

Целеполагание и синтез плана поведения когнитивным агентом*

А.И. Панов

ИСА ФИЦ ИУ РАН, г. Москва, Россия

Аннотация. В работе представлен знаковый подход к проблеме моделирования процесса целеполагания и его интеграции с методами синтеза плана поведения. На основе психологически правдоподобной модели знаковой картины мира когнитивного агента предложен алгоритм GoalMAP, который реализует итеративный процесс иерархического планирования с выделенным этапом постановки или выбора новой цели. Рассмотрена оценка сложности алгоритма планирования поведения. Проведены модельные эксперименты с программной реализацией построенных алгоритмов, демонстрирующей ключевые особенности используемого подхода.

Ключевые слова: когнитивный агент, планирование, знак, знаковая картина мира, теория деятельности, культурно-исторический подход, целеполагание, планирование поведения, GoalMAP.

DOI 10.14357/20718594180202

Введение

Вопрос моделирования процесса целеполагания в интеллектуальных системах является одним из наиболее насущных проблем в области повышения степени автономности робототехнических систем. Этим вопросом занимаются как в рамках классического искусственного интеллекта [1-3], так и в новых направлениях моделирования когнитивных функций человека [4,5]. Расширение областей применимости робототехнических систем приводит к тому, что расширяется и спектр возможных ситуаций, в которых может оказаться робот. Становится все более трудным заранее предугадать и заложить информацию о возможном развитии событий при достижении поставленной роботу цели. Существует большое количество примеров, когда робот или агент не может достигнуть первоначальной цели и ему необходимо вырабатывать или выбирать из имеющегося списка

новую цель. Примером может служить автономный глубоководный аппарат, перед которым ставится цель исследования определенного участка морского дна. Аппарат может оказаться в ситуации, когда на данном участке обнаружился посторонний объект (затонувший корабль), который препятствует выполнению первоначальной миссии. Методы целеполагания должны позволить агенту прекратить выполнение прежнего плана и поставить новую цель, например, заключающуюся в отправке сообщения в командный центр о наличии неопознанных объектов на дне.

Актуальность решения задачи моделирования целеполагания подчеркивалась достаточно давно не только специалистами по искусственному интеллекту, но и психологами, для которых феномен целеполагания (или целеобразования) до сих пор остается малоизученным: «Энтузиастам метода математического моделирования мы предлагаем использовать феномен целеобразования в качестве критерия для оцен-

*Работа выполнена при частичной поддержке РФФИ, проекты №№ 17-07-00281, 16-37-60055.

ки качества предлагаемых математических моделей психической деятельности: если модель воссоздает процесс целеобразования, то она «хорошая», если нет, то «плохая» [6].

Так как единственная реализация процесса целеполагания, поддающаяся подробному изучению, наблюдается только у человека, то достижения психологической науки в этой области должны служить первоисточником теорий и экспериментальных данных для построения моделей целеполагания. В настоящей статье мы используем отечественные работы в области теории деятельности и культурно-исторического подхода [7,8], которые постулируют, что целеполагание является неотъемлемой частью деятельности человека или отдельным ее видом и его моделирование должно быть так или иначе связано с моделированием процесса планирования поведения.

Принятое в искусственном интеллекте представление о цели, как о формально описанной некоторой конечной ситуации, является лишь частным случаем сложного психологического процесса работы с целями. В рамках теории деятельности выделяется несколько механизмов целеполагания [6], среди которых: преобразование побочного результата действия в цель на основе его осознания и связывания с мотивом, переформулирование целей при недостижении первоначально предвосхищавшегося результата, выбор одной из множества задаваемых целей, образование иерархии и временной последовательности целей и др. В настоящей работе будет рассмотрена модель одного из этих механизмов целеполагания, связанная с этапом планирования, а именно – выделение промежуточных целей как функции препятствия и усвоение заданной цели путем связывания ее с мотивом. В работе будет представлен психологически правдоподобный алгоритм синтеза плана поведения когнитивных агентом и предложена его модификация для реализации описанных выше механизмов целеполагания.

В современных работах, посвященных вопросу целеполагания и в более общей постановке – рассуждениям о целях, можно выделить два основных направления: когнитивное и классическое. Оба подхода определяют цель как класс состояний внешней среды, описываемый агентом на основе некоторого способа представления знаний и удовлетворяющий определенным условиям. В рамках когнитивно-

го направления работа с целями относится к так называемым метакогнитивным процессам, которые регулируют протекание других когнитивных процессов, например, планирования.

В рамках классического подхода к вопросу рассуждения о целях вводятся формальные определения цели с позиций теории планирования и выделяются несколько типов рассуждений, использующих понятие цели. Здесь стоит в первую очередь отметить работы по так называемой целенаправленной автономности [2,9,10], в которой выделяется четыре основных шага по работе с целями: определение рассогласования в знаниях или наблюдениях, объяснение рассогласования, генерация цели и управление целью. Первый шаг связан с наблюдением так называемой аномалии или особым событием, в качестве которого может выступать наблюдение невыполнимости текущей цели, новые непредвиденные обстоятельства, влияющие на выполнение плана и т.п. На втором шаге агент может использовать некоторый логический вывод для пополнения информации о возникшем событии, например, с использованием абдуктивных рассуждений, или накопленный опыт в виде отображения событие-цель. В данном случае подразумевается, что цели появляются и изменяются в процессе исполнения плана. Это важный момент, так как подход целенаправленной автономности не предполагает этапа постановки цели до того момента, когда мы начнем выполнять план и можем столкнуться с новой непредвиденной ситуацией. То есть в этом подходе некоторая начальная цель в любом случае ставится извне, что тем самым сужает сам смысл процесса целеполагания.

Когнитивное направление в моделировании целеполагания рассмотрим на примере когнитивных архитектур, в которых ставится цель построения системы психологически или биологически правдоподобного управления поведением агента. В некоторых архитектурах, рассматривающих понятие цели в одной из своих подсистем, имеется структура списка целей [11], в которую и из которых они по определенным правилам помещаются и выбираются для исполнения, что можно считать простейшей работой с целями. В других архитектурах часть эпизодической памяти предназначена для построения предсказания будущего развития событий и поддерживает возможность смены

цели [12, 13]. В некоторых когнитивных архитектурах [14] предпринимается попытка построить отдельный метакогнитивный уровень принятия решений, в котором присутствует как основанная на памяти о предыдущих состояниях среды определение гипотетических причин возникновения событий и анализ непредвиденных обстоятельств на основе выявленных причинно-следственных связей. Несмотря на глубокую проработку связей между подсистемой выдвижения и приоритизации целей и остальными подсистемами системы управления поведением (памятью, моторно-перцептивными механизмами, подсистемой мотивации и оценки), подход, развиваемые в когнитивных архитектурах, остаются сугубо теоретическим, а практическую реализацию находят лишь простые списки целей и реализация частных, похожих на рассматриваемые в целенаправленную автономности, случаев.

Статья организована следующим образом. В Разделе 1 приведено формальное описание модели знаковой картины мира, которая лежит в основе предлагаемого в Разделе 2 психологически правдоподобного метода планирования поведения когнитивным агентом. В конце Раздела 2 приводятся оценки сложности алгоритма синтеза плана. В Разделе 3 приведены две модели, реализующие два типа целеполагания: эмпирический («внутренний») и сценарный («внешний»). В Разделе 4 приведены экспериментальные результаты, носящие в основном теоретико-модельный характер, но, тем не менее, демонстрирующие особенности представленных алгоритмов, и обсуждаются некоторые результаты.

