



NATURA CUPIDITATEM INGENUIT HOMINI VERI VIDENDI
Marcus Tullius Cicero
(Природа наделила человека стремлением к познанию истины)

Мысли Об Истине

Альманах «**МОИ**»
Электронное издание, ISBN 9984-688-57-7

Альманах «Мысли об Истине» издается для борьбы с лженаукой во всех ее проявлениях и в поддержку идей, положенных в основу деятельности Комиссии РАН по борьбе с лженаукой и фальсификацией научных исследований. В альманахе публикуются различные материалы, способствующие установлению научной истины и отвержению псевдонаучных заблуждений в человеческом обществе.

Альманах издается с 8 августа 2013 года
Настоящая версия тома выпущена **2017-09-10**

© 2017 Марина Ипатьева (оформление и комментарии)

Эгле В. и Марьясов С. Переписка о Теории интеллекта

VEcordia

Извлечение R-POTI-3

Открыто: 2011.03.20 21:33

Закрето: 2011.08.22 11:41

Версия: 2014.04.22 14:10

Дневник «VECORDIA»

ISBN 9984-9395-5-3

© Valdis Egle, 2014

Переписка о теории интеллекта, том 3
ISBN

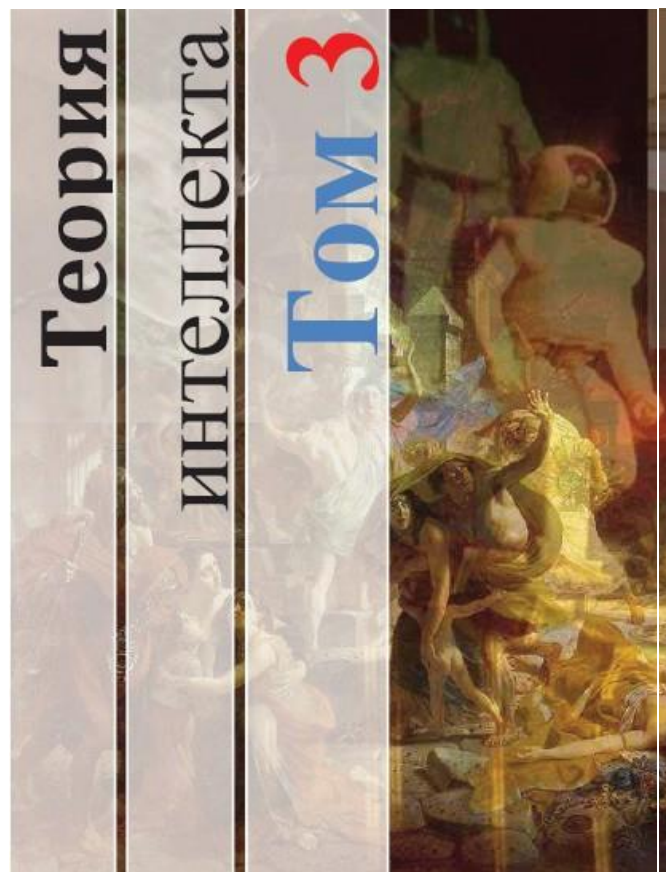
© Эгле В., Марьясов С., 2011

Валдис Эгле и Сергей Марьясов
ПЕРЕПИСКА о Теории Интеллекта
адресованная Комиссии РАН по борьбе с
лженаукой и фальсификацией научных ис-
следований

начата 7 июня 2010 года

Impositum

Grīziņkalns 2014



Не обгорят рябиновые кисти,
От желтизны не пропадет трава.
Как дерево роняет тихо листья,
Так я роняю грустные слова.

*Сергей Есенин,
«Рябиновый костер», 1924*

Третья переписка в ПОТИ

Глава 1. Как завязалась эта переписка

Она началась 24 февраля 2011 года с приведенного ниже письма, за которым последовали и дальнейшие.

§1. Первые письма

no Галина Сухова <galina.suhova@gmail.com>
kam valdis.egle@gmail.com, TALEO6@inbox.lv
datums 2011. gada 24. februāris 12:01
temats ask friend
piegādātājs gmail.com
parakstītāja gmail.com

Валдис, наш приятель Сергей, который уехал в Америку, вдарился в философию. Сам кончал вроде как физмат. Это Сашин приятель – работал в банке и был директором Карточного центра. На мой взгляд человек флегматичный, но с сангвинистическими нотками – контактирует легко. Кстати, ты же его должен помнить – ты преподавал ему латышский.

В общем, он решил составить свое представление об устройстве мира – стал читать, что попадется об этом, и призапутался, когда дошел до Блаватской. Начал со мной беседы вести на эту тему, но я ведь имею весьма поверхностные познания (только имена да общие слова), так что я пообещала ему, что спрошу у тебя совета как у ЯРОГО МАТЕРИАЛИСТА. Пусть сначала почитает тебя, а потом уже всяких идеалистов.

То есть, у меня просьба: пришли мне (я не знаю его майл – мы общаемся в скапе – потом спрошу и дам тебе для прямого сообщения) последовательные ссылки на тему устройства мира (модель, строение, какую часть занимает душа), чтобы ты посоветовал ему почитать.

Может только не стоит начинать с пенхауса¹ – может сначала свои работы и последовательное развитие твоих мыслей, а потом и до вашего Хауса доберетесь.

Спасибо заранее. Можно ли ему сразу дать твой майл?

¹ В.Э.: Видимо, имеется в виду Роджер Пенроуз.

no Valdis Egle <valdis.egle@gmail.com>
kam Галина Сухова <galina.suhova@gmail.com>
datums 2011. gada 24. februāris 15:39
temats Re: ask friend
piegādātājs gmail.com

Ну разумеется, я помню Сергея, только фамилию сейчас не помню: была какая-то нестандартная на «М» (надо бы посмотреть его визитную карточку, только долго искать, не знаю, где они лежат).

Вообще, если он хочет «пофилософствовать», то я приглашаю его в дискуссию ПОТИ на сайт <http://ve-poti.narod.ru/>. Там в книгах {VIEWS = МОИ № 100} и {DVESA = МОИ № 50} можно увидеть также и основы моего понимания «устройства мира». Там на сайте виден также и мой адрес e-почты. Но, конечно, вся переписка должна быть публичной, она будет помещаться на сайт, и к своим словам надо подходить с ответственностью.

Блаватскую (и биографию, и «учение») я знаю хорошо, потому что в 1999–2003 гг. у меня была одна очень активная читательница, которая до этого 10 лет провела среди «эзотериков», а потом всё допытывалась, чтобы я объяснил ей всё, с Блаватской связанное, и в результате добилась, чтобы я всё это изучил.

Елену Блаватскую я теперь считаю гениальной – но гением ИЗДЕВАТЕЛЬСТВА. Всё, что она делала и писала, было сплошным издевательством над дураками (или, если помягче, то – розыгрышем простачков). Она начала в 1848 году, в бум спиритизма, с издевательств над спиритистами, и розыгрыши и насмешки проходят через всю ее жизнь, но вершиной ее издевательств можно считать то, что она Учителем ее личным и всего человечества объявила Морию (!). (А Мория – по-гречески означает Глупость: главный персонаж сочинения Эразма Роттердамского «Похвала Глупости»). И вот – даже сегодня еще миллионы одураченных Блаватской простачков на полном серьезе поклоняются «Листьям Сада Мории» (т.е. – Царству Глупости).

В истории человечества было много шутников подобного рода (и Орден Розенкрейцеров придумывали, и многое другое), но, во-первых, большинство потом признавались в своих розыгрышах, – а Блаватская не призналась (лишь очень прозрачно намекала). А во-вторых, никому из них не удавалось достигнуть ТАКОГО масштаба, известности и количества последователей, как у Блаватской.

В.Э.

от Galina.Suhova@gmail.com
кому sevm@rambler.ru
дата 24 февраля 2011 г. 22:25
тема Fwd: Re: ask friend
отправлено через gmail.com

Сергея,
вот тебе ответ Валдиса – давай читай и общайся с ним. Я тоже посмотрю.

GS

no sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
kopija Galina.Suhova@gmail.com
datums 2011. gada 7. marts 06:08
temats Путь идей
piegādātājs rambler.ru

Валдис, добрый день!

С удовольствием погрузился в чтение твоей книги «Путь идей». Давно не читал ничего с таким интересом и практически использовал всё своё свободное время для этого.

Остаётся только сожалеть, что твоя книга была недоступна для широкого круга читателей в 1980–81 годах, когда я учился в Латвийском университете и в течение целого семестра изучал курс «Диалектический материализм». Мои отрывочные знания о философии вызывали интерес к курсу. Хотелось понять, что же это за такая наука философия. Однако надежды на получение

систематизированных знаний рассыпались почти сразу же после обсуждения «главного вопроса философии».

Прошло много лет, а я хорошо помню семинар, на котором обсуждалась «идеальная» природа марксистско-ленинской мысли. Аудитория в большинстве своём состояла из вчерашних десятиклассников и активно апеллировала к материалистическому пониманию природы мысли, базируясь на знаниях, полученных на школьных уроках и подчеркнутых в научно-популярной литературе. В понимании будущих физиков преподаватель не защищал свой тезис, а просто любые примеры аудитории путём неуклюжих манипуляций (весьма далёких от «манипуляций фокусника») превращал в «идеальность мысли». Честно говоря, тогда мне было очень неловко за нашего преподавателя. Естественно, как образованный человек, он не мог сказать, что на уроках биологии нас учили неправильно, но, защищая честь мундира, он не мог признать неверность тезиса и пытался внушить нам (как это часто потом происходило в жизни с руководителями, особенно партийными), что мы якобы чего-то там недопонимаем и т.д...

Ещё в курсе «Диалектический материализм» меня, и наверное не только меня, поразило огромное количество «воды», как если бы кто-то, например, теорему Пифагора излагал на 20 страницах, и далее ещё пять или шесть относительно полезных утверждений, вытекающих из этой теоремы, размазал по двумстам страницам учебного пособия. В этой воде спорность «решения основного вопроса философии» утонула окончательно ☺. М.б. в этом и была основная задача курса?

Книга «Путь Идей» отличается чёткой структурой и ясностью излагаемого материала, что вызывает желание подумать самостоятельно. Так, я решил посмотреть, что есть в интернете об отличиях в ДНК разных народов и наткнулся на книгу Л.Н. Гумилёва «От Руси до России» и его теорию пассионаризма. В частности, Гумилёв в предисловии связывает активность в науке и искусстве с проявлением пассионарности...

В результате мне захотелось посмотреть, над чем ты работаешь в настоящее время и, честно говоря, я несколько увлёкся теперь уже «Теорией интеллекта».

Должен подвести некоторый промежуточный итог. Спасибо, с книгами Елены Блаватской я расстался. Хотя от мысли, что где-то на Земле хранится или в прошлом могло храниться некоторое «абсолютное знание», привнесённое извне, отказаться совсем пока не смог.

Но и по этому поводу дискутировать пока не хочется. Думаю, что сначала дочитаю до конца твою книгу «Путь идей». И тогда попробую сформулировать вопросы, которые толкнули меня на чтение книг Елены Блаватской, а затем и твоих работ.

Спасибо за развёрнутый ответ о работах Елены Блаватской. Общее впечатление от того, что я успел прочитать в её книге, в принципе совпадает с твоими выводами.

С наилучшими пожеланиями,

Сергей Марьясов.

2011.03.08 7:51, вторник

Здравствуй, Сергей! Спасибо за письмо, и рад, что ты принял мое приглашение участвовать в переписке ПОТИ. Как видишь, на сайте этой переписки уже сейчас довольно много различных материалов (одного Пенроуза можно обсуждать годами). Когда ты затронешь конкретные вопросы, то можно будет о них поговорить.

А сейчас, пользуясь случаем, некоторая, так сказать, «общая ориентировка». Издание «Векордия», начатое в 2006 году в канун моего 60-летия как чисто электронно-интернетовское взамен предыдущим бумажным изданиям, таким как «Ведда», «Сидиоуэм» и другим, представляет собой нечто вроде моего «полного собрания сочинений», которое я хочу оставить после себя, и рассчитано оно больше на посмертное, чем прижизненное чтение, что в более или менее завуалированной форме выражает нижняя надпись на титульном листе каждого тома Векордии.

А верхняя надпись (тоже в скрытой форме) означает: «Дуракам не читать!». (Впервые эта надпись применялась еще в советское время в «Сидиоуэме» и тогда она пародировала стоящий на каждой газете и журнале лозунг; теперь это звучит уже не так, но я решил всё же сохранить в Векордии эту давнюю традицию моих изданий, зародившуюся в другую эпоху).

