritérií v takovém rozsáhlém kolektivním systému.

Využití hybridní kombinace reaktivních a intencionálních výkonných jednotek se ostatně nemusí nakonec omezovat jen na Kolektiv robotů.

Vzhledem k předchozímu považuji cíle práce za splněné a proto si dovoluji předkládat tuto práci k obhajobě.

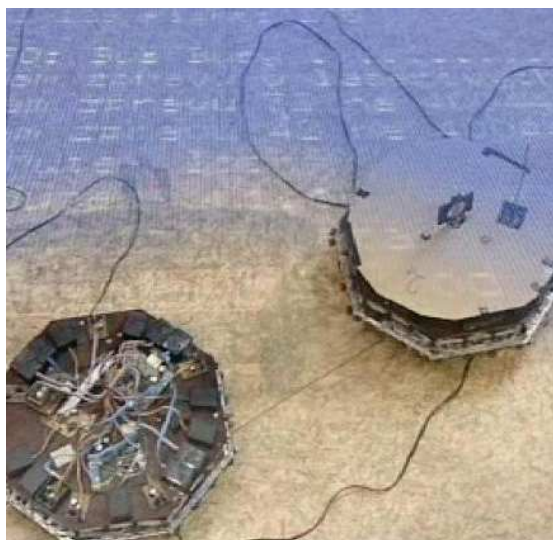
Dodatek A

Fotodokumentace

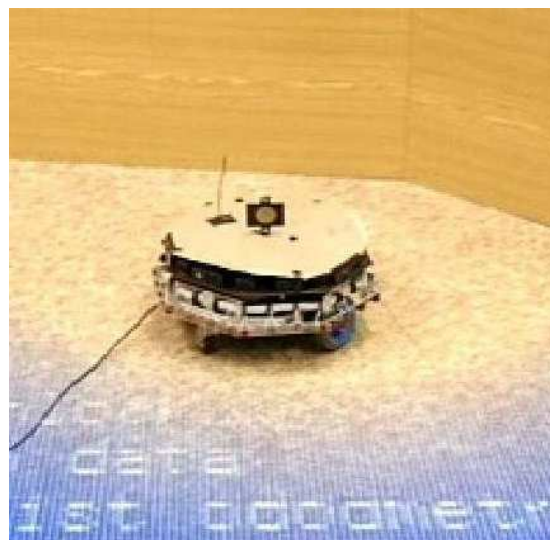
Příkládaná fotodokumentace ilustruje experimentální prostředí experimentů uváděných v kapitole 4. Jednotlivé snímky představují výseky z příkládaných videosekvencí, pro plnou představu proto odkazují na zdrojové videosekvence. Natočené sekvence ve formátu MPEG jsou součástí přiloženého CD.

První z videosekvencí byla pořízena jako záznam mimo experimenty. Přidávám ji zde proto, že plně vystihuje podstatu experimentů, a pro její autentičnost (snímek pochází z archivu České televize).

Následující videonahrávky poté pochází z archivu autora. Zachycují situace experimentů, které výzkum provázely, a to včetně experimentálních ověření popisovaných v této práci.



(a) Startovní pozice



(b) Individuální trajektorie robota

Obrázek A.1: Videosekvence sledování stěn. Snímek (a) zachycuje oba dva roboty ve startovní pozici této typové úlohy, uprostřed experimentální arény. Snímek (b) pak zachycuje vůdčího robota při nájezdu robotů podél stěny v rohu skutečné arény. Díky střihu tvůrců naleznete v příkládaném šotu na CD: *KratkyFilm.avi*



(a) Formace



(b) Sledování pohyblivého předmětu

Obrázek A.2: Koordinační úlohy. Snímek (a) ilustruje možné řízení kolektivu robotů ve formaci. Snímek dokresluje úvodní úlohu jak celého tématicky zaměřeného videozáznamu, ale i mých experimentů. Na příloženém CD se tohoto experimentu týkají pohyby robotů ve stopáži 00:00:17÷00:00:34. Jako úvodní předpoklad se o této práci v kapitole 4 ovšem nezmiňuji. Experimentální situace je na snímku (b). Na odpovídající části videonahrávky (stopáž 00:14:30÷00:16:30) je zachycena perioda akumulace nálad k sledování vůdčího objektu. Zdroj na CD: *MRG.mpeg*

Dodatek B

Realizace mechanismu výběru akce

B.1 Algorimizace experimentů

Při programovém řešení úloh jsem vytvořil nadstavbu nad jazykem makroúrovně fyzického řízení a monitoringu stavů robotů. Jeho hardwarová část byla implementována v assembleru jednočipových procesorů řady Atmel51 [Novák, Svatoš, 1999]. Následně jsem vyvíjel programovou část komunikačního protokolu pod operačními systémy Microsoft DOS a Microsoft Windows 3.1 ÷ Windows 2000. Programový a uživatelský manuál komunikačních rozhraní byl podrobně specifikován v [Svatoš, 1998] a pro studenty v [Nahodil a kol., 2001].

Pro simulaci paralelního chování několika agentů jsem vytvořil testovací platforma s pomocí vývojového nástroje Borland C++ 3.1 a 4.5 a Inprise Delphi™ 4.0. Pro komunikaci mezi pracovními stanicemi a mobilními roboty jsem použil komerční programové knihovny třetích stran [web AsyncPro].

Analytické možnosti mechanismu výběru akce jsem odvozoval v prostředí MathWorks Matlab®. Pro vyzkoušení koordinace více agentů jsem odladěný analytický model implementoval v klientské aplikaci robotického fotbalu [Svatoš, 2001]. Zdrojový kód klienta je v programovacím jazyku C++, při vývoji jsem využil prostředí Microsoft Visual C++ 6.0.

Celý programový balík výše popsanych programových knihoven je též několik let součástí znalostní báze skupiny Mobilní robotiky, pracující v doméně Umělého života.

Geometrie experimentů

Většina laboratorních experimentů probíhala v obdélníkové aréně $5,5 \times 3,5$ m. Jiné geometrie arény jsem zkoušel v dílčích modifikacích úloh. Takové modifikace v práci detailně neanalyzuji.

Záznam experimentů

Pro popis jednotlivých pravidel reakční báze užívám notaci, formulovanou v kapitole ?? vztahem (3.1). Experimentální výsledky jsem publikoval v interních výzkumných zprávách:

- Popis srovnávacího experimentu je publikován v [Svatoš, 2000c].
- Experiment autostimulace je publikován v [Svatoš, 2000a].
- Průběh experimentů s eko-gramatickým systémem je uveden v [Svatoš, 2000b].

B.2 Pravidla chování v experimentální simulaci

Obecně pro popis situace obě soupeřící skupiny vyhodnocují globální seznam absolutních poloh všech objektů. Z nich dále odvozují jednoduché pseudopodmínky, jaké bych v experimentech na reálných robotech očekával od senzorů.