1. Картина мира когнитивного агента

В данной работе мы будем использовать модель знаковой картины мира [4, 15, 16] и ее сетевую формализацию, изложенную в работах [17, 18] и введем основные понятия и определения. Основным элементом картины мира является знак, который может представлять как статический объект, так и действие. Знак задается именем и содержит компоненты образа, значения, личностного смысла. Компонента образа кодирует характерные признаки представляемого объекта или процесса. Компонента значения представляет доступные коллективу агентов обобщенные сценарии использования

объектов. Компонента личностного смысла знака определяет роль объекта в действиях субъекта, совершаемых им с этим объектом. Личностные смыслы знака формируются в процессе деятельности субъекта и являются конкретизацией сценариев из значений этого знака. Личностные смыслы раскрывают предпочтения субъекта деятельности, отражают мотив и эмоциональную окраску действий, формируют опыт действования.

Введем специальную структуру – каузальную матрицу, которую будем использовать для формального описания компонент знака.

Определение 1. *Каузальной матрицей будем называть следующую структуру $z = \langle e_1, e_2, \dots, e_t \rangle$ – кортеж длины t событий e_i – битовых векторов (колонок) длины h . Каждому индексу j вектора событий e_i (строке матрицы z) мы будем ставить в соответствие некоторый кортеж, возможно пустой, каузальных матриц Z_j , такой, что $z \notin Z_j$. С помощью специальной процедуры $\Lambda : 2^Z \rightarrow 2^{\mathbb{N}} \times 2^{\mathbb{N}}$, которая некоторому подмножеству каузальных матриц $Z' \subseteq Z$ из всего множества матриц Z ставит в соответствие два непересекающихся подмножества индексов столбцов $I^c \subset \mathbb{N}, \forall i \in I^c i \leq h$ и $I^e \subset \mathbb{N}, \forall i \in I^e i \leq h : \Lambda(Z') = (I^c, I^e)$ таких, что $I^c \cap I^e = \emptyset$. Множество I^c для матрицы z будем называть индексами столбцов условий, а множество I^e – индексами столбцов эффектов матрицы z .*

Каузальная матрица выступает основным составляющим элементом компонентов знака. Определенная выше структура позволяет единым образом кодировать, как статическую информацию и признаки объекта, так и динамические процессы. Встроенная возможность задания причин и эффектов позволяет кодировать базовое отношение, выделяемое по данным о внешней среде – причинно-следственное.

Определение 2. *Знаком будем называть четверку $s = \langle n, p, m, a \rangle$, где n – имя знака, $p = Z^p$, $m = Z^m$, $a = Z^a$ – кортежи каузальных матриц, которые соответственно называются образом, значением и личностным смыслом знака s .*

Исходя из Определений 1 и 2, все множество каузальных матриц распадается на три под-

множества (для образов, значений и личностных смыслов знаков), которые организованы в так называемые каузальные сети.

Определение 3. *Каузальная сеть $W = \langle V, E \rangle$ является помеченным ориентированным графом, в котором:*

- *каждому узлу $v \in V$ ставится в соответствие кортеж каузальных матриц $Z^p(s)$ образа некоторого знака s , что будем обозначать как $v \rightarrow Z^p(s)$;*

- *дуга $e = (v_1, v_2)$ принадлежит множеству дуг графа E , если $v_1 \rightarrow Z^p(s_1), v_2 \rightarrow Z^p(s_2)$ и $s_1 \in S_p(s_2)$, т.е. если знак s_1 является элементом образа s_2 ;*

- *каждой дуге графа $e = (v_1, v_2), v_1 \rightarrow Z^p(s_1), v_2 \rightarrow Z^p(s_2)$ ставится в соответствие метка $\epsilon = (\epsilon_1, \epsilon_2, \epsilon_3)$ - кортеж трех натуральных чисел:*

- ϵ_1 - индекс исходной матрицы в кортеже $Z^p(s_1)$, может принимать специальное значение 0, если исходными могут служить любые матрицы из кортежа;
- ϵ_2 - индекс целевой матрицы в кортеже $Z^p(s_2)$, строка которой ставится в соответствие признаку s_1 ;
- ϵ_3 - индекс столбца в целевой матрице, в которой в соответствующей признаку s_1 строке стоит 1, может принимать положительные значения (столбцы условий) и отрицательные (столбцы эффектов).

Каузальная сеть служит основной для определения основных отношений на множестве компонент знака [17] и некоторых процессов, моделирующих базовые когнитивные функции (обобщения, замыкания по значениям, агглютинации).

Определение 4. *Моделью знаковой картины мира будем называть семиотическую сеть $\Omega = \langle W_m, W_a, W_p, R, \Theta \rangle$, где W_m, W_a, W_p - каузальные сети на значения, личностных смыслах и образах соответственно, $R = \langle R^m, R^a, R^p \rangle$ - семейство отношений на компонентах знака, Θ - семейство операций на множестве знаков.*

Для описания процесса планирования в знаковой картине мира нам понадобятся понятия активности в семиотической сети и процесса его распространения. Введем метку активности для каузальных матриц сети W_x ($x \in \{p, m, a\}$) и будем называть активным множество Z_x^* матриц, обладающих этой меткой. Процесс распространения активности представляет собой изменение состава множества Z_x^* с течением времени (каждый дискретный момент) и описывается для каждого типа каузальной сети своей функцией: $\varphi_a, \varphi_m, \varphi_p$. Процесс распространения активности является итерационным, т.е. на каждом шаге новый состав множества активных матриц порождается на основе предыдущего состава добавлением новых матриц, связанных по дугам сети с текущими активными. В простейшем случае мы будем рассматривать такой процесс, в котором каждая матрица не влияет на ход распространения активности от другой матрицы и поэтому будем считать, что функции φ_x принимают на вход одну активную матрицу и выдают новое подмножество активных матриц.

В связи с тем, что ребра каузальных сетей имеют направление, будем различать распространение активности вверх по сети, когда используются исходящие от узла дуги (функции $\varphi_x^\uparrow$), и распространение активности вниз по сети, когда используются входящие в узел дуги (функции $\varphi_x^\downarrow$). В дальнейшем, при описании алгоритма планирования, нам понадобятся только функции на сети значения и личностных смыслов. Каждую функцию $\varphi_x^\uparrow, \varphi_x^\downarrow$ будем параметризовать глубиной распространения активности d_x , которая указывает на какую глубину просматриваются дуги в данном направлении (вверх или вниз).

2. Синтез плана поведения когнитивным агентом

Алгоритм синтеза плана поведения агента опирается на следующие основные положения теории деятельности. Во-первых, предлагаемый алгоритм реализует иерархическое планирование, в котором имеется отделение действий и операций и присутствует иерархия целей. Каж-

дое действие имеет свой операционный состав, который либо известен заранее (как составляющие его образ поддействия), либо может быть получен в результате отдельного алгоритма планирования (один из вариантов целеполагания Раздел 3). Во-вторых, действия всегда предметны, т.е. на сети значений связаны со знаками, представляющими предметы, выполняющие в действии ту или иную роль. Разделение множества доступных агенту процедур на действия, направленные на достижение цели, и операции, отражающие специфику конкретной ситуации, также согласуется с теорией Канемана по «быстрым» и «медленным» действиям [19].