Книга «Путь Идей» {VIEWS = **МОИ** № 100} была написана, во-первых, для публики типа Галины Васильевны, а, во-вторых, тоже очень очень давно, когда даже и Веданская теория еще не зародилась, и теперь она даже и мне самому кажется юношеской. Но в «полном собрании

сочинений», конечно, должна входить, и те пути, по которым я приближался к Веданской теории, она все-таки иллюстрирует (как и книга «Двеса» {DVESA = МОИ № 50}).

Что же касается собственно дискуссии ПОТИ, то моя главная задача в ней: запротokolировать, зафиксировать, засвидетельствовать «всё, что было». Я уже не рассчитываю, что мне удастся перевернуть математику (и другие науки) в соответствии с постулатами Веданской теории. Я не могу заставить людей мыслить логически. Но я могу оставить после себя документ, в котором зафиксировано, что они мыслили нелогично. Что я и сделаю.

Я практически не боюсь, что они опровергнут Веданскую теорию. Не смогут они этого сделать. Никаких ошибок в Веданской теории нет, и она представляет собой последовательно научный подход к проблеме интеллекта – результат минимизации постулатов. И никакой доктор наук, никакой профессор или академик ничего не может в этом изменить, как бы ему этого ни хотелось. Легко защищать правильный вывод, а вот неправильный – очень трудно. Так что мне легче, чем им.

Хорошо, – этого для первого раза хватит; читай дальше и задавай конкретные вопросы.

В.Э.²

no sevm@rambler.ru
kam Галина Сухова <galina.suhova@gmail.com>
kopija Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 8. marts 00:39
temats Re: Путь идей
piegādātājs rambler.ru

Галя, привет!

Это тебе спасибо за то, что подсказала, или вернее направила меня к Валдису. Есть мне теперь над чем подумать ☺.

Мне его книги интересны. Надо сказать, что их у него написано немало, и это радует. Значит, впереди ещё немало интересных встреч... Читая его книгу «Переписка о Теории интеллекта», по ссылке о том, что такое «интуиция», перескочил на его комментарии к книге Роджер Пенроуз «НОВЫЙ РАЗУМ КОРОЛЯ»³, и похоже, что это тоже для меня интересно.

Но, продолжая наш разговор о времени, увлечениях и модном чтиве, должен заметить, что не всем будут интересны работы Валдиса. Люди тратят своё свободное время (оставшееся от обязательных дел) в соответствии с своими предпочтениями, а они настолько различны! И очень редко лежат в области философии и искусственного интеллекта. И мне кажется, это это нормально.

Я, конечно же, попытаюсь порекомендовать работы Валдиса (те что уже читаю) своему племяннику и м.б. ещё нескольким знакомым молодым людям, но в какой области лежат их интересы, как глубоко они пытаются разобраться в той информации, которую им дают в высших учебных заведениях, мне не известно. Но так у них будет шанс, которого у меня не было в 1981 году ☺.

С наступающим праздником!

Сергей Марьясов.

P.S. Поскольку речь шла о книгах Валдиса Эгле, и мы немного знакомы лично, пошлю копию и ему. Пусть знает, по крайней мере, что у его работ появился ещё один почитатель ☺. Надеюсь, что он не отнесётся к этому, как к спаму ☺☺.

² В.Э.: Это письмо было написано утром 8 марта, но не было отправлено, так как при открытии почтовых ящиков обнаружилось приведенное ниже переправленное письмо, с которым я не знал, что делать, а также другие письма от других людей, которые привели меня в смятение и к желанию пока приостановить все переписки, и тогда я (в 14:02) ограничился просто вопросом о статусе переправленного письма.

³ В.Э.: Это ссылка в {ПОТИ-1 = МОИ № 41} на {PENRO1 = МОИ № 14}.

no Valdis Egle <valdis.egle@gmail.com>
kam sevm@rambler.ru, galina.suhova@gmail.com
datums 2011. gada 8. marts 14:02
temats Re: Путь идей
piegādātājs gmail.com

Спасибо за письма, ответ готовлю, но пока что у меня один маленький организационный вопрос: каков статус тех писем, которые мне присылаете как копии? Я их тоже могу публиковать – или это непубликуемая личная переписка?

Галине поздравления с 8 марта!

В.Э.

no sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 9. marts 05:12
temats Re: Путь идей
piegādātājs rambler.ru

По «административной» привычке стараюсь посылать копии всем, о ком в письме идёт речь, или кому содержание письма может быть интересно/полезно. Поэтому, Валдис, вы вправе опубликовать весь текст целиком или с сокращениями.

В свою очередь, у меня тоже есть вопрос. Я заметил в тексте несколько опечаток. Имеет ли смысл указать Вам их для исправления? Вопрос об исправлении опечаток не такой простой: как известно, одиночные опечатки не мешают чтению и пониманию текста. У одного новоявленного *entrepreneur*-а в книге о том, как успешно начать и развить новый бизнес, я вычитал, что ему дешевле выпустить книгу с некоторым количеством опечаток, чем оплачивать профессионального корректора, а идеи, изложенные в книге, от этого вроде как и не пострадают. Утверждение спорное, но я прекрасно понимаю, сколько времени требует коррекция одиночных опечаток.

СМ

no sevm@rambler.ru
kam Valdis.egle@gmail.com
datums 2011. gada 15. marts 18:24
temats R-POTI-2, «золотая теорема»
piegādātājs rambler.ru

Валдис, добрый день!

Пишу на электронную почту, т.к. ссылка на Форум Векордии не работает (удалена владельцами сайта)⁴, а Гостевая книга в моём понимании, это односторонняя связь от посетителя к посещаемому, кроме того, сообщение от 27.01.2009 не удалено (я бы определил его как спам и удалил, т.е. м.б. не просматривалось⁵?)

С большим интересом прочитал книгу R-POTI-1 [= МОИ № 41]. Появилось несколько вопросов / комментариев, но пришло и понимание, что тематика, затронутая в книге, гораздо сложнее, и цели, которые Вы ставите, гораздо важнее ответов на мои вопросы, которые я бы определил, как учебные. Поэтому я решил (пока:) не злоупотреблять Вашим вниманием.

Тем не менее, не могу удержаться от ещё одного комментария. Добрался до теоремы «золотого сечения» (R-POTI-2 [= МОИ № 42] стр.15), и, честно говоря, «забуксовал». Про эту теорему из университетских времён ничего не помню, м.б. и не проходили. Поэтому самостоятельно попытался понять её содержание. Вот что я прочитал:

«...Доказательство. Обозначим через Y^X множество всех отображений.
Мы должны доказать:

⁴ В.Э.: Да, *Narod.ru* закрыл форумы. Но у меня там особого обсуждения и не было.

⁵ В.Э.: У меня на тех сайтах стоит премодерация сообщений, т.е. нельзя сообщения поместить без моего ведома. Тот «спам» я два года назад почему-то пропустил (разрешил поместить). Видимо, они рекламировали что-то такое, что я решил поддержать. Мне постоянно (каждые несколько дней) помещают рекламу (московских) проституток – ту я не пропускаю. Лезут также японцы со своими иероглифами – вообще не знаю, о чем там у них речь (среди иероглифов иногда попадаются английские аббревиатуры, и тогда, судя по ним, – вроде о каких-то программах).

- 1) Существует 1–1 отображение X на подмножество Y^X .
- 2) Не существует 1–1 отображения X на всё Y^X ...»

Форма описания в 1) «подмножество Y^X » с точки зрения языка правильная, но, тем не менее, для человека, читающего текст впервые, вызывает ощущение, что в обозначения могла вкратце ошибка: Y^X – это «множество всех отображений», или это «подмножество»? Особенно, когда в 2) детализируется «на всё Y^X ». Т.е. в 1) используется определение Y^X из предыдущей строки, а в 2) это определение вдруг уточняется, будто бы м.б. ещё какое-то другое Y^X , кроме как «всё» (как, например, «подмножество Y^X » в 1).

Спасибо за то, что на стр.30 Вы вернулись к обсуждению этой теоремы и тут формулировка её оказалась более «удобоваримой»:

«... Доказательство. Обозначим через Y^X множество всех отображений множества X в множество Y .

В соответствии с определением неравенства мощностей мы должны доказать два утверждения:

- 1) Существует взаимно однозначное отображение множества X на **некоторое подмножество множества Y^X** .
- 2) Не существует взаимно однозначного отображения множества X на **всё множество Y^X** ...»

Моё внутреннее раздражение снято! (красным я выделил то, что мне не хватало для чёткости понимания текста).⁶

К сожалению, по этой причине я всегда недолюбливал математику. Такое впечатление (так мне кажется теперь, особенно после прочтения R-POTI-1), что когда люди пишут (или говорят) о вещах, им хорошо понятных, они забывают, что в общении с другими людьми надо употреблять максимально точные языковые конструкции. Особенно этим недостатком грешили учебники математики, и без пояснения преподавателя формулировки оставались совершенно нечеткими.

Интересно и другое. Обе стороны вступили в спор, не обращая внимания на сжатую форму формулировки теоремы. Т.е. вступили в спор с позиций, которые давно уже сложились у них в головах о содержании теоремы и её выводах, особенно не погружаясь в точность формулировок (языковых конструкций). Естественно, что при таком подходе путь к согласованному (тождественному) пониманию будет долгим. Но впрочем, на это, Валдис, именно Вы и указываете, объясняя природу непонимания Вашей теории другими специалистами. Когда «парадигма» сформирована, изменить мнение человека очень трудно. Один раз некоторое утверждение как-то состыковалось (у кого с большим трудом, у кого с меньшим), и каждый раз перестраивать эту «мозговую программу» (в Вашей терминологии) – это дополнительные усилия/затраты, которые непонятно что дают с точки зрения индивидуума, и естественная защита (программа защиты) от цикла (опять же Ваш пример R-POTI-1 [= **МОИ № 41**] стр.53) быстро снимает с выполнения программу перестыковки уже сложившихся понятий.

Наверное где-то тут рядом лежат принципы формирования идеологии отдельно взятого человека и влияние идеологии на процессы познания в целом и ответ на вопрос, почему открытия/изобретения не случаются каждый день. В списке книг видел, что у Вас есть книги, в которых, возможно, есть ответы и на эти вопросы, так что собираюсь прочесть и их.

С наилучшими пожеланиями,

Сергей Марьясов.

⁶ В.Э.: Когда я готовил эту часть книги POTI-2, у меня не было книги Александра, поэтому я не мог занести теорему прямо из книги, а вынужден был пользоваться своими собственными записями (30-летней давности). Я нашел теорему в лекции, которую готовил в 1981 году для Университета, а прочел тогда только в Институте электроники. Но, спустя 30 лет, я не вспомнил, что для той лекции теорема была выписана на плакатике в сокращенном виде: только узловые моменты, и предназначалось это для устного рассказа, когда я говорю и показываю указкой на то или иное место. Потому и получился такой слишком сжатый вариант. А потом я нашел полный текст теоремы у себя в одном конспекте и тогда повторил ее еще раз в книге POTI-2. (А еще потом нашел в Интернете и саму книгу Александра в формате Djvu и прикачал ее, так что теперь у меня есть всё).

no soukhova <natalja.soukhova@free.fr>
kam galina suhova <g_suhova@yahoo.com>, asmelov52@mail.ru, valdis.egle@gmail.com
datums 2011. gada 18. marts 15:36
temats Fwd: interesanti

Интересный факт, не могла не переслать, все-таки год кризиса ☺☺.

Nezinu, kas tas ir bet matemātika ir cool.⁷

Интригующие новости из Китая. В этом году мы собираемся испытать четыре необычные даты: 1.1.11; 1.11.11; 11.1.11; 11.11.11 и это ещё не всё...

Возьмите последние две цифры года, в котором вы родились, теперь добавьте ваш возраст этого года, и результат будет 111 для всех!!! Например, Гарри родился в 1957 году, в этом году ему исполняется 54 года: Итак $57 + 54 = 111$. Ольга родилась в 1974 году, и в этом году ей исполнится 37 лет: $74 + 37 = 111$. Как Вам это нравится⁸? Согласно китайскому фен-шуй – это год денег.

В этом году есть знаменательный месяц октябрь. Он будет иметь 5 воскресений, 5 понедельников и 5 суббот. Такое происходит один раз в 823 года.⁹ Именно эти годы называют «денежными мешками». Если вы отправите сегодня это письмо 8 хорошим друзьям, то, как гласит и обещает китайский фен-шуй, в ближайшие четыре дня у вас появятся деньги. Не прерывайте цепочку хорошей новости, перешлите её другим, и пусть всем тоже будет ХОРОШО!

no sevm@rambler.ru
kam galina suhova <g_suhova@yahoo.com>
kopija natalja.soukhova@free.fr, asmelov52@mail.ru, valdis.egle@gmail.com
datums 2011. gada 18. marts 17:08
temats Re: Fwd: interesanti
piegādātājs rambler.ru

Привет!