Také není smyslem této simulace není pracovat s autostimulací. Autostimulaci mezi akcemi nepoužívám, pouze ponechávám pravidelné zvyšování a jednorázové anulování zobecněné charakteristiky akcí (viz tabulka B.1). Vzájemné stimulování probíhá pouze v mezích jednotlivých zobecněných charakteristik podnětu a jako takové může být i předmětem adaptace.

Chování a taktika obránců

Po startu simulace obránci implicitně (pouze na základě rozložení spoluhráčů okolo míče) přepínají mezi dvěma rozdílnými chováními – Rozehrávka a Zálaha:

Rozehrávka. Robot během ní drží míč do doby, než další spoluhráč zaujme výhodnější pozici (hodnoceno jeho ϵ -kritériem).

Zálaha. Smyslem chování je uvolnit robota na výhodnější pozici, kde pak očekává přihrávku.

Pro obě chování jsou charakteristické nejednoznačnosti ve výběru akce pouze na základě pseudopodmínek. O to více se pak uplatňuje v mechanismu výběru akce zobecněná charakteristika podnětů a hodnoty kritérií. Je třeba si pouze uvědomit, že v konečném důsledku se každé chování musí vyhodnocovat vůči jiné poloze. Cílem chování Rozehrávka je míč, cílem chování Zálaha je volná oblast.

A. Iniciální parametry

Na podkladě pseudopodmínek robot odvozuje výběr aktuální akce z pevného seznamu (do něj zahrnutých bazálních akcí). Úvodní výčet podnětů, včetně mapování do akčního a perceptuálního schématu pro simulaci, uvádí tabulka B.1. Každá skupina ovšem volbu mezi těmito bazálními akcemi podřizuje svým vlastním ASM. V něm přirozeně mohou být některé akce zatlumené.

B. Adaptace ve hře obránců

Řídící systém obránců se optimalizuje vůči pevné strategii protihráčů a obdobně optimalizované strategii spoluhráčů. Vlastní postup adaptačního učení se metodou pokusu a omylu (Algoritmus 5) byl pro tuto simulaci podstatně zkrácen. Vzhledem k dynamice objektů Robotického fotbalu jsem se omezil pouze na detekci selhání, následovanou volbou nové akce (bez kroku zpět, protože fyzickým krokem tentokrát robot nemůže návratu k původnímu stavu dosáhnout).

Výsledky adaptace uvádím v kontextu s chováním útočníků v tabulce B.2 níže.

Akce	$S[N, E]$	W	σ_1	σ_2	ϱ
otáčení robota	$[-0.5; -0.7]$	0 po aktivaci $W(t-1)+1$ jindy	0.05	0.05	0.75
posun vpřed	$[1; 0]$	0 po aktivaci $W(t-1)+1$ jindy	0.05	0.05	0
zachycení míče	$[0; 0]$	0 po aktivaci 0.2 pro míč v pohybu	0.005	0.005	0
odkopnutí míče	$[-0.5; 0.5]$	0 po aktivaci $\max(\epsilon)$ při držení míče	0.05	0.05	0
krátká přihrávka (se záměrem dále následovat míč)	$[0.3; -0.3]$	0 po aktivaci 0.1 při držení míče	0.05	0.05	-0.75
vyčkávání na místě	$[0; 0]$	0 po aktivaci $W(t-1)+1$ jindy	0.05	0.05	0
Pseudopodmínka	$S[N, E]$	W	σ_1	σ_2	ϱ
míč v dosahu	$[0; 0.3]$	0	0.1	0.1	1
spoluhráč čelně	$[-0.2; 0.45]$	0	0.1	0.1	0.5

Tabulka B.1: Tabulka bazálních akcí použitých v simulaci. Tabulka se vztahuje k obrázku B.1. V tabulce se uvádí iniciálně nastavené rozložení intencionálního modelu. Rozložení vychází z rutinního nastavení robota ve srovnávacím experimentu.

Útočníci a jejich taktika

Cíle útočníků jsou mnohem více svázané s dynamikou prostředí. Hybridní architekturu řídicího systému v jejich případě totiž zužují jen na reaktivní model. Částečně je tak odvozené chování útočníků závislé i na chování obránců. Reaktivním ASM rozhoduje útočník mezi třemi akcemi

- *otáčení* robota
- *posun* vpřed
- *zachycení* míče.

Přičemž hlavním cílem útočníka je zachycení míče. Tato reakce se spouští, jakmile je míč detekován v dosahu útočníka.

A. Chování a významová metrika

Reaktivita robota se vyhodnocuje vždy proti vztažnému cíli. Obdobně jako u obránce útočník v závislosti na vzdálenosti od cíle předvádí dvě chování:

Útok Útočník blíže k míči se snaží dosáhnout míč.

Středopole. Vzdálenější útočník pak zamezuje přihrávkám mezi soupeři.

Z popisu chování je zřejmá i poloha cíle takového chování. Zatímco pro chování Útok je cíl určen polohou míče, v případě Středopole je cílovou polohou oblast mezi protihráči bez míče. Pro tato chování se zobecněné charakteristiky podnětů nepoužívají, ASM je čistě reaktivní.

Akce	$S[N, E]$	W	σ_1	σ_2	ϱ
otáčení robota	$[-0.5; -0.7]$	0 po aktivaci $W(t-1)+1$ jindy	0.485	0.221	0.75
posun vpřed	$[1; 0]$	0 po aktivaci $W(t-1)+1$ jindy	0.274	0.156	0
zachycení míče	$[0; 0]$	0 po aktivaci 0.2 pro míč v pohybu	0.005	0.005	0
odkopnutí míče	$[-0.5; 0.5]$	0 po aktivaci $\max(\epsilon)$ při držení míče	0.259	0.120	0
krátká přihrávka (se záměrem dále následovat míč)	$[0.3; -0.3]$	0 po aktivaci 0.1 při držení míče	0.1766	0.156	-0.75
vyčkávání na místě	$[0; 0]$	0 po aktivaci $W(t-1)+1$ jindy	0.207	0.152	0

Tabulka B.2: Rozložení intencionálního modelu po adaptaci. Tabulka se vztahuje k obrázku 4.3(a). Oproti iniciálnímu nastavení jsou zde udávány skutečné rozptyly zobecněných charakteristik podnětů po 100 krocích simulace adaptace intencionálního modelu.