Основная идея планирования в знаковой картине мира заключается в использовании локальности организации знаний, которая обеспечивается за счет того, что процедурные и декларативные знания сгруппированы друг с другом в такие группы (знаки), которые объединены на основе опыта взаимодействия с некоторым объектом или выполнения действий в некоторой ситуации. А так как целенаправленная деятельность агента является предметной и ситуативной, т.е. условия действий определяются текущей ситуацией, то и объединение информации о ключевых особенностях объекта, о тех действиях, в которых он может участвовать и о том опыте, который уже накопился у агента по взаимодействию с ним, позволяет локализовать или ограничить поиск необходимой информации для определения следующего действия.

Вначале опишем обычный алгоритм планирования в знаковой картине мира [18], а в следующем разделе опишем этап целеполагания, который оказывается одним из этапов этого алгоритма. Мы будем рассматривать случай сим-

вольного планирования, когда задачи по распознаванию объектов и ситуаций внешней среды не стоят, поэтому компоненты образа всех знаков будут пустыми

Процесс планирования в знаковой картине мира реализуется с помощью MAP-алгоритма и идет в обратном направлении: от конечной ситуации к начальной. Кратко опишем основные этапы его работы. На вход алгоритма поступает описание задачи:

$$T = \langle N_T, S, Sit_{start}, Sit_{goal} \rangle,$$

где N_T - идентификатор задачи, S - множество знаков семиотической сети, $Sit_{start} = \langle \emptyset, \emptyset, a_{start} \rangle$ - знак начальной ситуации планирования со смыслом $a_{start} = \{z_{start}^a\}$ и пустыми значением и образом, $Sit_{goal} = \langle \emptyset, \emptyset, a_{goal} \rangle$ - целевая ситуация со смыслом $a_{goal} = \{z_{goal}^a\}$ и пустыми значением и образом.

Значения знаков ситуации пусты так как мы рассматриваем случай, когда с поставленной агенту задачей не ассоциированы обобщенные в группе, к которой принадлежит агент, стандартные действия, т.е. для этой задачи в группе агентов не выработаны схемы решения. В общем случае задача T является результатом процедуры «означивания» - формирования картины мира по исходным описаниям домена планирования D , задающему списки возможных действий и типов объектов, и задачи планирования P , включающей в себя определение стартовых условий и конечной цели (шаг 1 алгоритма 1).

Результатом MAP-алгоритма является $Plan = \{\langle z_{s1}^a, z_{p1}^a \rangle, \langle z_{s2}^a, z_{p2}^a \rangle, \dots, \langle z_{sn}^a, z_{pn}^a \rangle\}$ - после-

Input: описание домена планирования D , описание задачи планирования P , максимальная глубина итераций i_{max}

Output: план $Plan$

```

1:  $T = \langle N_T, S, Sit_{start}, Sit_{goal} \rangle := \text{GROUND}(P)$ 
   //  $N_T$  - идентификатор задачи,  $S$  - множество знаков,  $Sit_{start} = \langle id_{start} \emptyset, \emptyset, \{z_{start}^a\} \rangle$  -
   начальная ситуация со смыслом  $a_{start}$ ,  $Sit_{goal} = \langle id_{goal} \emptyset, \emptyset, \{z_{goal}^a\} \rangle$  - целевая ситуация со
   смыслом  $a_{goal}$ 
2:  $Plan := \text{MAP\_SEARCH}(T)$ 
3: function MAP_SEARCH( $T$ )
4:    $z_{cur} := z_{goal}^a$ 
5:    $z_{start} := z_{start}^a$ 
6:    $Plans := \text{MAP\_ITERATION}(z_{cur}, z_{start}, \emptyset, 0)$ 
7:    $\{Plan_0, Plan_1, \dots\} = \text{SORT}(Plans)$ 
8:   return  $Plan_0$ 
```

Алгоритм 1. MAP-алгоритм планирования поведения: общая схема

довательность длины n пар $\langle z_{si}^a, z_{pi}^a \rangle$, где z_{si}^a - каузальная матрица некоторого узла сети на личностных смыслах, представляющая i -ую ситуацию планирования, а z_{pi}^a - каузальная матрица некоторого личностного смысла, представляющая применяемое в ситуации z_{si}^a действие. При этом ситуация z_{si+1}^a является результатом выполнения действия z_{pi}^a , в том смысле, который раскрывается далее при обсуждении алгоритма, $z_{s1}^a := z_{start}^a$ - каузальная матрица, соответствующая смыслу начальной ситуации, $z_{sn}^a = z_{goal}^a$ - каузальная матрица, соответствующая смыслу целевой ситуации.

Процесс планирования является иерархическим и состоит из повторения МАР-итерации, включающей в себя четыре этапа (Рис. 1).

- S-этап - поиск прецедента совершения действий в текущей ситуации,
- М-этап - поиск применимых действий на сети значений,
- А-этап - генерация личностных смыслов, соответствующих найденным значениям,
- Р-этап - построение новой текущей ситуации по множеству признаков условий найденных действий,

Кратко, МАР-алгоритм осуществляет итеративную генерацию новых каузальных матриц z_{next} личностных смыслов на основе текущей активной матрицы z_{cur} до тех пор, пока не будет достигнуто предельное количество шагов i_{max}

(шаг 10) или не будет целиком активирована начальная матрица z_{start} (шаг 41), соответствующей личностному смыслу a_{start} начальной ситуации. В качестве текущей активной каузальной матрицы для первой итерации выступает матрица, соответствующая личностному смыслу целевой ситуации z_{goal}^a (шаг 4). После завершения выполнения всех итераций, найденные планы сортируются по длине (шаг 7) и самый короткий из них является решением задачи планирования в знаковой картине мира (шаг 8).

Первым этапом в МАР-итерации является S-этап. Его суть состоит в том, что в картине мира интеллектуального агента производится поиск прецедентов, т.е. поиск действий, которые совершались в текущих условиях z_{cur} . Для этого просматриваются все знаки в картине мира S и их личностные смыслы $a(s)$ (шаги 13-16). Если текущие условия z_{cur} удовлетворяются матрицей z_a ($z_a \geq z_{cur}$), то список прецедентов $\hat{A}_{case}$ пополняется результатом распространения активности по сети личностных смыслов от знака s на расстояние d_a (шаг 16). Так как на сети личностных смыслов дуги связывают каузальные матрицы, представляющие действия и объекты, являющиеся участниками этих действий, то в активированные от знака s матрицы попадут действия, которые ранее выполнял агент со знаком s . Это единственный этап глобального поиска в алгоритме МАР, который будет заменен локальным поиском в случае