Я вам тоже желаю денег побольше и купюрами покрупнее ☺.

А сам я наткнулся на такую вот статью www.planetfuture.info. Много про закон Дарвина, искусственный интеллект, «в чём смысл жизни» и ничего про фен-шуй ☺☺.

Сергей Марьясов

no Valdis Egle <valdis.egle@gmail.com>
kam sevm@rambler.ru
datums 2011. gada 21. marts 14:37
temats Предварительный ответ
piegādātājs gmail.com

Здравствуй, Сергей!

Я думаю, мы можем остаться на «ты», как в первом письме.

Тогда, 8 марта, я начал писать ответ тебе с описанием некоторых общих положений Векордии, но, во-первых, я был очень измотан длительной подготовкой перед этим девяти томов Пенроуза, первых томов книг РОТІ и первой частью переписки с Маниным, и, во-вторых, к тому

⁷ В.Э.: «Не знаю, что это такое, но математика – cool» (по-латышски и на английском жаргоне; две данные зеленым цветом строки написаны двумя разными авторами, не совпадающими с автором собственно пересланного письма, но не знаю, кем именно; одна из строк, видимо, Наташи Суховой).

⁸ В.Э.: Возраст человека на любой год X есть X минус P (где P – год его рождения). Поэтому, если прибавить к P возраст, получишь обратно год X. Всегда. В любом году.

⁹ В.Э.: Козлы! (Это относится к авторам первоначального письма, утверждающим, что «*Такое происходит один раз в 823 года*»). Месяц октябрь имеет 31 день, т.е. 3 дня сверх 28-дневного, 4-недельного цикла. Поэтому октябрь (как и январь, март, май, июль, август и декабрь) каждый год имеет по пять каких-то подряд идущих дней недели (те три дня, с которых месяц начинается, повторяются 5-й раз в его конце). В этом (2011) году октябрь начинается с субботы (как и январь), поэтому в 2011 году оба эти месяца имеют по пять суббот, воскресений и понедельников. Предыдущий раз октябрь начинался с субботы (и, следовательно, имел такую же ситуацию) в 2005 году. Следующий раз октябрь начнется с субботы (и будет иметь 5 суббот, воскресений и понедельников) в 2016 году.

же в этот момент нагрянули угрозы от латвийских «ученых» подать на меня судебный иск за «оскорбление их чести и достоинства» в моих интернетовских сайтах, а Валентине (жене моей) пришло письмо от Налоговой службы Латвии с требованием уплатить большую неожиданную сумму в налогах (с чем, разумеется, разбираться приходилось мне). В результате у меня получилось нечто вроде «психологической перегрузки», и я на 10 дней полностью отключился от всех дел по переписке ПОТИ: не подходил к Интернету, не открывал почтовые ящики, а то время, которое оставалось от разбора налоговых деклараций, читал стихи Пушкина, Лермонтова, Есенина и смотрел фильмы. Словом – устроил себе отпуск.

Извини, что это получилось именно в тот момент, когда ты вошел в эту переписку, но это никак не было связано с тобой. Просто я «дошел до ручки» и не мог думать ни о каких переписках.

Когда же я вышел из «отпуска», то первым делом разобрался с Маниным (три дня писал ему ответ); в принципе хоть сейчас могу выставить это в Интернет (в книге {ПОТИ-2 = МОИ №42}), но немножко подожду и выставлю вместе с ответами тебе. Для переписки с тобой я (вчера) открыл новую книгу (ПОТИ-3), в которой – если ты не возражаешь, конечно, – авторы будут «Валдис Эгле и Сергей Марьясов». Там тогда и можно будет обсуждать разные вопросы, более или менее связанные с проблемой интеллекта, а ПОТИ-2 пусть останется для «чистого Кантора».

Сегодня я, наконец, открыл тот почтовый ящик, на который писал ты (с Маниным переписка идет в другом ящике) и увидел, что ты мне тут понаслал уже очень много. Еще раз извини, что я так поздно за это взялся. Я еще не читал это внимательно – вот, буду переносить в книгу, разбирать и оформлять, тогда и изучу и отвечу.

Увидел только, что ты спрашиваешь: надо ли отмечать опечатки. РАЗУМЕЕТСЯ, НАДО!

Я буду очень благодарен тебе, если ты заведешь какой-нибудь *Word* файл и скопируешь туда ошибочные фразы с отметкой, из какой это книги (с таким расчетом, чтобы в фразе было достаточно характерных слов, по которым эту фразу можно в книге найти *Find*-ом), а собственно ошибки пометишь красным. Так накопится список ошибок, а через некоторое время (скажем, раз в месяц) пришлешь этот файл мне, и я исправлю, что надо. (Так работали и предыдущие читатели из тех, что со мной переписывались, и очень помогли улучшить качество текстов).

Если ты не возражаешь против того, что я предложил о книге ПОТИ-3, то было бы желательно, чтобы ты отобрал какую-нибудь свою фотографию, которую поместить в книге. (Моих фотографий в Векордии предостаточно, и надо было бы иметь портрет и второго автора).

Пока всё – начну оформлять книгу ПОТИ-3.

В.Э.

no sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 22. marts 16:33
temats Re: Предварительный ответ
piegādātājs rambler.ru

Здравствуй, Валдис!

Честно говоря, зная твою обязательность, уже начал немного беспокоиться, и, как вижу из твоего письма, не зря. Если нужен знающий и надёжный юрист для разбирательства с латвийскими «учеными», могу переговорить с моим хорошим товарищем – профессиональным юристом.

Фотографию выбираю ☺.

С.М.

no Valdis Egle <valdis.egle@gmail.com>
kam sevm@rambler.ru
datums 2011. gada 22. marts 18:22
temats Re: Предварительный ответ
piegādātājs gmail.com

Спасибо, – еще же иска нет – есть только угрозы. И даже если они эти угрозы претворят в жизнь, то всё равно мне адвокат не нужен: я сам проведу дело лучше любого адвоката (потому что для них это чужое дело, а для меня – мое, и никто из них не будет ТАК вникать). Может быть ты не знаешь (это описано в моих латышских книгах): я же в 2001–2003 гг. для одной моей читательницы выиграл в трех инстанциях дело, которое все рижские адвокаты считали безна-

дежным (и если браться за это, то требовали 50 латов в час). Так мы выиграли без адвокатов – все бумаги писал я и инструктировал ее, что говорить, а на суде сидел в зрительях. А у противоположной стороны был адвокат (и когда они проиграли в первых двух инстанциях, то поменяли адвоката, – но всё равно не помогло).

Так что, если дело вообще можно выиграть, то я его выиграю сам. (А мое можно выиграть; я, когда балансирую «на лезвии ножа» в своих письменах, то всегда обдумываю ход возможного судебного процесса: какова будет моя линия защиты, – и не пишу такого, что нельзя защищать). Даже Вайра Вике-Фрейберга ничего не добились против меня ☺. (Времени и нервов только жалко, если действительно придется судиться).

Ага, раз ты ищешь фотографию, то согласен на РОТИ-3. Хорошо, тогда я продолжаю ее делать.

В.Э.

no sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 24. marts 16:58
temats Re: Предварительный ответ
piegādātājs rambler.ru

Валдис, добрый день!

Фотографию подобрал, и это означает, что внутренне я готов к совместной работе над РОТИ-3, хотя, честно говоря, пока плохо представляю, как мы эту работу организуем.

Нужен ли нам какой-нибудь стартовый план, чтобы попытаться не выходить за рамки «разных вопросов, более или менее связанных с проблемой интеллекта»? Нужно ли наметить предварительный список целей (по той же причине)? По своему опыту знаю, что работа над проектом всегда вносит коррективы в его планирование, а цели, по мере развития работ, могут видоизменяться. Но с чего-то надо начинать ☺.

Более серьёзным основанием для некоторого предварительного планирования может являться желание сделать Переписку интересной для широкого круга читателей (иначе зачем нам выходить за рамки личной переписки или некоторого относительно закрытого форума). Как ты сам неоднократно замечал в своих работах, форма и порядок изложения информации важны для её осмысления и понимания. Т.е. для читателя хорошо, если вопрос задан «вовремя» ☺. В этом смысле мне нравится твоя книга R-VIEWS «ПУТЬ ИДЕЙ» (хотя знакомым я рекомендую начинать читать с 6-ой страницы¹⁰ ☺).

В РОТИ-1 меня заинтересовала возможность «рассмотреть проект искусственного разума» (стр.12)¹¹, и в связи с этим (или для этого?) посмотреть, «как это может быть устроено у человека».

Как только ты определил, что *«интеллект есть не что иное, как работа системы обработки информации в организме»*, мне захотелось задать вопрос, ответ на который я нашёл на стр.65 (РОТИ-1) *«...достаточно понять такой (конкретный) проект искус-*



Сергей Марьясов

¹⁰ В.Э.: Люди, умеющие работать с литературой, никогда не читают всё подряд, а находят то, что важно в данный момент, будь то в середине или в конце книги и т.д. Я сам так делаю и ожидаю от читателей, что они тоже на это способны.

¹¹ МОИ 2017-09-08: Здесь и далее номера страниц по оригиналу книги; при перепечатке ее в Альманахе текст может быть сдвинут на страницу вперед.

ственного интеллекта, чтобы разрешенными оказались многие «трудные» проблемы в математике, логике и других областях, связанных с (естественным) интеллектом. В разрешении таких проблем при помощи анализа ПРОЕКТА системы программ интеллекта и состоит значение Веданской теории. Конечно, «теоретически» это может понять и сделать «любой»...» Могли бы мы в РОТИ-3 описать сам Проект?

Затем было бы интересно посмотреть (на уровне анализа Проекта) на будущие взаимоотношения человека (точнее человеческого общества) с искусственным интеллектом. И тогда может быть мы подтвердим твоё замечание, сделанное на стр.95 (РОТИ-1): «Именно искусственный интеллект должен был спасти мир от бездонной глупости обезьян. (Теперь уже не надеюсь ни на что: – даже искусственный интеллект не спасет, не успеет спасти мир и цивилизацию).» ☺

Свою фотографию прилагаю.

С.М.

P.S. Замучился с редактированием текста в редакторе электронной почты, поэтому перешёл к подготовке текста в MS Word и затем скопировал его в поле письма, сам MS Word файл посылаю тебе тоже. М.б. в будущем я могу посылать только MS Word файл?

no Valdis Egle <valdis.egle@gmail.com>
kam sevm@rambler.ru
datums 2011. gada 25. marts 16:00
temats Re: Предварительный ответ
piegādātājs gmail.com

Здравствуй, Сергей!

По замыслу проекта «Переписка ПОТИ» она должна была быть просто протоколом реального общения с разными людьми на заданную – в общих чертах – тему, т.е. «хроникой текущих событий», хранилищем писем, записок, файлов и т.д. (где они хранятся в таком виде, чтобы на них можно было легко сослаться по номеру тома, найти FIND-ом и т.п.). Моя общая цель – обсудить эти вопросы с людьми и по ходу дела разобрать те или иные вопросы в связи с Веданской теорией. Более четко я план себе не делаю, но более конкретное направление обсуждению придают именно мои партнеры (корреспонденты). То есть, в случае с тобой – именно ты направляешь книгу РОТИ-3 туда, куда ты хочешь (так предыдущую книгу {РОТИ-2 = МОИ № 42} Дмитрий Манин направил к теоремам Кантора).

Раз ты хочешь более подробно обсудить сам Проект искусственного интеллекта и «будущие взаимоотношения человека (точнее человеческого общества) с искусственным интеллектом», то, значит, эти вопросы и обсудим. Только, ради Бога, не спеша, чтобы я успевал!

(То есть, ты, конечно, можешь высказывать свое мнение в каком угодно темпе и объеме, и я помешу это в книгу, но не требуй от меня слишком быстрых ответов; впрочем, ты можешь написать и такие фрагменты – изложение своих мыслей по данному поводу – которые вообще не требуют моего ответа: ты просто написал фрагмент книги, и всё. – И было бы хорошо, если бы ты написал такие – более или менее развернутые – рецензии на те материалы, которые уже сейчас стоят на сайте РОТИ).

Мне удобнее, когда ты присылаешь текст в Word файле – так легче перенести в книгу, не теряя форматирование (показатели степени, индексы, курсив, bold, подчеркивание, цвет, рисунки и т.д.). Сам я тебе отдельные Word файлы посылать не буду: либо просто e-письмо, либо тогда уже сама книга, доступная и в PDF, и в DOC форматах.