B.3 Pravidla chování ve srovnávacím experimentu

Reakční báze

Základní axiomy úlohy uvažují pro startovní krok – 0.1– a podmínku ukončení – 0.2:

- 0.1 : S : $\mapsto$ průzkum-naslepo : 1
0.2 : : sledování-zdi, $\epsilon \leq \epsilon_{STOP}$ $\mapsto$ STOP : 1

V místě bez překážek robot nejprve používá chování průzkum-naslepo a reaguje s minimální vahou 1.0:

- 1.0 : F : průzkum-naslepo $\mapsto$ $krok_+$: 0,1

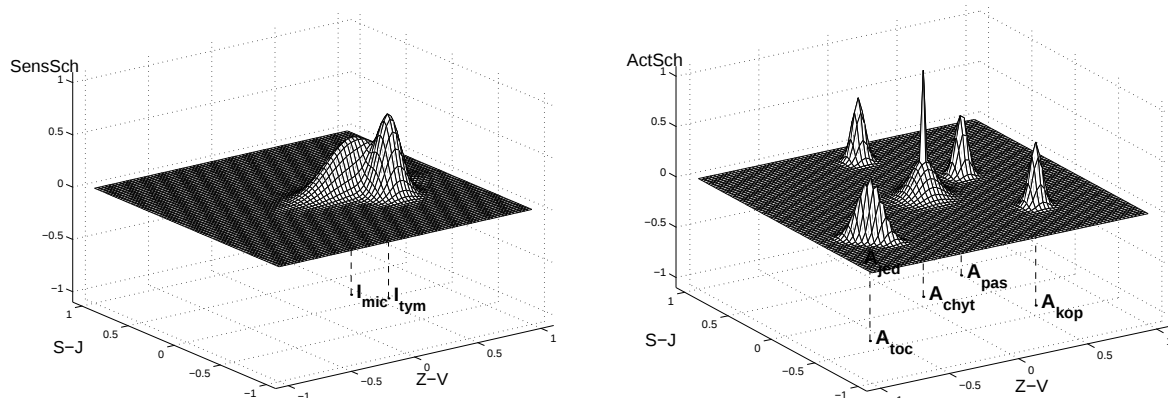
Do sledování-zdi se robot dostává po detekci nárazníky — pravidlo 1.1 řeší případ mechanické detekce, pravidlo 1.2 případ infračidel:

- 1.1 : B_n : průzkum-naslepo $\mapsto$ sledování-zdi : 0,9 $n \in 1..24$
1.2 : I_n : průzkum-naslepo $\mapsto$ sledování-zdi : 0,9 $n \in 1..16$

Jemnější reakce na blízkost překážky ošetřují pravidla nižší váhy 2.0 až 2.6:

- 2.0 : B_n : sledování-zdi $\mapsto$ $krok_-$: 0,9 $n \in \{1, 2, 3, 4\}$
2.1 : I_n : sledování-zdi $\mapsto$ $toč_-$: 0,8 $n \in \{1, 2, 3, 4, 5, 16\}$
2.2 : $I_{15}I_{16}$: sledování-zdi $\mapsto$ $toč_-$: 0,7
2.3 : $I_{14}I_{15}$: sledování-zdi $\mapsto$ $krok_+$: 0,6
2.4 : I_{14} : sledování-zdi $\mapsto$ $krok_+$: 0,5
2.5 : I_{13} : sledování-zdi $\mapsto$ $toč_+$: 0,4
2.5 : I_{13} : sledování-zdi $\mapsto$ $toč_+$: 0,3 $n \in \{6, 7, 8, 9, 10, 11, 12, 15\}$

Kritickým okamžikem je náraz. Řeším ho pravidlem 2.0 jako krok zpět. Kontextová pravidla 2.1-2.5 popisují odhad výskytu nárazových ploch podle kombinace vybuzených infračervených čidel. Pravidlo 2.6 shrnuje obecnou reakci na podněty na straně robota, odvrácené k překážce. Pravidlo plně pokrývá doplněk dosud nedefinovaného zpracování podnětů z infračervených nárazníků.



Obrázek B.1: Obránci a jejich taktika (intence) v různých místech hracího pole. Grafy přibližují počáteční rozložení podnětů v intencionální model obránců, při simulaci. v části (a) je zobrazen průběh akčního schématu a v části (b) průběh perceptuálního schématu. Osy v grafu jsou normované na rozměr hrací plochy. Na osu obránce se v tomto měřítku vynáší vzdálenost nejbližšího spoluhráče. Analogicky na osu útoku se vynáší v daném měřítku vzdálenost k nejbližšímu protihráči.

S nejnižší vahou pak robot přechází od sledování-zdi k chování průzkum-naslepo. Ke změně chování dochází u robota v klidu, kdy neregistruje žádný podnět z nárazníků:

2.7 : F : sledování-zdi $\mapsto$ průzkum-naslepo : 0,1

Transformace do pohybového schématu

term	rádus R [-]	azimut α [rad]	váha W [-]	rozptyl σ [-]	deformace $c_{\tilde{\alpha}}$ [-]
I_n	0,7	$n \cdot \frac{\pi}{12}$	-0,4	0,1	1
B_n	0,4	$n \cdot \frac{\pi}{6}$	-0,5	0,3	1
$krok_+$	0,8	0	0,35	0,1	5
$krok_-$	0,6	π	0,1	0,1	5
$toč_+$	0,7	$-\frac{\pi}{2}$	0,25	0,1	5
$toč_-$	0,7	$\frac{\pi}{2}$	0,25	0,1	5

Tabulka B.3: Transformace pravidel do pohybového schématu. Tabulka ukazuje korespondenci mezi pravidly reakční báze a odpovídajícím pohybovým schématem. Senzorické podněty robot vyhodnocuje jako repulsivní. Bazální akce jsou naopak hodnoceny atraktivně. Největší důraz je kladen na pohyb vpřed, nejmenší naopak na pohyb vzad.

Rádus podnětu odpovídá délce kroku pro dosažení signalizované události. Podle tohoto schématu robot při pohybu dělá větší krok než při pohybu vzad. Parametry rozptylu a deformace slouží k zobecnění podnětu. Rozptyl odpovídá váze podnětu s normálním rozložením $N([R, \alpha], \sigma)$. Toto rozložení se v případě deformace různé od 1 ještě přenáší ve stejném rádiu ještě v úhlové odchylce $\Delta\alpha = \pm \arctan(c_{\tilde{\alpha}} \frac{\sigma}{R})$. Nejširší význam tak má podnět z mechanických nárazníků, kterým se zamezuje pohyb v celém odpovídajícím sektoru. Podobně je do celého sektoru rozveden význam jednotlivých akcí, tentokrát ovšem s pevnou délkou kroku R .

Omezující podmínky

Práh pro ukončení algoritmu pravidlem 0.2 odpovídá podílu obvodu R^1 sledovaných překážek na ploše R^2 arény. V dokumentované úloze ze skutečného poměru $R^1 : R^2$ vylapnul koeficient $\epsilon_t = 5\%$.