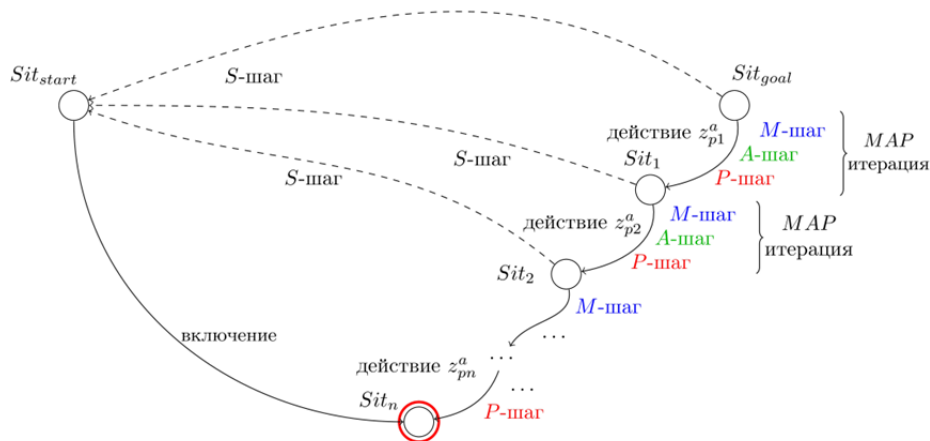


Рис. 1. Общая схема планирования поведения в знаковой картине мира

Двойным кружком отмечено достижение ситуации, включающей в себя стартовую (завершение итераций)

включения полноценного Р-этапа, в котором похожие на z_{cur} ситуации будут искаться в некоторой окрестности образа ситуации z_{cur} (Алгоритм 2).

Далее в МАР-алгоритме следует М-этап, на котором происходит распространение активности по сети личностных смыслов на расстояние d_a с целью активации всех знаков, связанных с текущей ситуацией (шаг 17). Элементы полученного множества каузальных матриц A^* служат отправными точками для распространения активности по сети значений: для каждой матрицы z_a с помощью функции связывания Ψ_a^m определяется необходимый узел на каузальной сети значений, от которого активность распространяется на расстояние d_m (шаг 20). Если активированные матрицы являются каузальными, то они добавляются в множество активных значений M^* (шаг 22). Так как значения задают ролевой состав действий и связывают классы и подклассы, в результате выполнения процедуры $\varphi_m^\uparrow$ в множество M^* попадут действия, которые выполняются с объектом $s(z_a)$ или с его надклассами. Оба шага распространения активности на сети ($\varphi_a^\downarrow$ и $\varphi_m^\uparrow$) задают область локального поиска в картине мира (Алгоритм 3).

Затем переходим к А-этапу, на котором происходит генерация каузальных матриц на сети личностных смыслов, которые представляют специфицированные относительно текущих условий z_{cur} действия, определяемые активными значениями из множества M^* . Для этой цели служат шаги 25-27, в которых распространение активности на каузальной сети значений на расстояние d_m приводит к формированию множества значений $\hat{M}^*$ знаков, связанных с ролевой структурой процедурной матрицы z_m , т.е. включает в себя объекты, которые потенциально могут замещать роли в действиях z_m . Далее, с помощью функции связывания Ψ_m^a происходит генерация новой каузальной матрицы на сети личностных смыслов, которая копирует значение z_m^* с замещением абстрактных знаков-ролей объективными знаками, связанными с ролями отношением класс-подкласс. Затем на

```

9: function MAP_ITERATION( $z_{cur}, z_{start}, Plan_{cur}, i$ )
10:   if  $i \geq i_{max}$  then
11:     return  $\emptyset$ 
12:    $\hat{A}_{case} := \emptyset$  // Список прецедентов
13:   // S-этап
14:   // Поиск прецедентов выполнения действий в текущих условиях
15:   for all  $s \in S$  do
16:     for all  $z_a \in a(s)$  do
17:       if  $z_a \geq z_{cur}$  then
18:          $\hat{A}_{case} = \hat{A}_{case} \cup \varphi_a^\uparrow(s, d_a)$ 

```

Алгоритм 2. МАР-алгоритм планирования поведения: S-этап

```

// M-этап
// Распространение активности вниз по сети личностных смыслов
17:  $A^* = \varphi_a^\downarrow(z_{cur}, d_a)$ 
18:  $M^* = \emptyset$ 
19: for all  $z_a \in A^*$  do
20:   // Распространение активности вверх по сети значений
21:   for all  $z_m \in \varphi_m^\uparrow(s(z_a), d_m)$  do
22:     if  $I^c(z_m) \neq \emptyset$  then
23:        $M^* := M^* \cup \{z_m\}$ 

```

Алгоритм 3. МАР-алгоритм планирования поведения: М-этап

А-шаге происходит отбор тех каузальных матриц, которые представляют действия, выполнимые в текущих условиях z_{cur} (шаги 29-32). Для этого удаляются все каузальные матрицы, эффекты z_{shift} которых не включены в текущую ситуацию $z_a \not\geq z_{shift}$ (напомним, что планирование осуществляется в обратном направлении).

В заключение А-этапа выполняется одна из операций в картине мира θ_a , осуществляющая в данном случае метарегулирование - проверку некоторой эвристики, которая может выражать, например, то правило, что нельзя повторять одинаковые действия, или лучше выполнить то действие, которое быстрее всего приближает к начальным условиям z_{start} (шаг 33). Это позволяет еще больше сократить пространство перебора. Любое эвристическое правило также представимо в виде каузальной матрицы личностного смысла знака, представляющего внутреннюю стратегию планирования своего поведения (Алгоритм 4).

Завершается МАР-алгоритм Р-этапом. Здесь для каждой сгенерированной каузальной матрицы z_a , представляющей некоторое действие, формируется новая ситуация Sit_{next} , которая является результатом обратного применения

```

// А-этап
23:  $\hat{A}_{gen} = \emptyset$ 
24: for all  $z_m \in M^*$  do
// Распространение активности вниз по сети значений
25:  $\hat{M}^* = \varphi_m^{\downarrow}(z_m, d_m)$ 
26: for all  $z_m^* \in \hat{M}^*$  do
27:  $\hat{A}_{gen} := \hat{A}_{gen} \cup \{\Psi_m^a(z_m^*)\}$ 
// Совмещение активности образованных смыслов и текущей ситуации
28:  $\hat{A} = \hat{A}_{gen} \cup \hat{A}_{case}$ 
29: for all  $z_a \in \hat{A}$  do
30:  $z_{shift} = (e_i | i \in I^e)$ 
31: if  $z_{cur} \not\supseteq z_{shift}$  then
32:  $\hat{A} = \hat{A} \setminus \{z_a\}$ 
// Метакогнитивная проверка эвристики
33:  $\hat{A} = \{\theta_a(z_a) | z_a \in \hat{A}\}$ 
34: if  $\hat{A} = \emptyset$  then
35: return  $\emptyset$ 

```

Алгоритм 4. MAP-алгоритм планирования поведения: А-этап

действия в текущих условиях z_{cur} . Обратное применение (шаг 39) заключается в формировании каузальной матрицы z_{next} , состоящей из событий, являющихся либо колонками-условиями действия $e_i \in \{e_k | e_k \in z_a, k \in I^e(z_a)\}$, либо принадлежащих текущей активной каузальной матрице и не являющихся колонками-эффектами действия $e_i \in z_{cur} \wedge e_i \notin \{e_j | e_j \in z_a, j \in I^e(z_a)\}$.

Так как здесь рассматривается случай символического планирования, Р-этап является усеченным и не включает в себя процесс построения образа новой ситуации за счет использования функции связывания Ψ_m^a . Построение образа позволило бы на S-этапе сузить область поиска прецедентов (Алгоритм 5).