(Книга может быть выпущена – с помещением в Интернет или без его – в предварительных версиях, когда на титульном листе стоит: «Закрыто: Не закрыто»).

В.Э.

Глава 2. Проект искусственного интеллекта

§2. Первая постановка задачи

2011.03.27 16:00 воскресенье

Итак, приступаем к обсуждению «Проекта искусственного интеллекта» – без намерения его на самом деле реализовать в каком-то «хардвере», а с намерением:

- 1) осознать, что для этого в принципе нужно было бы;
- 2) догадаться, «как это может быть устроено у человека»; и
- 3) уяснить, какие последствия из такого «устройства человека» вытекают для тех наук, которые давно существуют, касаются человеческого интеллекта, но учитывают его реальное устройство в недостаточной мере.

Такую задачу я уже ставил перед латвийскими «учеными»¹² в латышских книгах. Я здесь не буду прямо переводить или повторять латышские тексты, но буду держаться той же схемы изложения.

Схема там была такой. Представим, что мы имеем чрезвычайно тонко устроенную куклу по имени Долли, обладающую, с одной стороны, аппаратурой, способной уловить электромагнитные волны диапазона $4-8 \cdot 10^{14}$ герц, колебания воздуха частотой от 20 герц до 20 килогерц, а также ряд других датчиков, фиксирующих наличие различных химических веществ в окружающей куклу среде, температуру, давление на различные части ее тела, состояние различных аппаратов, встроенных в куклу, – а, с другой стороны, у нее есть множество мелких моторчиков, способных при наличии правильных управляющих сигналов осуществить различные движения частей кукольного тела с той же тонкостью и точностью, с какой это делают мышцы человеческого тела.

Наша задача как программистов – спроектировать операционную систему для управления этими куклами так, чтобы, однажды запущенная, она больше не требовала нашего вмешательства извне, а кукла годами и десятилетиями успешно действовала в окружающей ее среде и могла делать всё то, что могут делать люди.

Мы имеем право проектировать не только «софт» (программатуру), но и «хард» (аппаратуру – само устройство компьютера, необходимого для нашей операционной системы): производители построят нам такой компьютер, какой мы закажем.

Вот постановка задачи в самых общих чертах.

Назовем эту проектируемую операционную систему DollOS или, по-русски, Доллос (она – разновидность Витоса, описанного в {PENRO1.345 = МОИ № 14}, только Витосы могут быть и у лягушек, кур и собак, а Доллос – это уже исключительно уровень человека).

Что мы как (опытные) программисты можем сказать о (будущем) Доллосе, когда, почесывая затылки, выходим из кабинета начальника, поставившего нам ТАКУЮ задачу?

Ну, во-первых, очевидно, что в самых глобальных чертах задача сводится к тому, чтобы сигналы от всех имеющихся у Долли датчиков преобразовать в управляющие сигналы ее моторчикам. (Для удобства текста впредь вместо несклоняемой формы имени «Долли» будем употреблять склоняемую форму «Доллия»).

Во-вторых, очевидно, что мы не будем в состоянии предусмотреть все те ситуации, в какие в будущие десятилетия может попасть Доллия, и не сможем заранее обеспечить ее программами на все эти случаи. Следовательно, первое (но фундаментальное!) наше техническое решение состоит в том, что мы не обеспечиваем Доллию программами «на все случаи жизни», а вместо этого встраиваем в нее самопрограммирование.

¹² В том числе перед канадским профессором психологии, действительным членом Латвийской АН Вайрой Вике-Фрейбергой и ее мужем канадским профессором информатики Имантом Фрейбергом, которые в ответ в 2003 году пытались возбудить против меня уголовное дело в Полиции безопасности, и перед «ведущими философами» Латвии: директором Института философии и социологии ЛУ, действительным членом Латвийской АН Майей Куле и сотрудниками этого института, членами-корреспондентами АН Вилнисом Зариньшом и Айваром Табуном (последний из них и есть тот, кто пытался возбудить против меня уголовное дело в Прокуратуре ЛР, но когда я парировал преследование в уголовном порядке, то теперь я стою под угрозой гражданского иска).

То есть, мы не даем ей готовые наборы сигналов для ее моторчиков, а такие наборы она должна составить сама – перед тем, как послать эти сигналы по проводам своим моторчикам.

В связи с этим первым фундаментальным техническим решением Доллоса перед нами возникает задача вообще разработки принципов самопрограммирования (которое в «обычных» операционных системах почти не применяется, а если некоторые их элементы и можно считать относящимися к самопрограммированию, то это – лишь на зачаточном уровне).

Итак, как вообще должны быть устроены самопрограммирующиеся компьютерные системы?

§3. Письмо Сергея Марьясова от 26 апреля 2011 г.

no Sergey <sevm@rambler.ru>
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 26. aprīlis 18:06
temats R-POTI-3
piegādātājs rambler.ru

Валдис, здравствуй!

Прошу прощения за несколько затянувшееся молчание с моей стороны.¹³ Прежде чем что-то писать, хотелось посмотреть то, что уже есть в Векордии (VEcordia) и имеет прямое или косвенное отношение к обсуждаемым в POTI-3 вопросам. Комментарии к книгам Роджера Пенроуза (R-PENRO, R-PENRS) дали мне общее представление о твоих идеях в области теории искусственного интеллекта (ИИ). По одной из ссылок я перешёл на чтение книги L-ARTINT.¹⁴ Однако, недостаточное знание латышского языка несколько замедляет прочтение мною этой книги. Но то, что я уже успел прочитать, наводит меня на мысль, что ответы на мои вопросы частично (или полностью?)¹⁵ находятся в этой книге.

Одновременно я попытался просмотреть, какие материалы доступны в интернете на данную тему. Довольно много статей посвящено разным подходам к теории ИИ, отдельным элементам искусственного интеллекта или косвенным разработкам, но мне не удалось найти ничего похожего на блок-схему, приведённую в L-ARTINT (абзац .605, с. 112). Т.е. понятный и хорошо структурированный взгляд на «природу вещей» отсутствует. И это укрепило меня в мысли, что русскоязычному интернету такая книга нужна. (М.б. такая книга ещё более нужна англоязычному интернету, активность там по затронутой тематике повыше, чем на русскоязычных сайтах, но это вопрос отдельного исследования).

Дополнительную поддержку нашему проекту я нашёл в статье Пол М. Черчленд, Патриция Смит Черчленд. «Искусственный интеллект: может ли машина мыслить?» (судя по аннотации, статья попала в интернет <http://alt-future.narod.ru/Ai/sciam1.html>¹⁶ из журнала «В МИРЕ НАУКИ». (Scientific American. Издание на русском языке). 1990. № 3). Защищая возможность создания ИИ в ближайшем будущем, авторы пишут:

«...Каким образом мозг улавливает смысловое содержание информации, пока не известно, однако ясно, что проблема эта выходит далеко за рамки лингвистики и не ограничивается человеком как видом. Чтобы разработать теорию формирования смыслового содержания, мы должны больше знать о том, как нейроны кодируют и преобразуют сенсорные сигналы, о нейронной основе памяти, об обучении и эмоциях и о связи между этими факторами и моторной системой. Основанная на нейрофизиологии теория понимания смысла может потребовать даже¹⁷ наших интуитивных представлений, которые сейчас кажутся нам такими незыблемыми и которыми так свободно пользуется Сирл¹⁸ в своих рассуждениях. Подобные пересмотры – не редкость в истории науки...».

¹³ В.Э.: Такой темп переписки – как раз то, что надо: оставляет возможность позаниматься и другими делами.

¹⁴ <http://vekordija.narod.ru/L-ARTINT.PDF> , <https://yadi.sk/i/CW7p8j7a3K3JVL>.

¹⁵ В.Э.: Ну – полностью, конечно, нет.

¹⁶ В.Э.: Авторы обеих статей (Сирл и Черчленды) хорошо известны мне и уже фигурировали в Векордии (R-PENRO1, L-IDOM2 = <http://vekordija.narod.ru/L-IDOM-2.PDF>). Обе статьи я скачал и ниже включую в эту книгу, (слегка) прокомментировав.

¹⁷ В.Э.: Видимо, пропущено слово «пересмотра».

¹⁸ С.М.: Дж. Сирл – противник идеи ИИ. Его статья «Разум мозга – компьютерная программа?» опубликована в том же журнале.

Мне кажется, что «проект Искусственного Интеллекта», начатый в L-ARTINT, и работа над которым продолжена в R-POTI является таким пересмотром. А значит нужен и русскоязычному интернету.

Валдис, насколько я понимаю из вступительной части Главы 2 (R-POTI-3), ты не планируешь перевод L-ARTINT.¹⁹ М.б. я бы мог взяться за перевод наиболее важных для нашего проекта глав (с твоим редактированием), чтобы ссылки на них были понятны широкому кругу русскоговорящих читателей²⁰?

Возвращаясь к Главе 2 (R-POTI-3) я бы добавил в список наших намерений как более дальние цели:

4) попытаться определить характер взаимоотношений человеческого общества с искусственным интеллектом;

5) оценить, несёт ли в себе искусственное происхождение интеллекта (ИИ) потенциальное преимущество над биологическим (человеческим) интеллектом в будущем.²¹

Почему-то имя Долли (Доллия в нашем варианте) ассоциируется у меня с головой бедного профессора Доуля. Что это, желание дать возможность интеллекту профессора продолжить существование (жизнь?) в более достойном виде искусственного интеллекта е-человека, или просто совпадение²²?

Возможность перевода личностей отдельных людей в некоторый е-вид (е-людей) обсуждается в ряде статей (например, А. Жаров. «Будущее. Эволюция продолжается». Москва 2006–2007 гг (http://www.planetfuture.info/rus/r_index.htm)²³; Болонкин Александр. Сборник статей. «Бессмертие людей и электронная цивилизация». Библиотека учебной и научной литературы РГИУ. 2007 г. <http://www.twirpx.com/file/385376/>)²⁴. Этот вопрос мы могли бы обсудить в рамках пунктов 4) и 5).

Всего наилучшего,

Сергей Марьясов.

§4. Ответ

2011.04.29 15:42 пятница

Здравствуй, Сергей!

Ты говоришь о том, что «нужно» русскоязычному интернету, что англоязычному...

¹⁹ В.Э.: Там сказано: «Я ЗДЕСЬ не буду прямо переводить или повторять латышские тексты...». То есть, в серии POTI не буду. Но начата также и серия книг VETLA специально для таких переводов. (Только у меня всяких планов так много, что совершенно ясно: ВСЕ планы реализовать не удастся; – что-то будет реализовано, что-то останется нереализованным, а что именно попадет в первое и что во второе из этих множеств, я сейчас предсказать не берусь; это зависит от многих обстоятельств).

²⁰ В.Э.: Я думаю, что это будет лишней тратой энергии. Ссылки из русских текстов на латышские имеются; кто может их отследить – хорошо; кто не может – тоже не велика беда: ведь, например, в пенрозовских списках литературы (R-PENRO5, R-PENRS4) я тоже не могу отследить большинство ссылок (эти книги мне недоступны) – и ничего, всё нормально. Так же надо смотреть и на ссылки из русских книг на латышские.

²¹ В.Э.: Хорошо, добавляем.

²² В.Э.: «Doll» по-английски «кукла». В постановке задачи, о которой говорится выше в §2, дело начиналось так: «Представим, что мы имеем чрезвычайно тонко устроенную куклу...». И еще имя Долли несла овца, которая была (впервые в мире) клонирована в Англии в 1997 году (и стала тогда мировой сенсацией). Собственно от этой (в определенном смысле искусственной!) овцы наша Доллия и получила свое имя (как раз в то время, когда овца Долли гремела по всему миру).

²³ В.Э.: Этот ресурс я «скачал» к себе, и теперь он потенциально готов для включения в Векордию (для меня «читать» и «включить в Векордию» практически одно и то же). Но это целая книга, и для ее чтения-включения потребовались бы недели две чистого времени. Пока что повременю.

²⁴ В.Э.: Этот ресурс находится на сайте, который по-дурацки организован. Может я что-то не заметил, но, насколько я видел, они не дают «гостям» возможность «скачивать» (или даже просто посмотреть) их файлы: надо предварительно зарегистрироваться; зарегистрировавшись же я получаю возможность не только скачивать, но и выставлять на сайт свои и комментировать чужие файлы, т.е. становлюсь членом их корпорации (сайт организован преподавателями ВУЗов). Но я не хочу вступать в их корпорацию (я и так зарегистрирован на десятках сайтов и уже даже точно не помню, на каких именно, поэтому стараюсь в новые не влезать). Так что я счел, что для меня этот ресурс не существует.