B.4 Pravidla chování pro ověření autostimulace

Reakční báze

Naváděcí robot

Tohoto robota jsem ovládal pevnou pohybovou sekvencí, chováním

průzkum-cesty=

$120 \times \textit{krok}_+, 3 \times \textit{toč}_-, 40 \times \textit{krok}_+, 3 \times \textit{toč}_-, 30 \times \textit{krok}_+, 3 \times \textit{toč}_-, 20 \times \textit{krok}_+, 3 \times \textit{toč}_-,$
 $30 \times \textit{krok}_+, 3 \times \textit{toč}_-, 20 \times \textit{krok}_+, 3 \times \textit{toč}_-, 30 \times \textit{krok}_+, 3 \times \textit{toč}_-, 20 \times \textit{krok}_+.$

Sekvenci jsme určil jednoúčelově, s ohledem na experimentální situaci. Chování průzkum cesty je přímo spojené se startovním a cílovým stavem

$$\begin{array}{llll} 0.11 : S & : & \mapsto \textit{průzkum-cesty} & : 1 \\ 0.12 : B_n & : \textit{průzkum-cesty} \mapsto \textit{STOP} & : 1 & n \in 1..24 \end{array}$$

Sledující robot

Při návrhu pravidel chování sledujícího robota jsem vyšel z reakční báze použité ve srovnávacím experimentu. Jako mezistupeň mezi průzkumem-naslepo a sledováním-zdi zavádím pronásledování. Právě v novém chování využívám podnětů z autostimulace. Posloupnost přechodů mezi chováními vystihuje stavový diagram na obrázku 4.9(a). Přechod od průzkumu-naslepo ke sledování-zdi určují pravidla 1.5 a 1.6.

$$\begin{array}{llll} 1.5 : B_n & : \textit{průzkum-naslepo} \mapsto \textit{pronásledování} & : 0,9 & n \in 1..24 \\ 1.6 : I_n & : \textit{průzkum-naslepo} \mapsto \textit{pronásledování} & : 0,9 & n \in 1..16 \end{array}$$

V opačném směru zavádím zpoždění, ukončené změnou hodnoty infračerveného čidla. Pro další rozhodnutí je klíčové právě čidlo, na kterém proběhla změna. V následujících pravidlech je označuji F_n .

$$\begin{array}{llll} 3.0 : B_n & : \textit{pronásledování} \mapsto \textit{krok}_- & : 0,9 & n \in \{1, 2, 3, 4\} \\ 3.1 : I_n & : \textit{pronásledování} \mapsto \textit{čekej}(n) & : 0,7 & n = 1..16 \\ 3.2 : F_n & : \textit{pronásledování} \mapsto \textit{toč}(n), \textit{průzkum-naslepo} & : 0,7 & n = 1..16 \end{array}$$

Obdobně se pro ošetření mimořádné situace v pravidle 3.0 omezují pouze na čelní náraz. Speciálně pro tento experiment také doplňuji literál *toč*. Abych nastavil orientaci robota ve směru změny, volám akci *toč* s vázaným parametrem – indexem čidla, kde byla detekována změna.

Autostimulace

U sledujícího robota se autostimulace dotýká bazální akce *čekej*. Při výběru akce se autostimulace projevuje ve váze této akce v poměru k váze změny chování na sledování-zdi. Váhy obou chování $W = [W_{\textit{čekej}}, W_{\textit{sledování-zdi}}]^T$ jsem nastavil bez vzájemného vlivu:

$$W(t+1) = \begin{bmatrix} 1 & 0 \\ 0 & 0,9 \end{bmatrix} W(t)$$

Zatímco váha sledování-zdi je konstantní, váha akce *čekej* exponenciálně doznívá díky tlumícím účinkům použité autostimulace ($A_{22} < 1$). Na dostatečně rychlou změnu IR nárazníků tak robot stihne ještě reagovat pravidly chování pronásledování. Jinak v rozhodování sledujícího robota převáží chování sledování-zdi.

B.5 Pravidla chování pro ověření interakcí EG systému

Naváděcí robot

Robot R_1 se řídí zcela pravidly chování průzkum-plochy. Startovní a ukončovací axiomy vývoje společného atraktoru jsou pak v případě jediného sledujícího robota triviální. Stop stav robota R_1 je určen robotem R_2 .

0.6: $S \mapsto \text{průzkum-plochy} : 1$

0.7: $\text{průzkum-plochy} : \text{STOP}_2 \mapsto \text{STOP}_1 : 1$

Pravidla pohybu robota R_1 formuloval na základě pravděpodobností:

4.0: $B_n : \text{průzkum-plochy} \mapsto \text{krok}_- : 0,9 \quad n \in \{1, 2, 3, 4\}$

4.1: $I_n : \text{průzkum-plochy} \mapsto \text{toč}_{180} : 0,9 \quad n \in \{1, 2, 3, 4\}$

4.2: $F : \text{průzkum-plochy} \mapsto \text{krok}_+ : p_1(t)$

4.3: $F : \text{průzkum-plochy} \mapsto \text{toč}_+ : p_2(t)$

4.4: $F : \text{průzkum-plochy} \mapsto \text{toč}_- : p_3(t)$

Sledující robot

Robot R_2 určuje směr z intenzity světla, snímané a akumulované ve čtyřech kvadrantech. Robot se s takovouto znalostí rozhoduje už od začátku experimentu pravidly cíleného postupu.

0.8: $S \mapsto \text{cílený-postup} : 1$

V následujících pravidlech jednoduchým znakem $C_x(\alpha)$ vyjadřuji aktuální intenzitu světla, detekovanou odpovídajícím senzorem. Parametr α představuje aktuální měřenou hodnotu intenzity světla. Zdvojené $\mathbb{C}_x(\alpha)$ pak udržuje akumulovanou hodnotu od doby posledního rozhodování. Následujícím pravidlem pro přepis vjemu do vnitřní reprezentace formalizují diskrétní akumulaci snímané intenzity světla:

0.9: $C_n(\alpha) : \text{akce},$
 $\mathbb{C}_n(\beta) \mapsto \mathbb{C}_n(\alpha + \beta) : 0,9 \quad n \in \{s, j, v, z\};$
 $\text{akce} \in \{\text{sledování-zdi}, \text{cílený-postup}\}$

přičemž veličiny ve směru pohybu robota označují s . Analogicky veličiny měřené na zádi robota označují j a veličiny měřené v postranních směrech v , respektive z . Změnu rozložení světla v tomto popisu robot R_2 vyhodnocuje jako změnu vázaného parametru α v pravidlech chování cílený-postup. Směr určují goniometrickou aritmetikou

5.1: $\mathbb{C}_s(\alpha_s)\mathbb{C}_j(\alpha_j)\mathbb{C}_v(\alpha_v)\mathbb{C}_z(\alpha_z), \alpha_v \neq$
 α_z
 $\text{cílený-postup} \mapsto \text{toč}\left(\arctan\left(\frac{\alpha_s - \alpha_j}{\alpha_v - \alpha_z}\right)\right) \text{reset}, \text{krok}_{-\bar{\alpha}}$

5.2: $\mathbb{C}_s(\alpha_s)\mathbb{C}_j(\alpha_j)\mathbb{C}_v(\alpha_v)\mathbb{C}_z(\alpha_z) : \text{reset} \mapsto \mathbb{C}_s(0)\mathbb{C}_j(0)\mathbb{C}_v(0)\mathbb{C}_z(0) : 0,9$

V závislosti na akumulované intenzitě robot může i chování cílený-postup ukončit. Robot R_2 při mimořádně intenzivním vjemu dostává do cílového stavu. Naopak po nedostatečných podnětech robot přechází na sledování-zdi.