В текущий $Plan_{cur}$ добавляется применимое действие $\langle z_{cur}, z_a \rangle$. Если новая ситуация не покрывает стартовую ситуацию (шаг 43), то итерации продолжают с новой текущей ситуацией, пополняя все множество генерируемых планов $Plans_{fin}$.

Константы d_a, d_m , которые определяют глубину распространения активности в каузальных сетях, являются параметрами алгоритма и задают внутреннюю характеристику носителя картины мира, различаясь от агента к агенту. Обычно в модельных экспериментах эти параметры не превышают 5. Увеличение значений этих констант приводит к существенному росту сложности алгоритма MAP, в соответствии с Теоремой 1.

```

// Р-этап
36:  $Plans_{fin} := \emptyset$ 
37: for all  $z_a \in \hat{A}$  do
38:  $Plan_{cur} = Plan_{cur} \cup \{\langle z_{cur}, z_a \rangle\}$ 
// Генерация новой ситуации - применение действия
39:  $z_{next} := (e_i | (e_i \in z_{cur} \wedge e_i \notin \{e_j | e_j \in z_a, j \in I^e(z_a)\}) \vee e_i \in \{e_k | e_k \in z_a, k \in I^e(z_a)\})$ 
40:  $Sit_{next} = \langle id_{next}, \emptyset, \emptyset, \{z_{next}\} \rangle$ 
41: if  $z_{next} \geq z_{start}$  then
42:  $Plans_{fin} = Plans_{fin} \cup \{Plan_{cur}\}$ 
43: else
44:  $Plans_{rec} := MAP\_ITERATION(z_{next}, z_{start}, Plan_{cur}, i + 1)$ 
45:  $Plans_{fin} = Plans_{fin} \cup Plans_{rec}$ 
46: return  $Plans_{fin}$ 

```

Алгоритм 5. MAP-алгоритм планирования поведения: Р-этап

Теорема 1. Сложность итерационного шага *MAP_ITERATION* алгоритма MAP в символьной постановке (алгоритмы 1-5) в худшем случае равна $O(N + C_1^{d_a} C_2^{2d_m})$, где N - количество знаков в картине мира; C_1, C_2 - некоторые константы; d_a, d_m - параметры алгоритма, определяющие глубину распространения активности в каузальных сетях на личностных смыслах и значения, соответственно.

Доказательство. Пусть сложность процедуры распространения активности по каузальным сетям ограничена константой c_ϕ , сложность проверки вложенности каузальных матриц – константой c_{in} , сложность генерации новой матрицы – c_{next} , сложность работы метакогнитивной процедуры θ_a – c_{meta} , сложность работы функций связывания Ψ – c_ψ . Тогда сложность S-этапа (шаги 13-16) ограничена величиной $c_\phi c_{in} N$, сложность M-этапа (шаги 17-22) – величиной $c_\phi N_a^{\downarrow d_a} N_m^{\uparrow d_m}$, где $N_a^{\downarrow}$ – максимальное количество исходящих ребер для узла на сети личностных смыслов, $N_m^{\uparrow}$ – максимальное количество входящих ребер узла на сети значений. Так как мощность множества M^* не более $N_a^{\downarrow d_a} N_m^{\uparrow d_m}$, то сложность A-этапа (шаги 23-32) оценивается как $c_\phi c_\psi N_a^{\downarrow d_a} N_m^{\uparrow d_m} N_m^{\downarrow d_m}$, где $N_m^{\downarrow}$ – максимальное число исходящих ребер для узла на сети значений. Метакогнитивная проверка (шаг 33) прибавляет к сложности величину $c_{meta} N_a^{\downarrow d_a} N_m^{\uparrow d_m} N_m^{\downarrow d_m}$, т.к. мощность множества $\hat{A}$ не больше $N_a^{\downarrow d_a} N_m^{\uparrow d_m} N_m^{\downarrow d_m}$. Наконец, P-этап по сложности ограничен величиной $c_{next} c_\phi N_a^{\downarrow d_a} N_m^{\uparrow d_m} N_m^{\downarrow d_m}$, а величина $N_a^{\downarrow d_a} N_m^{\uparrow d_m} N_m^{\downarrow d_m}$ задает степень ветвления всего рекурсивного алгоритма MAP. Суммируя, получаем, что сложность всего итерационного шага *MAP_ITERATION* равна:

$$\begin{aligned} & O(c_\phi c_{in} N + c_\phi N_a^{\downarrow d_a} N_m^{\uparrow d_m} + c_\phi c_\psi N_a^{\downarrow d_a} N_m^{\uparrow d_m} N_m^{\downarrow d_m} + \\ & c_{meta} N_a^{\downarrow d_a} N_m^{\uparrow d_m} N_m^{\downarrow d_m} + c_{next} c_\phi N_a^{\downarrow d_a} N_m^{\uparrow d_m} N_m^{\downarrow d_m}) = \\ & = O(c_\phi c_{in} N + (c_\phi c_\psi + c_{meta} + c_{next} c_\phi) N_a^{\downarrow d_a} N_m^{\uparrow d_m} N_m^{\downarrow d_m}) \\ & < O(N + C_1^{d_a} C_2^{2d_m}), \end{aligned}$$

где $C_1 = N_a^{\downarrow}$, $C_2 = \max(N_m^{\downarrow}, N_m^{\uparrow})$.

Сложность всего MAP-алгоритма зависит от степени ветвления процедуры *MAP_ITERATION*, которую можно оценить сверху как $N_a^{\downarrow d_a} N_m^{\uparrow d_m} N_m^{\downarrow d_m}$ (мощность множества $\hat{A}$) и, в целом, глубины получающейся рекурсии, которая, в свою очередь, определяется поставленной задачей. Общая сложность алгоритма, таким образом, оценивается как $O((N_a^{\downarrow d_a} N_m^{\uparrow d_m} N_m^{\downarrow d_m})^r (N + C_1^{d_a} C_2^{2d_m}))$, где r – максимальная глубина рекурсии.

Необходимо отметить, что в случае рассмотрения под-этапа по построению образа текущей ситуации на P-этапе вместо полного количества знаков N на S-этапе перебор бы шел только по некоторому подмножеству мощности порядка $N_p^{d_p}$, где N_p – максимальное количество ребер узла на сети образов, а d_p – глубина распространения активности на сети образов.

3. Этап целеполагания

Как уже было сказано в предыдущем параграфе, результатом построения плана является цепочка действий, последовательность каузальных матриц, которые позволяют перейти от начального состояния среды к целевому. Найденная и, возможно, выполненная цепочка действий может быть сохранена в картине мира агента в виде знака $s_{plan} = \{p_{plan}, \emptyset, a_{plan}\}$ со следующими компонентами: образ плана $p_{plan} = \{z_{plan}^p\}$ включает каузальную матрицу z_{plan}^p , состоящую из столбцов, в каждом из которых присутствует только одна ссылка на знак действия, который выполняется на текущем шаге плана, т.е. если $Plan = \{\langle z_{s1}^a, z_{p1}^a \rangle, \langle z_{s2}^a, z_{p2}^a \rangle, \dots, \langle z_{sn}^a, z_{pn}^a \rangle\}$, то $z_{plan}^p = (e_1, e_2, \dots, e_n)$, где n – количество шагов плана, а столбец $e_i = \langle 0, \dots, 0, s_{pi}, 0, \dots, 0 \rangle$ содержит только одну ссылку на знак действия s_{pi} с i -го шага плана. Каузальная матрица личностного смысла плана $a_{plan} = \{z_{plan}^a\}$ содержит две колонки со ссылками на знаки начальной ситуации $s_{s1} = Sit_{start}$ и целевой $s_{sn} = Sit_{goal}$. Таким образом, агент в общем случае в качестве опыта действо-


```

41:   if  $z_{next} \geq z_{start}$  then
42:      $Plans_{fin} = Plans_{fin} \cup \{Plan_{cur}\}$ 
43:     *   if  $\Psi_a^p(z_a) = \emptyset$  then
44:       *     MAP_SEARCH('sub' +  $N_T, S, Sit_{start}, \{\emptyset, \emptyset, \{z_{cur}\}\}$ )
45:   else
46:      $Plans_{rec} := \text{MAP\_ITERATION}(z_{next}, z_{start}, Plan_{cur}, i + 1)$ 
47:      $Plans_{fin} = Plans_{fin} \cup Plans_{rec}$ 