У меня нет столь оптимистических настроений. (По-моему, им ничего не нужно, и люди вообще почти никогда не хотят знать правду...). Изложение Веданской теории и других своих идей я хочу оставить после себя в общедоступном виде в Интернете скорее не потому, что это кому-то нужно или в расчете на какой-то широкий резонанс, а потому, что это – мой долг (не перед миром, нет, – перед моей собственной жизнью: зачем я тогда жил, если не оставляю зафиксированным в письменном виде то, что я придумал?).

Но если ты находишься в более оптимистическом настроении, то я не хочу тебя от этого отговаривать.

Если ты не хочешь оставаться в рамках одной лишь этой переписки ПОТИ, то (вместо переводов с латышского) может быть тебе лучше выступить на англоязычных сайтах? (Не прямо сейчас, а через какое-то время, когда уже будет много что сказать).

В надстройке Векордии (файл A-ENTRY.htm, глава «2. Введение») даны общие установки Векордии относительно языков: *«В Векордии имеются тексты на латышском, русском и английском языках. Сам я пишу на латышском и русском, но читаю и цитирую (цитирую без перевода) также на английском».*

Таким образом, твоя переписка (на наши темы) на английских сайтах может быть включена в Векордию (в том числе в серию книг POTI) без всякого перевода. (И со временем у нас могли бы появиться не только русскоязычные, но и англоязычные участники и оппоненты). Мне кажется, что такая твоя роль (в качестве посредника-переводчика на английский) была бы ценнее, чем роль переводчика с латышского.

Если же при таких предполагаемых дискуссиях на английских сайтах потребовалось бы и мне что-то сказать (например, ответить на какой-то адресованный непосредственно мне вопрос), то я это сделал бы по-русски, а ты бы для них перевел.

Однако это, разумеется, по твоему желанию и лишь в отдаленной перспективе. Пока что я включаю в эту книгу две статьи (Сирла и Черчлендов), на которые ты указал своим первым URL. Статьи, правда, довольно старые (21 год только русскому переводу, не учитывая дату английского оригинала), но мне кажется, что в мире мало что изменилось в этом вопросе за это время (а если и изменилось, то в худшую сторону – в сторону усиления всякого мракобесия).

«Китайскую комнату» Сирла я разобрал уже 11 лет назад в ответе профессору Тамбергу {PENRO1 = МОИ № 14}. В теперешней статье мы видим оригинальный (сравнительно поздний) текст самого Сирла (а не пересказ Пенроуза, как в PENRO1). Если смотреть глобально, то точка зрения Сирла базируется на том, что он разграничивает «синтаксис» и «семантику»; синтаксис по его мнению доступен компьютерам, а семантика – нет. Но что такое эта «семантика», Сирл не знает; для него это нечто таинственное, неизвестное. Однако на самом деле «семантика» – просто некоторый (правда, довольно обширный) набор данных, связанный с данным «синтаксисом символов» и без проблем доступный компьютерам. Поэтому вся аргументация Сирла – впустую.

Супруги Черчленды – это философы (университетские преподаватели философии)²⁵. Но они постоянно тесно сотрудничали с нейрофизиологами и публиковались в их журналах (см., например, {IDOM-1}²⁶). Поэтому в данной ниже статье они тоже всё преподносят с точки зрения нейрофизиолога (а не программиста!). В целом их установки, конечно, правильные, но все-таки отсутствие программистской точки зрения не позволяет им дойти до тех результатов, которые имеет Веданская теория (например, объяснить, как из «нейронных сетей» вывести природу математики).

Это были общие оценки, а несколько подробнее – в комментариях к собственно статьям, которые следуют прямо здесь ниже:

Глава 3. Американские философы

§5. Передовица на сайте

<http://alt-future.narod.ru/Ai/sciam1.html>

В МИРЕ НАУКИ. (Scientific American. Издание на русском языке). 1990. № 3

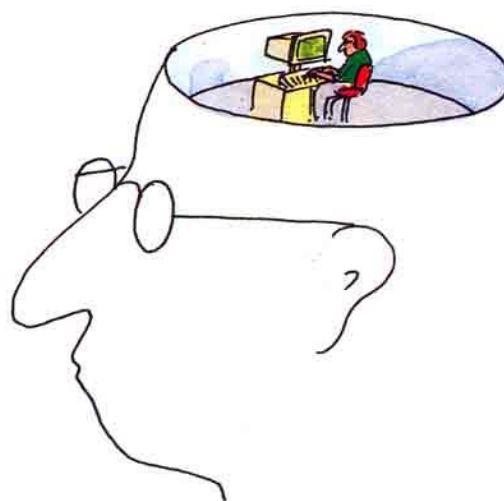
Искусственный интеллект: различные взгляды на проблему

²⁵ The Philosophy Department, University of California San Diego, La Jolla. California 92093, USA.

²⁶ <http://vekordija.narod.ru/L-IDOM-1.PDF>, <https://yadi.sk/i/v6gaGvpX3KEM3N>.

Последние 35 лет попыток создать думающие машины были полны и удач, и разочарований. «Интеллектуальный» уровень современных компьютеров довольно высок, однако для того, чтобы компьютеры могли «разумно» вести себя в реальном мире, их поведенческие способности не должны уступать способностям по крайней мере самых примитивных животных. Некоторые специалисты, работающие в областях, не связанных с искусственным интеллектом, говорят, что компьютеры по своей природе не способны к сознательной умственной деятельности.

В этом номере журнала в статье Дж.Р. Сирла утверждается, что компьютерные программы никогда не смогут достичь разума в привычном для нас понимании. В то же время в другой статье, написанной П.М. Черчлендом и П.С. Черчленд, приводится мнение, что с помощью электронных схем, построенных по образу и подобию мозговых структур, возможно, удастся создать искусственный интеллект. За этим спором по существу скрывается вопрос о том, что такое мышление. Этот вопрос занимал умы людей на протяжении тысячелетий. Практическая работа с компьютерами, которые пока не могут мыслить, породила новый взгляд на этот вопрос и отвергла многие потенциальные ответы на него. Остается найти правильный ответ.²⁷



§6. Джон Сирл. «Разум мозга – компьютерная программа?»

Нет. Программа лишь манипулирует символами, мозг же придает им смысл
ДЖОН СИРЛ

СПОСОБНА ли машина мыслить? Может ли машина иметь сознанные мысли в таком же смысле, в каком имеем их мы? Если под машиной понимать физическую систему, способную выполнять определенные функции (а что еще под ней можно понимать?), тогда люди – это машины особой, биологической разновидности, а люди могут мыслить, и, стало быть, машины, конечно, тоже могут мыслить. Тогда, по всей видимости, можно создавать мыслящие машины из самых разнообразных материалов – скажем, из кремниевых кристаллов или электронных ламп. Если это и окажется невозможным, то пока мы, конечно, этого еще не знаем.

Однако в последние десятилетия вопрос о том, может ли машина мыслить, приобрел совершенно другую интерпретацию. Он был подменен вопросом: способна ли машина мыслить только за счет выполнения заложенной в нее компьютерной программы? Является ли программа основой мышления? Это принципиально иной вопрос, потому что он не затрагивает физических, каузальных (причинных) свойств существующих или возможных физических систем,²⁸ а скорее относится к абстрактным, вычислительным свойствам формализованных компьютерных программ, которые могут быть реализованы в любом материале, лишь бы он был способен выполнять эти программы.

Довольно большое число специалистов по искусственному интеллекту (ИИ) полагают, что на второй вопрос следует ответить положительно; другими словами, они считают, что, составив правильные программы с правильными входами и выходами,²⁹ они действительно создадут разум. Более того, они полагают, что имеют в своем распоряжении научный тест, с помощью которого можно судить об успехе или неудаче такой попытки. Имеется в виду тест Тьюринга,³⁰

²⁷ В.Э.: Здесь воспроизводятся также и 7 рисунков, находившихся на сайте вместе со статьями.

²⁸ В.Э.: Это неверно. Выполнение программы ВСЕГДА есть физический, «каузальный» процесс. А «абстрактная программа» (или, лучше сказать, алгоритм) – это то общее, что имеется у двух (или более) «каузальных процессов» (физическая природа которых может быть совершенно разной).

²⁹ В.Э.: Сущность не во входах и выходах, а в алгоритмах программ: ЧТО они делают.

³⁰ В.Э.: По-моему, тест Тьюринга уже давно никто всерьез не воспринимает. Я не воспринимаю {PENRO1.333 = МОИ № 14}, Черчленды не воспринимают, Сирл не воспринимает, Пенроуз не воспринимает – а кто воспринимает?

изобретенный Аланом М. Тьюрингом, основоположником искусственного интеллекта. Тест Тьюринга в том смысле, как его сейчас понимают, заключается просто в следующем: если компьютер способен демонстрировать поведение, которое эксперт не сможет отличить от поведения человека, обладающего определенными мыслительными способностями (скажем, способностью выполнять операции сложения или понимать китайский язык), то компьютер также обладает этими способностями. Следовательно, цель заключается лишь в том, чтобы создать программы, способные моделировать³¹ человеческое мышление таким образом, чтобы выдерживать тест Тьюринга. Более того, такая программа будет не просто моделью разума; она в буквальном смысле слова сама и будет разумом, в том же смысле, в котором человеческий разум – это разум.³²

Конечно, далеко не каждый специалист по искусственному интеллекту разделяет такую крайнюю точку зрения. Более осторожный подход заключается в том, чтобы рассматривать компьютерные модели как полезное средство для изучения разума, подобно тому, как они применяются при изучении погоды, пищеварения, экономики или механизмов молекулярной биологии.³³ Чтобы провести различие между этими двумя подходами, я назову первый «сильным ИИ», а второй – «слабым ИИ». Важно понять, насколько радикальным является подход сильного ИИ. Сильный ИИ утверждает, что мышление – это не что иное, как манипулирование формализованными символами,³⁴ а именно это и делает компьютер: он оперирует формализованными символами.³⁵ Подобный взгляд часто суммируется примерно следующим высказыванием: «Разум по отношению к мозгу – это то же, что и программа по отношению к аппаратуре компьютера».

СИЛЬНЫЙ ИИ отличается от других теорий разума по крайней мере в двух отношениях: его можно четко сформулировать, но также четко и просто его можно опровергнуть. Характер этого опровержения таков, что каждый человек может попробовать провести его самостоятельно. Вот как это делается. Возьмем, например, какой-нибудь язык, которого вы не понимаете. Для меня таким языком является китайский. Текст, написанный по-китайски, я воспринимаю как набор бессмысленных каракулей. Теперь предположим, что меня поместили в комнату, в которой расставлены корзинки, полные китайских иероглифов. Предположим также, что мне дали учебник на английском языке, в котором приводятся правила сочетания символов китайского языка, причем правила эти можно применять, зная лишь форму символов, понимать значение символов совсем необязательно. Например, правила могут гласить: «Возьмите такой-то иероглиф из корзинки номер один и поместите его рядом с таким-то иероглифом из корзинки номер два».



³¹ В.Э.: Ну вот, Сирл уже ушел от настоящего предмета разговора (который был: «Разум мозга – компьютерная программа?»). Всё! Сирл уже рассуждает не о том, может ли компьютерная программа реализовать разум (разумную деятельность), а о том, как моделировать (имитировать) разумную деятельность, подражать ей. (Его будущий ответ заложен уже в стартовой постановке проблемы).

³² В.Э.: Если они выполняют одинаковую работу по обработке информации, то они одинаково разумны. Если работа разная – значит, не одинаково разумны.

³³ В.Э.: Перечисленные приложения «компьютерных моделей» показывают, что включенная в этот же ряд «модель разума» не будет иметь никакого отношения к реальному ИИ.

³⁴ В.Э.: Нет, не утверждает – это Сирл так утверждает.

³⁵ В.Э.: Нет, компьютер не оперирует формализованными символами – это Сирл думает, что компьютер так работает. На самом деле в компьютере (как и в мозге) происходят определенные физические процессы (преимущественно электрической природы). Эти процессы представляют собой обработку информации.

Представим себе, что находящиеся за дверью комнаты люди, понимающие китайский язык, передают в комнату наборы символов и что в ответ я манипулирую символами согласно правилам и передаю обратно другие наборы символов. В данном случае книга правил есть не что иное, как «компьютерная программа». Люди, написавшие ее, – «программисты», а я играю роль «компьютера». Корзинки, наполненные символами, – это «база данных»³⁶; наборы символов, передаваемых в комнату, это «вопросы», а наборы, выходящие из комнаты, это «ответы».