5.3: $\mathbb{C}_s(\alpha_s)\mathbb{C}_j(\alpha_j)\mathbb{C}_v(\alpha_v)\mathbb{C}_z(\alpha_z) :$
 $\text{balacovaná-navigace},$
 $\sum \alpha_i < 0,1 \mapsto \text{balacovaná-detekce} : 0,9$
0.10: $: \text{cílený-postup}, \epsilon \leq \epsilon_t \mapsto \text{STOP}_2 : 1$

Dodatek C

Obsah příloženého CD

soubor	stopáž	popis
KratkyFilm.avi	celé	sledování zdi. Váže se ke srovnávacímu experimentu v kapitole 4.3.
Robocup.m1v	celé	interakce robota s míčem
MRG.avi	0:17 - 0:34	skupina čtyř robotů - použití formace
	1:00 - 1:40	skupina čtyř robotů
	1:40 - 5:40	sledování zdi. Váže se ke srovnávacímu experimentu v kapitole 4.3.
	5:40 - 6:10	světelná navigace - 1 robot
	6:10 - 6:30	světelná navigace - 2 roboti. Váže se k ověření EG systému v kapitole 4.5.
	6:30 - 7:30	světelná navigace - 1 robot
	14:30 - 16:30	sledování pohybu (člověk/robot). Váže se k experimentu autostimulace v kapitole 4.4.

Tabulka C.1: Videosekvence na CD. Soubory jsou v adresáři VIDEO

soubor	popis
Skupina.jpg	skupina čtyř mobilních robotů - experimentální platforma Gerstnerovy laboratoře.
Seminar.jpg	ze seminářů Aplikované Robotiky

Tabulka C.2: Fotografie na CD. Soubory jsou v adresáři FOTO

soubor	popis
Thesis.pdf	dizertační práce
Teze.pdf	Teze k dizertační práci

Tabulka C.3: Dokumenty na CD. Soubory jsou v adresáři DOC

Rejstřík

- agent, 12
 - animát, 7, 13
 - hybridní, 41
 - inteligentní, 12
 - intencionální, 17, 20
 - reaktivní, 17, 18
 - sociální, 18
- akce
 - bazální, 4, 10, 12, 39
 - zatlumená, 42
- architektura
 - BSA, 24
 - hybridní, 9, 38
 - vrstvená, 19, 38, 58
- ASM, *viz* mechanismus výběru akce
- autostimulace, 22
- cíle práce
 - formulace, 4
 - splnění, 110
- emergence, 1, 96
- gramatika, 8
 - formalismus
 - aktivační podmínka, 42
 - odvozovací krok, 43
 - pravděpodobnostní faktor, 42
 - pravděpodobnostní faktor, 41, 42
 - generativní síla, 8
 - kolonie, 31
 - regulární gramatika, 30
- chování, 36, 39
- intence, *viz* záměr
- kritérium, 16
 - diferenciální, 36
 - integrální, 36
- mechanismus výběru akce, 14
 - intencionální, 55
 - reaktivní, 44
 - vítěz-bere-vše, 49
- model
 - intencionální, 20, 49
 - pohybové schéma, 19, 53
 - akční schéma, 20
 - perceptuální schéma, 20
 - reaktivní, 18, 40
- modularita, 13
 - modul adaptace, 59
 - neomodulární imaginace, 15
 - neomodulární adaptace, 15
 - neomodulární relace, 15
- podnět
 - atraktivní trend, 21
 - exponenciální zapomínání, 51
 - zobecněná charakteristika, 55
- robot, 8, 13
 - behaviorální robotika, 7, 9
- řízení
 - distribované, 8
 - neomodulární, 15
 - přístup
 - bottom-up, 1, 7, 21, 72
 - hybridní, 36
 - top-down, 7
- sociální inteligence, 8
- stavový prostor, 36
- systém
 - eko-gramatický, 52
 - Lindenmayerův, 8, 31
 - parametrický, 43
 - s pamětí, 36
- vzor, 36, 39

atraktivní, 36

repulsivní, 36

záměr, 1, 14, 17, 21

Literatura

- [Arkin, 1989a] Arkin, R. C., (Ed.) (1989a). *Behavior-based robotics*. The MIT Press, Boston.
- [Arkin, 1989b] Arkin, R. C. (1989b). Motor schema-based mobile robot navigation. *Journal of Robotics Research*, 8(4):92–112.
- [Arkin, 1992] Arkin, R. C. (1992). Cooperation without communication: Multi-agent schema based robot navigation. *Journal of Robotic Systems*, 9(3):351–364.
- [Arkin, 1997] Arkin, R. C., (Ed.) (1997). *Robot Colonies*. Kluwer Academic Publishers, Boston, MA.
- [Balch, 2000] Balch, T. (2000). Lifelong adaptation in heterogeneous teams: Response to continual variation in individual robot performance. *Autonomous Robots*, 8(3):133–156.
- [Balch, Hybinette, 2000] Balch, T., Hybinette, M. (2000). Social potentials for scalable multi-robot formations. In *IEEE International Conference on Robotics and Automation (ICRA '00)*, volume 1, str. 73–80.
- [Balch, 1997] Balch, T. R. (1997). Behavioral diversity in robot teams. In *AAAI 97 Workshop on Multiagent Learning*, str. 67–75. Providence.
- [Balch, Arkin, 1995a] Balch, T. R., Arkin, R. C. (1995a). Communication in reactive multi-agent robotic systems. *Autonomous Robots*, 1(1):223–231.
- [Balch, Arkin, 1995b] Balch, T. R., Arkin, R. C. (1995b). Motor schema-based formation control for multiagent robot teams. *Autonomous Robots*, 1(1):64–72.
- [Barnes, 1996] Barnes, D. P. (1996). A behavior synthesis architecture for co-operant mobile robots. In Gray, J. O., Caldwell, D. G., (Eds.), *Advanced Robotics and Intelligent Machines*, number 51 in IEE Control Engineering Series, str. 295–314. The Institution of Electrical Engineers, London, UK.
- [Bensaid, Mathieu, 1997] Bensaid, N., Mathieu, P. (1997). A hybrid architecture for hierarchical agents. str. 91–95. Griffith University, Gold-Coast, Australia.
- [Bowling a kol., 2004] Bowling, M., Browning, B., Chang, A., Veloso, M. (2004). Plays as team plans for coordination and adaptation. In Polani, D., Browning, B., Bonarini, A., Yoshida, K., (Eds.), *RoboCup 2003: Robot World Cup VII*, number 3020 in LNAI, str. 686—693. Springer-Verlag, Berlin, Germany.