```

Алгоритм 6. GoalMAP-алгоритм планирования поведения: Р-этап

Звездочками обозначены шаги, дополнительные к оригинальному Р-этапу алгоритма MAP

В качестве примера рассмотрим процесс планирования действий по достижению цели «находиться в Санкт-Петербурге» из начальной ситуации «находиться в Москве». На некотором этапе может быть найдена такая цепочка действий «купить билет на поезд», «приехать на Ленинградский вокзал», «ехать на поезде», «прибыть на Московский вокзал», где крайнее действие (с точки зрения обратного планирования) применимо в начальной ситуации, но является схематическим, т.е. известно, что ранее субъект покупал билет на поезд, но конкретный операционный состав этого действия либо не известен, либо не рассматривается в текущем плане и может быть уточнен в результате отдельного процесса планирования с новой целевой ситуацией «иметь билет на поезд».

В отличие от эмпирического, сценарное целеполагание происходит еще до запуска основного алгоритма планирования MAP (соответствующий переход на Рис. 2 отмечен правой жирной линией). В общем случае MAP-алгоритм применяется для начальных и целевых ситуаций, в которых условиях заданы конкретными объектами (определенный кубики “a” и “b”) и специфическими только для данного случая отношениями на их множестве (именно кубик “a” лежит на кубике “b”). Однако в некоторых случаях, планирование

может осуществляться в абстрактном домене, когда, например, условия ситуаций заданы не объектами, а их классами (как переместить множество любых кубиков из одной позиции в другую), а ситуация является стандартной для той группы, к которой принадлежит агент. В таком случае генерация личностных смыслов для определения применимого действия не обязательна и возможно построение новой текущей ситуации по сети значений, т.е. с использованием схемы действия вместо самого действия. Получаемая ситуация становится новой целью, план по достижению которой строится на основе отдельного процесса планирования. Текущий план оказывается состоящим из одного обобщенного действия и, возможно, потребует уточнения, когда очередь дойдет до его исполнения (Алгоритм 7).

Дополнительные шаги в алгоритме GoalMAP (Алгоритм 7) включают в себя просмотр возможных процедурных матриц z_m , которые принадлежат значению знака целевой ситуации Sit_{goal} или каком-либо ее классу (результат распространения активности $\phi_m^{\uparrow}$). Соответствующие значения каузальные матрицы личностных смыслов (специфицированные действия в текущем контексте) проверяются на возможное решение текущей задачи. Так как

```

3: function MAP_SEARCH(T)
4:    $z_{cur} := z_{goal}^a$ 
5:    $z_{start} := z_{start}^a$ 
6:   ** for all  $z_m \in \varphi_m^{\uparrow}(Sit_{goal}, d_m)$  do
7:     **    $z_a = \Psi_m^a(z_m)$ 
8:     **   if  $I^e(z_a) \neq \emptyset \wedge \exists k : k \in I^e(z_a), e_k \in z_a, Sit_{goal} \in e_k$  then
9:       **      $z_{cur} := (e_i | e_i \in z_a \wedge e_i \notin \{e_j | e_j \in z_a, j \in I^e(z_a)\})$ 
10:    $Plans := \text{MAP\_ITERATION}(z_{cur}, z_{start}, \emptyset, 0)$ 
11:    $\{Plan_0, Plan_1, \dots\} = \text{SORT}(Plans)$ 

```

Алгоритм 7. GoalMAP-алгоритм планирования поведения: сценарный случай

Двумя звездочками обозначены шаги, дополнительные к оригинальному алгоритму MAP

целевая ситуация может оказаться стандартной, то могут найтись действия, в которых она участвует в качестве эффектов $Sit_{goal} \in e_k, k \in I^e(z_a)$. Условия найденных действий и составляют новую цель z_{cur} , для которой уже и будет строиться текущий план $MAP_ITERATION$. Необходимо отметить, что ситуация оказывается типичной для группы агентов не только в связи с тем, что в ней участвуют обобщенные понятия, но и тогда, когда часто встречаемая ситуация согласуется в группе агентов и становится для них типичной.

В качестве примера такого типа целеполагания может выступать размышление по достижению цели «стать известным». Субъект, даже не имея опыта выполнения конкретных действий в этом направлении, знаком с принятым в той культурной группе, в которой он принадлежит, сценарием, в котором для того, чтобы быть известным принято становиться писателем, т. е. писать романы или детективы. Таким образом, субъект использует исходную ситуацию «быть писателем» действия «становиться известным» в качестве новой целевой ситуации для нового процесса планирования своего поведения.

4. Модельный эксперимент

Для демонстрации обоих типов целеполагания (эмпирического и сценарного) рассмотрим теоретико-модельные примеры. В первом случае будет продемонстрировано то, как опыт решения задачи сохраняется в схематичном ви-

де и служит основой для запуска процедуры планирования с новой целью. Во втором случае в картину мира агента заранее будет добавлен сценарий достижения цели, применимый в новом задании и используемый для генерации новой цели. Оба модельных эксперимента будут проведены в широко распространенном домене «Мир кубиков», который использовался и для демонстрации работы основного алгоритма MAP [18]. Описание домена на языке PDDL [22] состоит из определения типа «кубик», четырех предикатов («на столе», «свободный кубик», «свободный манипулятор», «держат в манипуляторе»), и четырех действий («поднять», «положить», «поставить», «снять»). Описание домена транслируется в знаковую картину мира агента, в которой для каждого перечисленных компонент создается свой знак, а на сети значений строятся соответствующие каузальные матрицы.

В обоих случаях будем рассматривать задачу построения башни из пяти кубиков A, B, C, D, E, которые изначально находятся на столе. Для начальной и целевой ситуаций создаются соответствующие каузальные матрицы на сети личностных смыслов. На Рис. 3 изображен фрагмент каузальной сети для целевой ситуации (башня ABCDE).