Предположим далее, что книга правил написана так, что мои «ответы» на «вопросы» не отличаются от ответов человека, свободно владеющего китайским языком.³⁷ Например, люди, находящиеся снаружи, могут передать непонятные мне символы, означающие: «Какой цвет вам больше всего нравится?» В ответ, выполнив предписанные правилами манипуляции,³⁸ я выдам символы, к сожалению, мне также непонятные,³⁹ и означающие, что мой любимый цвет синий, но мне также очень нравится зеленый. Таким образом, я выдержу тест Тьюринга⁴⁰ на понимание китайского языка. Но всё же на самом деле я не понимаю ни слова по-китайски.⁴¹ К тому же я никак не могу научиться этому языку в рассматриваемой системе,⁴² поскольку не существует никакого способа, с помощью которого я мог бы узнать смысл хотя бы одного символа. Подобно компьютеру, я манипулирую символами, но не могу придать им какого бы то ни было смысла.

Сущность этого мысленного эксперимента состоит в следующем: если я не могу понять китайского языка только потому, что выполняю компьютерную программу для понимания китайского, то и никакой другой цифровой компьютер не сможет его понять таким образом.⁴³ Цифровые компьютеры просто манипулируют формальными символами согласно правилам, зафиксированным в программе.⁴⁴

То, что касается китайского языка, можно сказать и о других формах знания. Одного умения манипулировать символами еще недостаточно, чтобы гарантировать знание, восприятие, понимание, мышление и т.д. И поскольку компьютеры как таковые – это устройства, манипулирующие символами, наличия компьютерной программы недостаточно, чтобы можно было говорить о наличии знания.⁴⁵

Этот простой аргумент имеет решающее значение для опровержения концепции сильного ИИ.⁴⁶ Первая предпосылка аргумента просто констатирует формальный характер компьютерной

³⁶ В.Э.: Нет, это не база данных для разума; это всего лишь коробки металлических литер в старой, докомпьютерной типографии; это «алфавит».

³⁷ В.Э.: Красным подчеркнул я. Это центральная, стержневая предпосылка Сирла, и к ней мы еще будем не раз возвращаться. Пока что отметим, что такую «книгу» невозможно написать заранее; она должна «писаться» сама, при этом накапливая где-то у себя подлинную базу данных разума и тем самым овладевая китайским языком. (Назовем эту базу данных «Базой знаний», в отличие от тех корзинок с иероглифами, которые Сирл считал базой данных, и которые на самом деле лишь литеры алфавита).

³⁸ В.Э.: «Правила» смогут предписать определенные манипуляции с иероглифами только в том случае, если обратятся к подлинной базе данных разума (к «Базе знаний») и выполнят действия, эквивалентные действиям мозга, владеющего китайским языком. Если же обращения к «Базе знания» нет, то правила не могут ничего предписывать. То есть, имеет место одно из двух: либо предполагаемых Сирлом в его предпосылке «правил» вообще не существует, либо они представляют собой компьютер с программой и «Базой знаний», достаточной для понимания китайского языка.

³⁹ В.Э.: Но понятные «книге» (которая на самом деле не книга, а компьютер с интеллектом и соответствующей базой данных).

⁴⁰ В.Э.: Нет, не ты выдержишь, а «книга» выдержит – если кто-то в той комнате вообще выдержит.

⁴¹ В.Э.: Зато «книга» понимает – если там вообще кто-то понимает и может выдержать тест.

⁴² В.Э.: Разумеется – интеллектом, владеющим китайским языком-то обладает «книга», а не Сирл.

⁴³ В.Э.: Логическая ошибка! В «китайской комнате» компьютер (обозванный «книгой») понимает китайский язык – если выполняется заданное Сирлом в начале рассуждения (и подчеркнутое красным) условие, что «книга» способна предписывать ему, как нужно манипулировать иероглифами, чтобы за дверью комнаты существо, живущее в комнате, не могли отличить от китайца. Абсурдность этого «доказательства», данного Сирлом, особенно очевидна, если мы представим, что «книга» на самом деле не книга и не компьютер, а настоящий, живой китаец, который сидит и диктует Сирлу (по-английски), какие иероглифы брать. И вот, из того факта, что Сирл не понимает смысла передвигаемых им иероглифов, но, тем не менее правильно выполняет английские указания китайца, Сирл делает вывод, что китайца нет и быть не может!!! Сколь абсурдно это в отношении живого китайца, столь же абсурдно это и в отношении компьютерной программы с подлинным интеллектом.

⁴⁴ В.Э.: Компьютеры, как и мозг, обрабатывают информацию.

⁴⁵ В.Э.: Смотря, какой программы. Если это программа Доллоса, то достаточно.

⁴⁶ В.Э.: Этот аргумент уже опровергнут и не имеет никакого дальнейшего значения.

программы. Программы определяются в терминах манипулирования символами, а сами символы носят чисто формальный или «синтаксический» характер. Между прочим, именно благодаря формальной природе программы, компьютер является таким мощным орудием. Одна и та же программа может выполняться на машинах самой различной природы,⁴⁷ равно как одна и та же аппаратная система способна выполнять самые разнообразные компьютерные программы. Представим это соображение кратко в виде «аксиомы»:

Аксиома 1. Компьютерные программы – это формальные (синтаксические) объекты.⁴⁸

Это положение настолько важно, что его стоит рассмотреть несколько подробнее. Цифровой компьютер обрабатывает информацию, сначала кодируя ее в символических обозначениях, используемых в машине,⁴⁹ а затем манипулируя символами в соответствии с набором строго определенных правил. Эти правила представляют собой программу. Например, в рамках тьюринговской концепции компьютера в роли символов выступали просто 0 и 1, а правила программы предписывали такие операции, как «Записать 0 на ленте, продвинуться на одну ячейку влево и стереть 1». Компьютеры обладают удивительным свойством: любая представимая на естественном языке информация может быть закодирована в такой системе обозначений, и любая задача по обработке информации может быть решена путем применения правил, которые можно запрограммировать.

ВАЖНОЕ значение имеют еще два момента. Во-первых, символы и программы – это чисто абстрактные понятия: они не обладают физическими свойствами, с помощью которых их можно было бы определить и реализовать в какой бы то ни было физической среде.⁵⁰ Нули и единицы, как символы, не имеют физических свойств. Я акцентирую на этом внимание, поскольку иногда возникает соблазн отождествить компьютеры с той или иной конкретной технологией – скажем, с кремниевыми интегральными микросхемами – и считать, что речь идет о физических свойствах кремниевых кристаллов или что синтаксис означает какое-то физическое явление, обладающее, может быть, еще неизвестными каузальными свойствами аналогично реальным физическим явлениям, таким как электромагнитное излучение или атомы водорода, которые обладают физическими, каузальными свойствами. Второй момент заключается в том, что манипуляция символами осуществляется без всякой связи с каким бы то ни было смыслом.⁵¹ Символы в программе могут обозначать всё, что угодно программисту или пользователю. В этом смысле программа обладает синтаксисом, но не обладает семантикой.

⁴⁷ В.Э.: Неточно, и в общем-то «вверх ногами». Не «одна и та же программа может выполняться на машинах самой различной природы», а в устройствах различной природы могут происходить процессы обработки информации, имеющие между собой нечто общее (изоморфизм). Это общее называется программой (а лучше – алгоритмом).

⁴⁸ В.Э.: Аксиома ошибочна.

⁴⁹ В.Э.: Нет, не так. Информация вообще всегда есть некоторое соответствие между объектом *B* и объектом *A*, возникшее в результате определенного физического процесса. Например, если объект *A* – это электромагнитное поле, а объект *B* – сетчатка глаза, то в результате физического процесса (поглощения фотонов) в *B* возникают такие изменения, которые соответствуют (изоморфны) системе *A*, и тогда мы говорим, что *B* «содержит информацию» об *A*, а отдельные элементы системы *B* «кодируют» информацию и являясь «символами» (своих денотатов). От системы *B* информация (опять каким-то физическим процессом) может передаваться к системе *C*, *D*, ..., *Z* и т.д., по пути преобразовываясь, обобщаясь, «обрабатываясь». В компьютер информация попадает тоже (как и в мозг) не «с воздуха», а по цепочке определенных физических процессов. И в мозге (точно так же, как и в компьютере) системы, содержащие информацию, «кодируют» ее «символами» и обрабатывают «программами». С точки зрения информатики никакой принципиальной разницы между человеческим мозгом и промышленным компьютером нет. (Сирл, видимо, просто плохо представляет, как работает и то, и другое).

⁵⁰ В.Э.: Ну вот – до чего договорился Сирл! Для него сначала есть символы, которые «не обладают физическими свойствами», а потом их нельзя «реализовать в какой бы то ни было физической среде»... В реальном мире же всё обстоит «с точностью до наоборот». Сначала есть импульсы «кремниевых интегральных микросхем», а потом мы их начинаем обозначать как 0 или 1. (А можем и не обозначать – компьютер всё равно будет работать; или можем обозначить как-то совсем по-другому, например, FALSE и TRUE или даже вовсе непонятно: NĒ и JĀ). А для Сирла обозначения (символы) первичны, существуют сами по себе, и связь их с физическими объектами отсекается.

⁵¹ В.Э.: Абсолютно неверно! Только потому, что обозначения (символы) соответствуют чему-то вне компьютера, работа программы («манипулирование символами») может быть полезной и пригодной. Конечно, можно и бессмысленно манипулировать (так и делают компьютеры в моменты ожидания, например, в «скрин-сейверах»), но не для этого компьютеры созданы.

Следующая аксиома является простым напоминанием об очевидном факте, что мысли, восприятие, понимание и т.п. имеют смысловое содержание. Благодаря этому содержанию они могут служить отражением объектов и состояний реального мира. Если смысловое содержание связано с языком, то в дополнение к семантике, в нем будет присутствовать и синтаксис, однако лингвистическое понимание требует по крайней мере семантической основы. Если, например, я размышляю о последних президентских выборах, то мне в голову приходят определенные слова, но эти слова лишь потому относятся к выборам, что я придаю им специфическое смысловое значение в соответствии со своим знанием английского языка. В этом отношении они для меня принципиально отличаются от китайских иероглифов. Сформулируем это кратко в виде следующей аксиомы:

Аксиома 2. Человеческий разум оперирует смысловым содержанием (семантикой).⁵²

Теперь добавим еще один момент, который был продемонстрирован экспериментом с китайской комнатой. Располагать только символами как таковыми (т.е. синтаксисом) еще недостаточно для того, чтобы располагать семантикой. Простого манипулирования символами недостаточно, чтобы гарантировать знание их смыслового значения. Кратко представим это в виде аксиомы.

Аксиома 3. Синтаксис сам по себе не составляет семантику и его недостаточно для существования семантики.⁵³

На одном уровне этот принцип справедлив по определению. Конечно, кто-то может определить синтаксис и семантику по-иному. Главное, однако, в том, что существует различие между формальными элементами, не имеющими внутреннего смыслового значения, или содержания, и теми явлениями, у которых такое содержание есть.⁵⁴ Из рассмотренных предпосылок следует:

Закключение 1. Программы не являются сущностью разума, и их наличия недостаточно для наличия разума.⁵⁵

А это по существу означает, что утверждение сильного ИИ ложно.

Очень важно отдавать себе отчет в том, что именно было доказано с помощью этого рассуждения и что нет.

Во-первых, я не пытался доказывать, что «компьютер не может мыслить». Поскольку всё, что поддается моделированию вычислениями, может быть описано как компьютер, и поскольку наш мозг на некоторых уровнях поддается моделированию, то отсюда тривиально следует, что наш мозг – это компьютер, и он, разумеется, способен мыслить. Однако из того факта, что систему можно моделировать посредством мани-



⁵² В.Э.: Верно, но это означает просто то, что программа обращается к той базе данных, которую мы выше назвали «Базой знаний».

⁵³ В.Э.: Это неудачная формулировка, которая не способствует пониманию подлинной природы вещей. На самом деле вопрос заключается просто в том, какими базами данных пользуется программа. Если она не использует «Базу знаний», то Сирл будет говорить, что это «только синтаксис», а если программа использует «Базу знаний», то Сирл будет говорить, что «здесь семантика».

⁵⁴ В.Э.: Конечно, есть разница между той программой, которая к «Базе знаний» не обращается, и той программой, которая к ней обращается.

⁵⁵ В.Э.: Заключение ошибочное. Правильное заключение было бы таким: «Для программ разума необходимо обращение к «Базе знаний»; без этой базы интеллект не построить».

пулирования символами и что она способна мыслить, вовсе не следует, что способность к мышлению эквивалентна способности к манипулированию формальными символами.

Во-вторых, я не пытался доказывать, что только системы биологической природы, подобные нашему мозгу, способны мыслить. В настоящее время это единственные известные нам системы, обладающие такой способностью, однако мы можем встретить во Вселенной и другие способные к осознанным мыслям системы, а может быть, мы даже сумеем искусственно создать мыслящие системы. Я считаю этот вопрос открытым для споров.