- [Brooks, 1991] Brooks, R. A. (1991). New approaches to robotics. *Science*, 253(5025):1227–1232.
- [Brooks, 1999] Brooks, R. A., (Ed.) (1999). *Cambrian Intelligence*. The MIT Press, Boston.
- [Chernova, Veloso, 2007] Chernova, S., Veloso, M. M. (2007). Confidence-based policy learning from demonstration using gaussian mixture models. In Durfee, E. H., Yokoo, M., Huhns, M. N., Shehory, O., (Eds.), *International Conference on Autonomous Agents and Multiagent Systems (AAMAS'07)*, str. 233–240, New York. IFAAMAS.
- [Christensen, Dorigo, 2006] Christensen, A., Dorigo, M. (2006). Evolving an integrated phototaxis and hole-avoidance behavior for a swarm-bot. In Rocha, L. M., Yaeger, L. S., Bedau, M. A., Floreano, D., Goldstone, R. L., Vespignani, A., (Eds.), *Artificial Life X: Proceedings of the Tenth International Conference on the Simulation and Synthesis of Living Systems*, str. 248–254. MIT Press, Cambridge, MA.
- [Csuha-j-Varjú, Jiménez-López, 1998] Csuha-j-Varjú, E., Jiménez-López, M. D. (1998). Cultural eco-grammar systems: A multi-agent system for cultural change. In *Proc. of the MFCS satellite workshop on Grammar Systems*, str. 165–182, Brno, Czech Republic.
- [Csuha-j-Varjú a kol., 1997] Csuha-j-Varjú, E., Kelemen, J., Kelemenová, A., Păun, G. (1997). Eco-grammar systems - a grammatical approach to lifelike interactions. *Artificial Life*, 3:1–28.
- [Donald a kol., 1997] Donald, B. R., Jennings, J., Rus, D. (1997). Information invariants for distributed manipulation. *International Journal of Robotics Research*, 16:673–702.
- [Espinoza, 2008] Espinoza, R. R. (2008). *Action Selection in Cooperative Robot Soccer using Case-Based Reasoning*. PhD thesis, Universitat Automa Barcelona, Barcelona, Spain.
- [Ferber, 1996] Ferber, J. (1996). *Reactive Distributed Artificial Intelligence. Principles and Applications*. John Wiley & Sons, New York.
- [Gat, Dorais, 1994] Gat, E., Dorais, G. (1994). Robot navigation by conditional sequencing. In *Proceedings of IEEE International Conference on Robotics and Automation*, volume 2, str. 1293–1299.
- [Goldberg, Matarić, 1997] Goldberg, D., Matarić, M. J. (1997). Interference as a tool for designing and evaluating multi-robot controllers. In *AAAI 97*, str. 637–642.
- [Havel, 2007] Havel, I. M. (2007). Causal domain and emergent rationality. In Smith, B., Broguard, B., (Eds.), *Rationality and Irrationality. Proceedings of the 23rd International Wittgenstein Symposium*, Kirchberg.
- [Holldobler a kol., 1974] Holldobler, B., Moglich, M., Machwitz, U. (1974). Communication by tandem running in the ant *Camponotus sericeus*. *Journal of Comparative Physiology*, 90:105–127.

- [Ieropoulos a kol., 2005] Ieropoulos, I., Melhuish, C., Greenman, J., Horsfield, I. (2005). Ecobot-ii: An artificial agent with a natural metabolism. *Advanced Robotics Systems*, 2(4):295–300.
- [Kadleček, Nahodil, 2001] Kadleček, D., Nahodil, P. (2001). New hybrid architecture in artificial life simulation. In *Advances in Artificial Life*, number 2159 in LNAI, str. 143–146, Berlin, Germany. Springer-Verlag.
- [Kelemen, 1989] Kelemen, J. (1989). A note on achieving low-level rationality from pure reactivity. *Journal of Experimental & Theoretical Artificial Intelligence*, 8:121–127.
- [Kelemen, 1993] Kelemen, J. (1993). Multiagent symbol systems and behavior-based robots. *Applied Artificial Intelligence*, 7:419–432.
- [Kelemen, Kelemenová, 1992] Kelemen, J., Kelemenová, A. (1992). A grammar - theoretic treatment of multi-agent systems. *Cybernetics and systems*, 23(1992):621–633.
- [Kelemen a kol., 2001] Kelemen, J., Kelemenová, A., Mitraná, V. (2001). Intelligence as an emergent behavior; or, the songs of Eden. In Martin-Vide, C., Mitraná, V., (Eds.), *Where Mathematics, Computer Science, Linguistic and Biology Meet*, str. 63–74. Kluwer Academic Press, Cambridge, MA.
- [Khatib, 1986] Khatib, O. (1986). Real-time obstacle avoidance for manipulators and mobile robots. *International Journal of Robotic Research*, 5(1):90–98.
- [Kristensen, 1995] Kristensen, S. (1995). *Sensor Planning Using Bayesian Decision Theory*. PhD thesis, University of Amsterdam, Amsterdam, Netherland.
- [Langton, 1992] Langton, C. G. (1992). Preface. In Langton, C. G., Taylor, C., Farmer, J. D., Rasmussen, S., (Eds.), *Artificial Life II*, volume X of *of the Santa Fe Institute Studies in the Sciences of Complexities*, str. 1–47. Addison-Wesley, Redwood City, CA.
- [Langton, 1984] Langton, L. G. (1984). Self-reproduction in cellular automata. *Physica*, 10D:135–144.
- [Laue, Röfer, 2005] Laue, T., Röfer, T. (2005). A behavior architecture for autonomous mobile robots based on potential fields. In *RoboCup 2004: Robot World Cup VIII*, number 3276 in LNAI, str. 122–133. Springer-Verlag, Berlin, Germany.
- [Lindenmayer, 1968a] Lindenmayer, A. (1968a). Mathematical models for cellular interaction in development. *Journal of Theoretical Biology*, 18:280–315.
- [Lindenmayer, 1968b] Lindenmayer, A. (1968b). Mathematical models for cellular interaction in development. *Journal of Theoretical Biology*, 18:280–315.
- [Lorenz, 1993] Lorenz, K. (1993). *Základy etologie*. Academia, Praha.
- [Lötzsch a kol., 2006] Lötzsch, M., Risler, M., Jüngel, M. (2006). Xabsl - a pragmatic approach to behavior engineering. In *Proceedings of IEEE/RSJ International Conference of Intelligent Robots and Systems (IROS)*, str. 5124–5129, Beijing, China.