Предположим, что агент имеет ранее полученный опыт планирования достижения целевой ситуации «башня из четырех кубиков ABCD» из целевой ситуации «все кубики на столе». Обычно опыт планирования сохраняется в картине мира агента в виде фрагментов

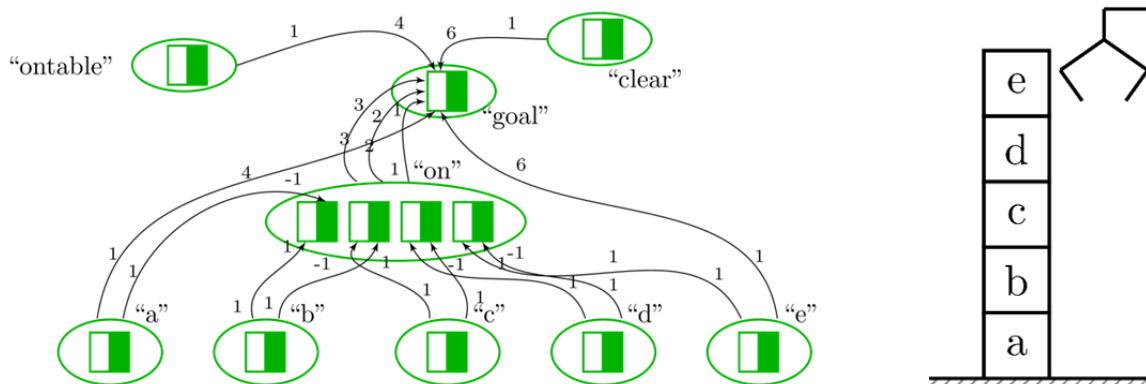


Рис.3. Целевая ситуация задачи планирования «башня из кубиков ABCDE» (слева – фрагмент каузальной сети на личностных смыслах, справа – схема ситуации)

Знак «ontable» означает предикат «на столе», «clear» - «свободный кубик», «goal» - целевую ситуацию.

Цифры при стрелках означают индексы ϵ_1 и ϵ_3 меток дуг каузальной сети (Определение 3)

каузальной сети на образах и личностных смыслах. На сети на образах сохраняется операционный состав действия (какие поддействия входят в действие «построить башню») (Рис. 4). На личностных смыслах сохраняется информация о начальной и конечной ситуациях, для которых был составлен план.

Вернемся к построению плана по достижению новой целевой ситуации «башня ABCDE». При использовании опыта предыдущего планирования агент составит следующий план: «поставить кубик е на d», «поднять кубик е со стола», «построить башню ACBD», где последнее действие представлено новым знаком, сохранившимся в картине мира после построения (и возможного успешного выполнения) плана по достижению ситуации «башня ABCD». Предположим, что, например, в связи с нехваткой памяти, операционный состав этого действия не сохранился, т.е. процедура связывания $\Psi_a^p(z_a)$ в Алгоритме 6 выдает пустое множество. В соответствии с алгоритмом эмпирического целеполагания агент, завершая формирование текущего плана, приступает к планированию достижения новой цели «построить башню ABCD». Если бы агент запоминал весь свой накопленный опыт, то, очевидно, что постановки новой цели в этом случае не потребовалось бы.

Теперь предположим, что действие по построению «башни ABCDE» либо повторялось много раз и было согласовано в некоторой группе агентов, к которой принадлежит носитель картины мира, либо информация об этом действии поступила агенту в виде некоторого сообщения. В этом случае у агента имеется лишь схема такого действия в виде фрагмента каузальной сети на значениях при этом решающего не всю задачу целиком, а начинающаяся с ситуации «башня ABCD» (Рис. 5). Для агента наличие такого сценария означает, что построение «башни ABCDE» из ситуации, когда есть «башня ABCD» гарантировано групповым опытом. Поэтому старая цель считается выполненной, а новой целью становится ситуация «башня ABCD», что и реализует сценарное целеполагание.

Заключение

В статье был представлен оригинальный метод целеполагания, основывающийся на представлениях о работе с целями в психологиче-

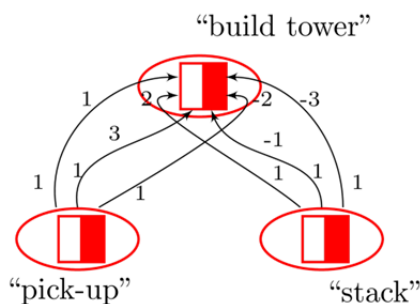


Рис. 4. Фрагмент каузальной сети на образах, демонстрирующий сохранение опыта по построению башни (знак «build tower»)

Знак «pick-up» обозначает действие «поднять»
«stack» - «поставить»

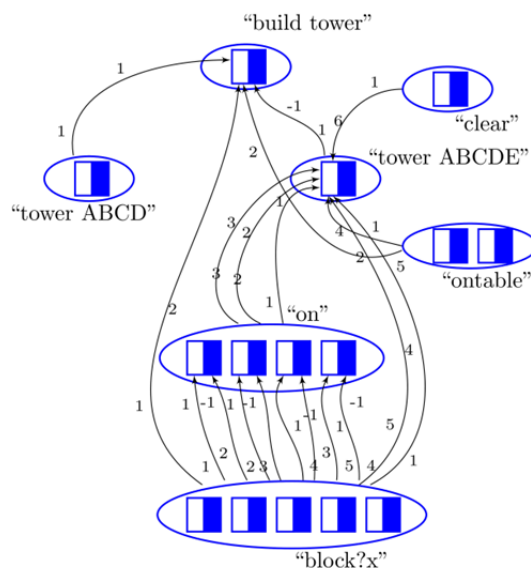


Рис. 5. Фрагмент каузальной сети на значениях, представляющий сценарий достижения цели «башня ABCDE» из ситуации «башня ABCD» и лежащий на столе кубик

ской теории деятельности. Было показано, что некоторые механизмы целеполагания интегрируются в процесс планирования поведения. Был представлен новый алгоритм GoalMAP синтеза плана поведения с этапом целеполагания, основанный на знаковой модели картины мира. Предложенный алгоритм был реализован в виде программной системы, с которой был проведен ряд модельных экспериментов, демонстрирующих основные особенности предложенного метода.