В-третьих, утверждение сильного ИИ заключается не в том, что компьютеры с правильными программами могут мыслить, что они могут обладать какими-то неизвестными доселе психологическими свойствами; скорее, оно состоит в том, что компьютеры просто должны мыслить, поскольку их работа – это и есть не что иное, как мышление.

В-четвертых, я попытался опровергнуть сильный ИИ, определенный именно таким образом.⁵⁶ Я пытался доказать, что мышление не сводится к программам, потому что программа лишь манипулирует формальными символами – а, как нам известно, самого по себе манипулирования символами недостаточно, чтобы гарантировать наличие смысла. Это тот принцип, на котором основано рассуждение о китайской комнате.

Я подчеркиваю здесь эти моменты отчасти потому, что П.М. и П.С. Черчленды в своей статье (см. Пол М. Черчленд и Патриция Смит Черчленд «Может ли машина мыслить?»), как мне кажется, не совсем правильно поняли суть моих аргументов. По их мнению, сильный ИИ утверждает, что компьютеры в конечном итоге могут обрести способность к мышлению и что я отрицаю такую возможность, рассуждая лишь на уровне здравого смысла. Однако сильный ИИ утверждает другое, и мои доводы против не имеют ничего общего со здравым смыслом.⁵⁷

Далее я скажу еще кое-что об их возражениях. А пока я должен заметить, что в противоположность тому, что говорят Черчленды, рассуждение с китайской комнатой опровергает любые утверждения сильного ИИ относительно новых параллельных технологий, возникших под влиянием и моделирующих работу нейронных сетей. В отличие от компьютеров традиционной архитектуры фон Неймана, работающих в последовательном пошаговом режиме, эти системы располагают многочисленными вычислительными элементами, работающими параллельно и взаимодействующими друг с другом в соответствии с правилами, основанными на открытиях нейробиологии. Хотя пока достигнуты скромные результаты, модели «параллельной распределенной обработки данных» или «коммутационные машины» подняли некоторые полезные вопросы относительно того, насколько сложными должны быть параллельные системы, подобные нашему мозгу, чтобы при их функционировании порождалось разумное поведение.

Однако параллельный, «подобный мозгу» характер обработки информации не является существенным для чисто вычислительных аспектов процесса. Любая функция, которая может быть вычислена на параллельной машине, будет вычислена и на последовательной.⁵⁸ И действительно, ввиду того, что параллельные машины еще редки, параллельные программы обычно всё еще выполняются на традиционных последовательных машинах. Следовательно, параллельная обработка также не избегает аргумента, основанного на примере с китайской комнатой.

⁵⁶ В.Э.: Ну, здесь тогда нужно разбираться с сирловскими определениями «сильного» и «слабого» ИИ. Там, где он эти понятия вводит (в начале этой своей статьи) – там действительно они определены так, что Веданская теория (т.е. теория подлинного ИИ) не попадает ни в одно из этих понятий и остается вне их (и, значит, сирловская классификация не полна). Но там, где Сирл эти понятия использует практически, – там Веданская теория уже попадает в разряд «сильного ИИ». Так, например, уже в следующем предложении Сирл будет утверждать, что «*мышление не сводится к программам*», т.е. пытаться опровергать не только свой «сильный ИИ» (т.е. сильный ИИ по его определению), но также и Веданскую теорию, потому что мышление все-таки сводится к программам.

⁵⁷ В.Э.: Переводчики что-то так напутали, что получается забавная двусмысленность.

⁵⁸ В.Э.: Вообще это не так. Если речь идет о вычислении значения какой-то математической функции $y = f(x)$ при заданном аргументе x , то, конечно, нет разницы между последовательным и параллельным вычислением. Но если речь идет об операционной системе реального времени, то разница есть – и фундаментальная. И дело не только (и не столько) в скорости вычислений, необходимой для того, чтобы система успела отреагировать в реальном времени. Дело в первую очередь в том, что параллельные процессоры посылают друг другу сигналы, которые для другого процессора являются исходной информацией для обработки («вычислений»). То есть, такой сигнал, поступивший в определенный момент, становится для этого процессора тем самым «аргументом x », от которого надо «вычислять y », а если сигнал поступит в другой момент, то он уже будет «другим x », и значение «функции y » будет другим. Эту (принципиальную и фундаментальную!) асинхронность невозможно отобразить линейным процессом.

Более того, параллельные системы подвержены своей специфической версии первоначального опровергающего рассуждения в случае с китайской комнатой. Вместо китайской комнаты представьте себе китайский спортивный зал, заполненный большим числом людей, понимающих только английский язык. Эти люди будут выполнять те же самые операции, которые выполняются узлами и синапсами в машине коннекционной архитектуры, описанной Черчлендами, но результат будет тем же, что и в примере с одним человеком, который манипулирует символами согласно правилам, записанным в руководстве. Никто в зале не понимает ни слова по-китайски, и не существует способа, следуя которому вся система в целом могла бы узнать о смысловом значении хотя бы одного китайского слова. Тем не менее при правильных инструкциях⁵⁹ эта система способна правильно отвечать на вопросы, сформулированные по-китайски.

У параллельных сетей, как я уже говорил, есть интересные свойства, благодаря которым они могут лучше моделировать мозговые процессы по сравнению с машинами с традиционной последовательной архитектурой. Однако преимущества параллельной архитектуры, существенные для слабого ИИ, не имеют никакого отношения к противопоставлению между аргументом, построенным на примере с китайской комнатой, и утверждением сильного ИИ. Черчленды упускают из виду этот момент, когда говорят, что достаточно большой китайский спортивный зал мог бы обладать более высокими умственными способностями, которые определяются размерами и степенью сложности системы, равно как и мозг в целом более «разумен», чем его отдельные нейроны. Возможно и так, но это не имеет никакого отношения к вычислительному процессу. С точки зрения выполнения вычислений последовательные и параллельные архитектуры совершенно идентичны: любое вычисление, которое может быть произведено в машине с параллельным режимом работы, может быть выполнено машиной с последовательной архитектурой.⁶⁰ Если человек, находящийся в китайской комнате и производящий вычисления, эквивалентен и той и другой системам, тогда, если он не понимает китайского языка исключительно потому, что ничего, кроме вычислений не делает, то и эти системы также не понимают китайского языка. Черчленды правы, когда говорят, что первоначальный довод, основанный на примере с китайской комнатой, был сформулирован исходя из традиционного представления об ИИ, но они заблуждаются, считая что параллельная архитектура делает этот довод неуязвимым. Это справедливо в отношении любой вычислительной системы. Производя только формальные операции с символами (т.е. вычисления) вы не сможете обогатить свой разум семантикой,⁶¹ независимо от того, выполняются эти вычислительные операции последовательно или параллельно; вот почему аргумент китайской комнаты опровергает сильный ИИ в любой его форме.

МНОГИЕ люди, на которых этот аргумент производит определенное впечатление, тем не менее затрудняются провести четкое различие между людьми и компьютерами. Если люди, по крайней мере в тривиальном смысле, являются компьютерами и если люди обладают семантикой, то почему они не могут наделить семантикой и другие компьютеры? Почему мы не можем запрограммировать компьютеры Vax или Cray таким образом, чтобы у них тоже появились мысли и чувства? Или почему какая-нибудь новая компьютерная технология не сможет преодолеть пропасть, разделяющую форму и содержание, или синтаксис и семантику? В чем на самом деле состоит то различие между биологическим мозгом и компьютерной системой, благодаря которому аргумент с китайской комнатой действует применительно к компьютерам, но не действует применительно к мозгу?

Наиболее очевидное различие заключается в том, что процессы, которые определяют нечто как компьютер (а именно, вычислительные процессы), на самом деле совершенно не зависят от какого бы то ни было конкретного типа аппаратной реализации. В принципе можно сделать компьютер из старых жестяных банок из-под пива, соединив их проволокой и обеспечив энергией от ветряных мельниц.

Однако когда мы имеем дело с мозгом, то хотя современная наука в значительной степени еще пребывает в неведении относительно протекающих в мозгу процессов, мы поражаемся чрезвычайной специфичности анатомии и физиологии. Там, где мы достигли некоторого понимания того, как мозговые процессы порождают те или иные психические явления, — например, боль, жажду, зрение, обоняние — нам ясно, что в этих процессах участвуют вполне

⁵⁹ В.Э.: Которые, как мы уяснили, либо не существуют, либо представляют собой программу, обращающуюся к «Базе знаний».

⁶⁰ В.Э.: Так думают люди, слабо знающие компьютеры.

⁶¹ В.Э.: То есть, не сможете обратиться к «Базе знаний»...

определенные нейробиологические механизмы. Чувство жажды, по крайней мере в некоторых случаях, обусловлено срабатыванием нейронов определенных типов в гипоталамусе, которое в свою очередь вызвано действием специфического пептида, ангиотензина II. Причинные связи прослеживаются здесь «снизу вверх» в том смысле, что нейронные процессы низшего уровня обуславливают психические явления на более высоких уровнях. В самом деле, каждое «ментальное» явление, от чувства жажды до мыслей о математических теоремах и воспоминаний о детстве, вызывается срабатыванием определенных нейронов в определенных нейронных структурах.

Однако почему эта специфичность имеет такое важное значение? В конце концов всевозможные срабатывания нейронов можно смоделировать на компьютерах, физические и химические свойства которых совершенно отличны от свойств мозга. Ответ состоит в том, что мозг не просто демонстрирует формальные процедуры или программы (он делает и это тоже), но и вызывает ментальные события благодаря специфическим нейробиологическим процессам. Мозг по сути своей является биологическим органом и именно его особые биохимические свойства позволяют достичь эффекта сознания и других видов ментальных явлений.⁶² Компьютерные модели мозговых процессов обеспечивают отражение лишь формальных аспектов этих процессов. Однако моделирование не следует смешивать с воспроизведением. Вычислительные модели ментальных процессов не ближе к реальности, чем вычислительные модели любого другого природного явления.

Можно представить себе компьютерную модель, отражающую воздействие пептидов на гипоталамус, которая будет точна вплоть до каждого отдельного синапса. Но с таким же успехом мы можем представить себе компьютерное моделирование процесса окисления углеводов в автомобильном двигателе или пищеварительного процесса в желудке. И модель процессов, протекающих в мозге, ничуть не реальнее моделей, описывающих процессы сгорания топлива или пищеварительные процессы. Если не говорить о чудесах, то вы не сможете привести свой автомобиль в движение, моделируя на компьютере окисление бензина, и вы не сможете переварить обед, выполняя программу, которая моделирует пищеварение. Представляется очевидным и тот факт, что и моделирование мышления⁶³ также не произведет нейробиологического эффекта мышления.

Следовательно, все ментальные явления вызываются нейробиологическими процессами мозга. Представим сокращенно этот тезис следующим образом:

Аксиома 4. Мозг порождает разум.

В соответствии с рассуждениями, приведенными выше, я немедленно прихожу к тривиальному следствию.

Закключение 2. Любая другая система, способная порождать разум, должна обладать каузальными свойствами (по крайней мере), эквивалентными соответствующим свойствам мозга.⁶⁴

Это равносильно, например, следующему утверждению: если электрический двигатель способен обеспечивать автомашине такую же высокую скорость, как двигатель внутреннего сгорания, то он должен обладать (по крайней мере) эквивалентной мощностью. В этом заключении ничего не говорится о механизмах. На самом деле, мышление – это биологическое явление: психические состояния и процессы обусловлены процессами мозга. Из этого еще не следует, что только биологическая система может мыслить, но это в то же время означает, что любая система другой природы, основанная на кремниевых кристаллах, жестяных банках и т.п., должна будет обладать каузальными возможностями, эквивалентными соответствующим возможностям мозга. Таким образом, я прихожу к следующему выводу:

Закключение 3. Любой артефакт, порождающий ментальные явления, любой искусственный мозг должен иметь способность воспроизводить специфические каузальные свойства мозга, и наличия этих свойств невозможно добиться только за счет выполнения формальной программы.

Более того, я прихожу к важному выводу, касающемуся человеческого мозга:

⁶² В.Э.: Ну, вот так Сирлу приходится выкручиваться в тех вопросах, которые неизбежно возникают после того, как он отверг идею о том, что разум есть работа программ.

⁶³ В.Э.: Но мы не моделируем мышление, мы его осуществляем в компьютерах.

⁶⁴ В.Э.: Это верно, только Сирл не знает, в чем заключаются эти «каузальные свойства», а мы знаем.