- [Maes, 1990] Maes, P. (1990). Situated agents can have goals. In Maes, P., (Ed.), *Designing Autonomous Agents: Theory and Practice from Biology to Engineering and Back*, str. 49–70. MIT Press, Cambridge, MA.
- [Maes, 1991] Maes, P. (1991). A bottom-up mechanism for action selection in a artificial creature. In Meyer, J.-A., Wilson, S. W., (Eds.), *Proceedings of the First International Conference on Simulation of Adaptive Behavior: From Animals to Animats*, str. 238–248. MIT Press.
- [Mařík a kol., 2003] Mařík, V., Štěpánková, O., Lažanský, J., (Eds.) (2003). *Umělá inteligence IV*. Academia, Praha.
- [Mařík a kol., 2007] Mařík, V., Štěpánková, O., Lažanský, J., (Eds.) (2007). *Umělá inteligence V*. Academia, Praha.
- [Matarić, 1992] Matarić, M. J. (1992). Integration of representation into goal-driven behavior-based robots. *IEEE Transaction on Robotics and Automation*, 8:204–312.
- [Matarić, 1994] Matarić, M. J. (1994). *Interaction and intelligent behavior*. PhD thesis, MIT EECS, Cambridge, MA. MIT AI Lab AIRT-1495.
- [Minski, 1985] Minski, M. (1985). *Society of Minds*. MIT Press, Boston.
- [Nahodil, 2007] Nahodil, P. (2007). Od jednoduchých robotů k anticipujícím stvořením. In Mařík, V., Štěpánková, O., Lažanský, J., (Eds.), *Umělá inteligence V*, str. 374–404. Academia, Praha.
- [Nahodil a kol., 2004] Nahodil, P., Slavík, P., Kadleček, D., Řehoř (2004). Dynamic analysis of agents behaviour combining alife visulaziation and ai. In *Engineering Societies in the Agents World IV*, number 3071 in LNCS, str. 346–359, Berlin, DE. Springer-Verlag, Heidelberg.
- [Nahodil, Svatoš, 2001] Nahodil, P., Svatoš, V. (2001). Sense-action encoding in hybrid decision system. In *Proceedings of IASTED International Conference on Modelling, Identification and Control*, volume 1, str. 414–419, Calgary, Canada. IASTED.
- [Nahodil, Svatoš, 2002] Nahodil, P., Svatoš, V. (2002). Low-band communication to dynamical role assignment in robotic soccer. In *Proceedings of Workshop 2002*, volume A, str. 348–349. CTU Prague.
- [Nahodil, Svatoš, 2007] Nahodil, P., Svatoš, V. (2007). Application of artificial life in robot navigation. In *Proceedings of Workshop 2007*, volume A, str. INF030, Czech Republic. CTU Prague.
- [Nahodil a kol., 2001] Nahodil, P., Svatoš, V., Petrus, M., Kurzveil, J. (2001). Aplikovaná robotika. Technical Report GL 52/01, CTU FEE, Department of Cybernetics, The Gerstner Laboratory, Prague, Czech Republic. 13 s.
- [Novák, Svatoš, 1999] Novák, P., Svatoš, V. (1999). Control system for experimental autonomous mobile robot. In *Proceedings of Workshop '99*, str. 102, Czech Republic. CTU Prague.

- [Parker, 1993] Parker, L. E. (1993). Designing control laws for cooperative agent teams. In *Proceedings of the 1993 International Conference on Robotics and Automation*, str. 582–587. IEEE.
- [Parker, 1998] Parker, L. E. (1998). Alliance: An architecture for fault-tollerant multi-robot cooperation. *IEEE Transactions on Robotics and Automation*, 14(2):220–240.
- [Parker, 2000] Parker, L. E. (2000). Lifelong adaptation in heterogeneous teams: Response to continual variation in individual robot performance. *Autonomous Robots*, 8(3):133–156.
- [Parker, Touzet, 2000] Parker, L. E., Touzet, C. (2000). Multi-robot learning in a cooperative observation tasks. In Parker, L. E., Bekey, G., Barhen, J., (Eds.), *Distributed Autonomous Robotic Systems 4*, str. 1391–401. Springer-Verlag.
- [Pasteels a kol., 1987] Pasteels, J. M., Deneubourg, J.-L., Gross, S. (1987). Self-organization mechanisms in ant societies. In Pasteels, J. M., Deneubourg, J.-L., (Eds.), *From Individual to Collective Behavior in Social Insects*, str. 155–175. Experientia Supplementum.
- [Peijnenburg a kol., 2002] Peijnenburg, A. T. A., Veldhuis, G. P., Warmerdam, T. P. H. (2002). Philips cft robocup platform selection. In *Proceedings of IEEE international symposium on intelligent control ISIC'02*, str. 811–814. Kluwer Academic Publishers.
- [Petrus a kol., 2001a] Petrus, M., Nahodil, P., Svatoš, V. (2001a). Environment influence into reactive agents behaviour adaptation. In *Proceedings of the International Carpatian Control Conference 2001*, volume 1, str. 545–550, Krakow, Poland. Wydział Inżynierii Mechanicznej i Robotyki.
- [Petrus a kol., 2001b] Petrus, M., Nahodil, P., Svatoš, V. (2001b). Interactions in community of behaviour-based agents. In *Proceedings of the IASTED International Conference EUNITE*, volume 1, str. 57–62, Aachen, Germany. Verlag Mainz.
- [Petrus, Svatoš, 2000] Petrus, M., Svatoš, V. (2000). Extended reactive behaviour in mobile robots. In *Proceedings of 4th International Student Conference on Electrical Engineering "Poster 2000"*, str. IC 19. CTU Prague.
- [Reynolds, 1994] Reynolds, C. W. (1994). Competition, coevolution and the game of tag. In Brooks, R. A., Maes, P., (Eds.), *Proceedings of the Fourth Workshop on Artificial Life*, str. 59–69, Boston, MA. MIT Press.
- [Röfer a kol., 2007] Röfer, T., Brose, J., Göhring, D., Jüngel, M., Laue, T., Risler, M. (2007). Germanteam 2007. In Visser, U., Ribeiro, F., Ohashi, T., Dellaert, F., (Eds.), *RoboCup 2007: Robot Soccer World Cup XI Preproceedings*. RoboCup Federation; <http://www.robocup.org/>.
- [Sebestyén, Sosík, 2004] Sebestyén, P., Sosík, P. (2004). Kolektivní učení robotů v prostředí eko-gramatického systému. In Kelemen, J., Kvasnička, V., (Eds.), *Kognice a umělý život VI.*, str. 449—462. Slezská universita Opava, Opava.