Литература

1. Alford R. et al. Hierarchical planning: Relating task and goal decomposition with task sharing // IJCAI Int. Jt. Conf. Artif. Intell. 2016. P. 3022–3028.
2. Molineaux M., Klenk M., Aha D.W. Goal-Driven Autonomy in a Navy Strategy Simulation // Proceedings of the Twenty-Fourth AAAI Conference on Artificial Intelligence. 2010. P. 1548–1554.
3. Cox M.T. Perpetual Self-Aware Cognitive Agents // AI Mag. 2007. Vol. 28, № 1. P. 32–45.
4. Осипов Г.С. и др. Знаковая картина мира субъекта поведения. М.: Физматлит, 2017. 259 с.
5. Осипов Г.С., Панов А.И., Чудова Н.В. Управление поведением как функция сознания. I. Картина мира и целеполагание // Известия Российской академии наук. Теория и системы управления. 2014. № 4. С. 49–62.
6. Психологические механизмы целеобразования / Под ред. О.К. Тихомиров. М.: Наука, 1977. 260 с.
7. Леонтьев А.Н. Деятельность. Сознание. Личность.. Изд. 2-е. М.: Политиздат, 1977. 304 с.
8. Выготский Л.С. Мышление и речь // Психология развития человека / Под ред. С. Бобко. Эксмо, 2005. С. 664–1019.
9. Muñoz-Avila H. et al. Goal-Driven Autonomy with Case-Based Reasoning // Case-Based Reasoning. 2010. P. 228–241.
10. Roberts M. et al. Iterative Goal Refinement for Robotics // Working Notes of the Planning and Robotics Workshop at ICAPS. 2014.
11. Anderson J.R. How Can the Human Mind Occur in the Physical Universe? Oxford University Press, 2007. 304 p.
12. Kurup U. et al. Using Expectations to Drive Cognitive Behavior // AAAI Conf. Artif. Intell. 2011. P. 221–227.
13. Laird J.E. The Soar Cognitive Architecture. MIT Press, 2012. 374 p.
14. Samsonovich A. V. Goal reasoning as a general form of metacognition in BICA // Biol. Inspired Cogn. Archit. Elsevier B.V., 2014. Vol. 9. P. 105–122.
15. Осипов Г.С., Панов А.И., Чудова Н.В. Управление поведением как функция сознания. II. Синтез плана поведения // Известия РАН. Теория и системы управления. 2015. № 6. С. 47–61.
16. Osipov G.S. Sign-based representation and word model of actor // 2016 IEEE 8th International Conference on Intelligent Systems (IS) / ed. Yager R. et al. IEEE, 2016. P. 22–26.
17. Осипов Г.С., Панов А.И. Отношения и операции в знаковой картине мира субъекта поведения // Искусственный интеллект и принятие решений. 2017. № 4.
18. Panov A.I. Behavior Planning of Intelligent Agent with Sign World Model // Biol. Inspired Cogn. Archit. 2017. Vol. 19. P. 21–31.
19. Kahneman D. Thinking Fast and Slow. New York: Penguin, 2011. 443 p.
20. Kiselev G.A., Panov A.I. Synthesis of the Behavior Plan for Group of Robots with Sign Based World Model // Interactive Collaborative Robotics / ed. Ronzhin A., Rigoll G., Meshcheryakov R. Springer, 2017. P. 83–94.
21. Панов А.И., Яковлев К.С. Взаимодействие стратегического и тактического планирования поведения коалиций агентов в динамической среде // Искусственный интеллект и принятие решений. 2016. № 4. С. 68–78.
22. Gerevini A.E. et al. Deterministic planning in the fifth international planning competition: PDDL3 and experimental evaluation of the planners // Artif. Intell. Elsevier B.V., 2009. Vol. 173, № 5–6. P. 619–668.

Панов Александр Игоревич. Старший научный сотрудник ИСА ФИЦ ИУ РАН. Доцент Высшей школы экономики. Кандидат физико-математических наук. Количество печатных работ: 50. E-mail: pan@isa.ru

Goal setting and behavior planning for cognitive agent

Aleksandr I. Panov

Federal Research Centre "Informatics and Control" of Russian Academy of Sciences, Moscow, Russia

The paper considers a sign-based approach to the problem of modeling the goal-setting process and its integration with the methods of synthesizing the plan of behavior. On the basis of the psychologically plausible model of the sign-based world model of the cognitive agent, the GoalMAP algorithm is proposed, which, based on the formal network model, implements the iterative process of hierarchical planning with the selected stage of setting or selecting a new goal. The estimation of the complexity of the behavior planning algorithm is considered, model experiments with software implementation of the constructed algorithms are performed, demonstrating the key features of the approach used.

Keywords: cognitive agent, sign based world model, sign, GoalMAP, behavior planning, goal reasoning, goal setting, activity theory.

DOI 10.14357/20718594180202

References

1. Alford R. et al. Hierarchical planning: Relating task and goal decomposition with task sharing // IJCAI Int. Jt. Conf. Artif. Intell. 2016. P. 3022–3028.
2. Molineaux M., Klenk M., Aha D.W. Goal-Driven Autonomy in a Navy Strategy Simulation // Proceedings of the Twenty-Fourth AAAI Conference on Artificial Intelligence. 2010. P. 1548–1554.

3. Cox M.T. Perpetual Self-Aware Cognitive Agents // *AI Mag.* 2007. Vol. 28, № 1. P. 32–45.
4. Osipov G.S. Signs-Based vs. Symbolic Models // *Advances in Artificial Intelligence and Soft Computing Lecture Notes in Computer Science.* / под ред. G. Sidorov, S.N. Galicia-Haro. : Springer International Publishing, 2015. С. 3–11.
5. Osipov G.S., Panov A.I., Chudova N. V. Behavior control as a function of consciousness. I. World model and goal setting // *J. Comput. Syst. Sci. Int.* 2014. T. 53. № 4. С. 517–529.
6. Psychological mechanisms of goal formulation / pod red. O.K. Tihomirov. M.: Izdatel'stvo «Nauka», 1977. 260 s. (In Russian).
7. Leontyev A.N. The Development of Mind. Kettering: Erythros Press and Media, 2009. 428 c.
8. Vygotsky L.S. Thought and Language. : MIT Press, 1986. 344 c.
9. Muñoz-Avila H. et al. Goal-Driven Autonomy with Case-Based Reasoning // *Case-Based Reasoning.* 2010. P. 228–241.
10. Roberts M. et al. Iterative Goal Refinement for Robotics // *Working Notes of the Planning and Robotics Workshop at ICAPS.* 2014.
11. Anderson J.R. How Can the Human Mind Occur in the Physical Universe? Oxford University Press, 2007. 304 p.
12. Kurup U. et al. Using Expectations to Drive Cognitive Behavior // *AAAI Conf. Artif. Intell.* 2011. P. 221–227.
13. Laird J.E. The Soar Cognitive Architecture. MIT Press, 2012. 374 p.
14. Samsonovich A. V. Goal reasoning as a general form of metacognition in BICA // *Biol. Inspired Cogn. Archit.* Elsevier B.V., 2014. Vol. 9. P. 105–122.
15. Osipov G.S., Panov A.I., Chudova N. V. Behavior Control as a Function of Consciousness. II. Synthesis of a Behavior Plan // *J. Comput. Syst. Sci. Int.* 2015. T. 54. № 6. С. 882–896.
16. Osipov G.S. Sign-based representation and word model of actor // *2016 IEEE 8th International Conference on Intelligent Systems (IS)* / ed. Yager R. et al. IEEE, 2016. P. 22–26.
17. Osipov G.S., Panov A.I. Relationships and operations in sign-based world model // *Iskusstvennyj intellekt i prinyatie reshenij.* 2017. № 4. (In Russian)
18. Panov A.I. Behavior Planning of Intelligent Agent with Sign World Model // *Biol. Inspired Cogn. Archit.* 2017. Vol. 19. P. 21–31.
19. Kahneman D. Thinking Fast and Slow. New York: Penguin, 2011. 443 p.
20. Kiselev G.A., Panov A.I. Synthesis of the Behavior Plan for Group of Robots with Sign Based World Model // *Interactive Collaborative Robotics* / ed. Ronzhin A., Rigoll G., Meshcheryakov R. Springer, 2017. P. 83–94.
21. Panov A.I., Yakovlev K.S. Psychologically Inspired Planning Method for Smart Relocation Task // *Procedia Comput. Sci.* 2016. T. 88. С. 115–124.
22. Gerevini A.E. et al. Deterministic planning in the fifth international planning competition: PDDL3 and experimental evaluation of the planners // *Artif. Intell.* Elsevier B.V., 2009. Vol. 173, № 5–6. P. 619–668.

Panov Aleksandr I. Senior Researcher of Federal Research Centre "Informatics and Control" of Russian Academy of Sciences. PhD in Computer Science. Number of publications: 50. E-mail: pan@isa.ru, apanov@hse.ru