- [Staffiotti a kol., 1995] Staffiotti, A., Konolige, K., Rusipini, E. H. (1995). A multivalued logic approach to integrating planning and control. *Artificial Intelligence*, 76:481–526.
- [Starý, 2000] Starý, J. (2000). Rozšiřující HW pro chování inteligentních mobilních robotů. Master's thesis, ČVUT FEL v Praze, Praha.
- [Stein, 1994] Stein, L. A. (1994). Imagination and situated cognition. *Journal of Experimental and Theoretical Artificial Intelligence*, 6:393–427.
- [Stone a kol., 2006] Stone, P., Kuhlmann, G., Taylor, M. E., Liu, Y. (2006). Keepaway soccer: From machine learning testbed to benchmark. In Noda, I., Jacoff, A., Bredendfeld, A., Takahashi, Y., (Eds.), *RoboCup-2005: Robot Soccer World Cup IX*, volume 4020, str. 93–105. Springer Verlag, Berlin.
- [Sudolský, 2006] Sudolský, S. (2006). Tímová stratégia v multiagentovom prostredí. In Kelemen, J., Kvasnička, V., (Eds.), *Kognice a umelý život VI.*, str. 347—354. Slezská universita Opava, Opava.
- [Svatoš, 1997] Svatoš, V. (1997). Rozšiřující SW pro evoluci vlastností mobilních robotů. Master's thesis, ČVUT FEL v Praze, Praha.
- [Svatoš, 1998] Svatoš, V. (1998). Návrh vrstvené hierarchie ŘS robota. Technical Report GLM 05/98, ČVUT FEL, Katedra řídicí techniky - Gerstnerova laboratoř, Praha. 11 s.
- [Svatoš, 2000a] Svatoš, V. (2000a). Environment-based co-operation among robots. Technical Report GL 96/00, CTU FEE, Department of Cybernetics, The Gerstner Laboratory, Prague, Czech Republic. 36 s.
- [Svatoš, 2000b] Svatoš, V. (2000b). Gradient navigation of robot. Technical Report GL 97/00, CTU FEE, Department of Cybernetics, The Gerstner Laboratory, Prague, Czech Republic. 26 s.
- [Svatoš, 2000c] Svatoš, V. (2000c). Sense-action encoding in hybrid decision system. Technical Report GL 117/00, CTU FEE, Department of Cybernetics, The Gerstner Laboratory, Prague, Czech Republic. 13 s.
- [Svatoš, 2001] Svatoš, V. (2001). Mnohonásobně motivované chování jako prostředek účelné navigace robota. Vyzvaná přednáška.
- [Svatoš a kol., 2001] Svatoš, V., Kurzveil, J., Nahodil, P. (2001). Imitation principle to navigation in robot group. In *IFAC Workshop on Mobile Robot Technology*, volume 1, str. 140–145, Seoul, Korea. IFAC.
- [Svatoš, Nahodil, 2008a] Svatoš, V., Nahodil, P. (2008a). Behavioral criterions in robot community. In *CTU Reports - Proceedings of WORKSHOP 2008*, volume 12, str. 86–87, Czech Republic. CTU Prague.
- [Svatoš, Nahodil, 2008b] Svatoš, V., Nahodil, P. (2008b). Skupinová kritéria určující chování jednotlivce. In Kelemen, J., Kvasnička, V., Pstružina, K., (Eds.), *Kognice a umelý život VIII.*, str. 307–312, Opava. Slezská universita Opava.

- [Svatoš a kol., 2000] Svatoš, V., Nahodil, P., Kadleček, D., Kurzveil, J., Maixner, V. (2000). Community of mobile robots in education. In *Proceedings of IFAC EPAC 2000 Conference: The 21st Century Education and Training in Automation and Control*, str. 53–57.
- [Svatoš, 2001] Svatoš, V. (2001). Behavioural formation management in robotic soccer. In *Advances in Artificial Life*, number 1600 in LNAI, str. 553–562, Berlin, DE. Springer-Verlag, Berlin.
- [Svatoš a kol., 2001] Svatoš, V., Petrus, M., Nahodil, P. (2001). Behavioural approach in gradient navigation of robot. In *Proceedings of the International Carpatian Control Conference 2001*, volume 1, str. 109–114, Krakow, Poland. Wydział Inżynierii Mechanicznej i Robotyki.
- [Takadama a kol., 1999] Takadama, K., Hajiri, K., Nomura, T., Shimohara, K., Nakasuka, S. (1999). Grammatical learning model for distributed computing: A grammatical approach. In Păun, G., Salomaa, A., (Eds.), *Grammatical Models of Multi-Agent Systems*, str. 343–355. Gordon and Breach, London.
- [Telenský a kol., 2006] Telenský, P., Jiroutek, P., Svoboda, J., Blahna, K., Bureš, J. (2006). Biologická evoluce versus evoluční systémy: porovnání učení potkana a robota v dynamických prostorových úlohách srovnatelného typu. In Kelemen, J., Kvasnička, V., (Eds.), *Kognice a umělý život VI.*, str. 373–392. Slezská universita Opava, Opava.
- [Tinbergen, 1953] Tinbergen, N. (1953). *Social Behaviour in Animals*. Methuen & Co. Ltd., London, UK.
- [Vail, Veloso, 2003] Vail, D., Veloso, M. (2003). Dynamic multi-robot coordination. In *Multi-Robot Systems: From Swarms to Intelligent Automata*, volume II, str. 87–100. Kluwer Academic Publishers.
- [Weiss, 1999] Weiss, G. (1999). *Multi-agent Systems. A Modern Approach to Distributed Artificial Intelligence*. MIT Press, Cambridge, MA.
- [Wooldridge, 2002] Wooldridge, M. R. (2002). *Wooldridge M.: An Introduction to Multi-Agent Systems*. John Wiley & Sons, New York.
- [Wooldridge, Jennings, 1994] Wooldridge, M. R., Jennings, N. J. (1994). Formalising the cooperative problem solving process. In *Proceedings of 13th International DAI Workshop*, str. 403–417.
- [Wooldridge, Jennings, 1995] Wooldridge, M. R., Jennings, N. J. (1995). Agent theories, architectures and languages: A survey. In Wooldridge, M. R., Jennings, N. J., (Eds.), *Intelligent Agents*, volume 890, str. 1–39. Springer-Verlag, Berlin. LNAI.
- [web AsyncPro] <http://sourceforge.net/projects/tpapro>. Přístup 15.11.2008.
- [web Keepaway] <http://www.cs.utexas.edu/users/AustinVilla/sim/keepaway/>. Přístup 15.11.2008.