Заключение 4. Тот способ, посредством которого человеческий мозг на самом деле порождает ментальные явления, не может сводиться лишь к выполнению компьютерной программы.⁶⁵

ВПЕРВЫЕ сравнение с китайской комнатой было приведено мною на страницах журнала «*Behavioral and Brain Sciences*» (Науки о поведении и мозге) в 1980 г. Тогда моя статья сопровождалась, в соответствии с принятой в этом журнале практикой, комментариями оппонентов, в данном случае свои соображения высказали 26 оппонентов. Откровенно говоря, мне кажется, что смысл этого сравнения довольно очевиден, но, к моему удивлению, статья и в дальнейшем вызвала целый поток возражений, и что еще более удивительно, этот поток продолжается и по сей день. По-видимому, аргумент китайской комнаты затронул какое-то очень больное место.⁶⁶

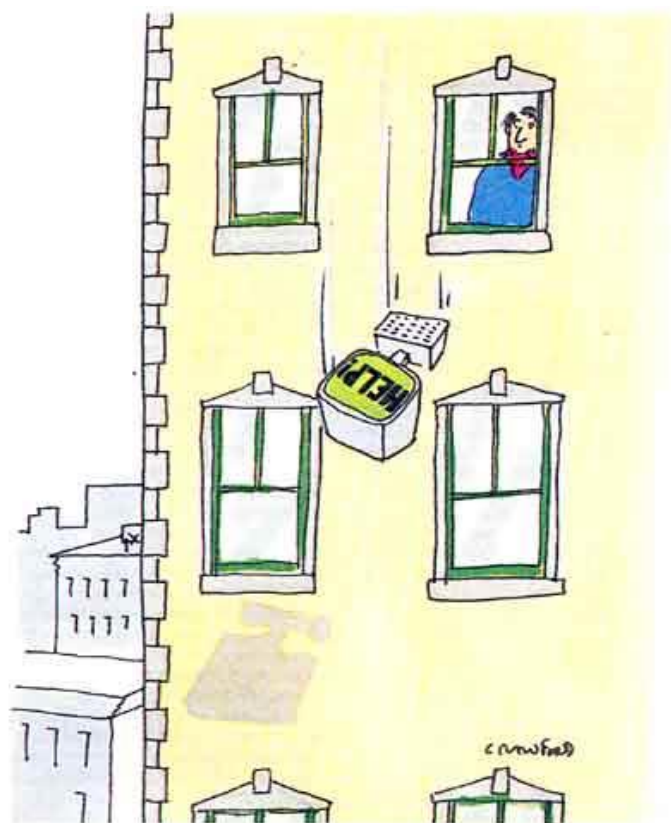
Основной тезис сильного ИИ заключается в том, что любая система (независимо, сделана ли она из пивных банок, кремниевых кристаллов или просто из бумаги) не только способна обладать мыслями и чувствами, но просто должна ими обладать, если только она реализует правильно составленную программу, с правильными входами и выходами.⁶⁷ Очевидно, это абсолютно антибиологическая точка зрения, и естественно было бы ожидать, что специалисты по искусственному интеллекту охотно откажутся от нее.

Многие из них, особенно представители молодого поколения, согласны со мной, но меня поражает, как много сторонников имеет эта точка зрения и как настойчиво они защищают ее. Приведу некоторые наиболее распространенные из высказываемых ими доводов:

а) В китайской комнате вы на самом деле понимаете китайский, хотя и не отдаете себе в этом отчета. В конце концов вы можете понимать что-то и не отдавая себе в этом отчета.⁶⁸

б) Вы не понимаете китайского, но в вас существует подсистема (подсознание), которая понимает. Существуют ведь подсознательные психические состояния, и нет причины считать, что ваше понимание китайского не могло бы быть полностью неосознанным.⁶⁹

в) Вы не понимаете китайского, но комната как целое — понимает. Вы подобны отдельному нейрону в мозгу, и нейрон как таковой не может ничего понимать, он лишь вносит свой вклад в то понимание, которое демонстрирует система в целом; вы сами не понимаете, но вся система понимает.⁷⁰



Кому могло прийти в голову, что компьютерное моделирование процесса мышления и сам мыслительный (ментальный процесс) — это одно и то же?

⁶⁵ В.Э.: Ну, то есть, Сирл не знает, как должна быть устроена такая программа обработки информации, которая осуществит его «каузуальные свойства мозга».

⁶⁶ В.Э.: Просто он неправильный, но на него часто ссылаются, будто он верен.

⁶⁷ В.Э.: С правильным алгоритмом. Ну вот, под это определение «сильного ИИ» Веданская теория попадает.

⁶⁸ В.Э.: Глупости.

⁶⁹ В.Э.: Глупости.

⁷⁰ В.Э.: То стартовое предположение рассуждения Сирла, которое он принял в начале (и которое подчеркнуто выше красным цветом), может выполняться только в том случае, если в комнате есть программа, понимающая китайский язык. Либо это условие не выполняется (и всё рассуждение Сирла пустое), либо «комната как целое понимает китайский язык» (благодаря той «книге», которая на самом деле не книга, а компьютерная программа ИИ или живой китаец).

г) Никакой семантики и не существует: есть только синтаксис. Полагать, что в мозгу есть какое-то загадочное «психическое содержание», «мыслительные процессы» или «семантика», это своего рода донаучная иллюзия. Всё, что на самом деле существует в мозгу, – это некоторое синтаксическое манипулирование символами, которое осуществляется и в компьютерах. И ничего больше.⁷¹

д) В действительности вы не выполняете компьютерную программу – это вам только кажется. Если существует некий сознательный агент, следующий строкам программы, то процесс уже вовсе не является простой реализацией программы.⁷²

е) Компьютеры обладали бы семантикой, а не только синтаксисом, если бы их входы и выходы были поставлены в причинные, каузальные зависимости – по отношению к остальному миру. Допустим, что мы снабдили робота компьютером, подключили телевизионные камеры к его голове, установили трансдюсеры, подводящие телевизионную информацию к компьютеру, и позволили последнему управлять руками и ногами робота. В таком случае система как целое будет обладать семантикой.⁷³

ж) Если программа моделирует поведение человека, говорящего по-китайски, то она понимает китайский язык. Предположим, что нам удалось смоделировать работу мозга китайца на уровне нейронов. Но тогда, конечно, подобная система будет понимать китайский так же хорошо, как и мозг любого китайца.⁷⁴

И так далее.

У всех этих доводов есть одно общее свойство: все они неадекватны рассматриваемой проблеме, потому что не улавливают самой сути рассуждения о китайской комнате. Эта суть заключается в различии между формальным манипулированием символами, осуществляемым компьютером, и смысловым содержанием, биологически порождаемым мозгом, – различии, которое я для краткости выражения (и надеюсь, что это никого не введет в заблуждение) свел к различию между синтаксисом и семантикой. Я не буду повторять своих ответов на все эти возражения, однако проясню, возможно, ситуацию, если скажу, в чем заключаются слабости наиболее распространенного довода моих оппонентов, а именно довода (в), который я назову ответом системы. (Очень часто встречается также и довод (ж), основанный на идее моделирования мозга, но он уже был рассмотрен выше.)

В ОТВЕТЕ системы утверждается, что вы, конечно, не понимаете китайского, но вся система в целом – вы сами, комната, свод правил, корзинки, наполненные символами, – понимает. Когда я впервые услышал это объяснение, я спросил высказавшего это объяснение человека: «Вы что же, считаете, что комната может понимать китайский язык?» Он ответил, да.⁷⁵ Это, конечно, смелое утверждение, однако, помимо того, что оно совершенно неправдоподобно, оно не состоятельно еще и с чисто логической точки зрения. Суть моего исходного аргумента была в том, что простое тасование символов еще не обеспечивает доступа к пониманию смысла этих символов. Но это в той же мере касается комнаты в целом, как и находящегося в ней

⁷¹ В.Э.: В мозге действительно происходит только такое же «манипулирование символами», как в компьютерах, но это вовсе не означает, что «никакой семантики и не существует». Семантика (обращения к «Базе знаний») есть и в мозге, и может быть встроено в компьютер.

⁷² В.Э.: Глупости. Программу Сирл в «китайской комнате», конечно, выполняет, но вопрос заключается в том, откуда эта программа взялась. Она не может быть дана заранее в настоящей, подлинной книге. Она должна возникнуть (в процессе самопрограммирования) уже после того, как в окошко был подан вопрос. Поэтому либо Сирлу вообще нечего будет выполнять (не будет никаких инструкций), либо «книга» должна представлять собой систему, способную генерировать программы (обращаясь к своей «Базе знаний»). Такой системой может быть либо компьютерная программа «сильного ИИ», либо живой китаец.

⁷³ В.Э.: Путанно, неточно и непрофессионально сформулированная идея куклы Доллии. И важно не то, что у нее есть телекамеры и т.д., а важно то, что у нее есть подходящая операционная система – «Доллос».

⁷⁴ В.Э.: Ах, опять эти туманные разговоры о моделировании! (Что здесь означает слово «моделировать»?)

⁷⁵ В.Э.: Надо было не просто ответить «да», а пояснить Сирлу, что его собственное стартовое условие (выше подчеркнутое красным цветом) может выполняться только в том случае, если «комната понимает китайский язык», т.е. если в ней есть соответствующая программа «сильного ИИ» или живой китаец.

человека.⁷⁶ В правоте моих слов можно убедиться, несколько расширив наш мысленный эксперимент. Представим себе, что я заучил наизусть содержимое корзинок и книги правил и что я провожу все вычисления в уме. Допустим даже, что я работаю не в комнате, а у всех на виду. В системе не осталось ничего такого, чего бы не было во мне самом, но поскольку я не понимаю китайского языка, не понимает его и система.

В своей статье мои оппоненты Черчленды используют одну из разновидностей ответа системы, придумав любопытную аналогию. Предположим, кто-то стал утверждать, что свет не может иметь электромагнитную природу, поскольку, когда человек перемещает магнит в темной комнате, мы не наблюдаем видимого светового излучения. Приведя этот пример, Черчленды спрашивают, а не является ли аргумент с китайской комнатой чем-то в том же роде? Не равносильно ли будет сказать, что когда вы манипулируете китайскими иероглифами в семантически темной комнате, в ней не возникает никакого просвета в понимании китайского языка? Но не может ли потом в ходе будущих исследований выясниться – так же, как было доказано, что свет все-таки целиком состоит из электромагнитного излучения, – что семантика целиком и полностью состоит из синтаксиса? Не является ли этот вопрос предметом дальнейшего научного изучения?

Аргументы, построенные на аналогиях, всегда очень уязвимы, поскольку, прежде чем аргумент станет состоятельным, необходимо еще убедиться, что две рассматриваемые ситуации действительно аналогичны. В данном случае, я думаю, что это не так. Объяснение света на основе электромагнитного излучения – это причинное рассуждение от начала и до конца. Это причинное объяснение физики электромагнитных волн. Однако аналогия с формальными символами не состоятельна, поскольку формальные символы не имеют физических причинных свойств. Единственное, что во власти символов как таковых, – это вызвать следующий шаг в программе, которую выполняет работающая машина. И здесь не возникает никакой речи о дальнейших исследованиях, которым еще предстоит раскрыть доселе неизвестные физические причинные свойства нулей и единиц. Последние обладают лишь одним видом свойств – абстрактными вычислительными свойствами, которые уже хорошо изучены.

Черчленды говорят, что у них «напрашивается вопрос», когда я утверждаю, что интерпретированные формальные символы не идентичны смысловому содержанию. Да, я, конечно, не тратил много времени на доказательство, что это так, поскольку я считаю это логической истиной. Как и в случае с любой другой логической истиной, каждый может быстро убедиться, что она справедлива, поскольку, предположив обратное, сразу приходишь к противоречию. Попробуем провести такое доказательство. Предположим, что в китайской комнате имеет место какое-то скрытое понимание китайского языка. Что же может превратить процесс манипулирования синтаксическими элементами в специфично китайское смысловое содержание? Подумав, я в конце концов пришел к выводу, что программисты должны были говорить по-китайски, коль скоро они сумели запрограммировать систему для обработки информации, представленной на китайском языке.⁷⁷

⁷⁶ В.Э.: Неужели Сирлу так никто и не объяснил, как обстоят дела в действительности? (Просто удивительно!) Хотел бы я услышать, что он ответил бы МНЕ.

⁷⁷ В.Э. (Примечание *1): Видно, как Сирл (и Черчленды тоже) ходят кругом до около, но никак не могут увидеть вещи такими, какие они есть на самом деле. Недостаточно знать китайский язык, чтобы написать ту программу, которую Сирл будет (слепо) выполнять в своей «китайской комнате». Здесь должны быть ДВЕ программы: одна (А) – это та, которую будет выполнять Сирл и которая предписывает манипулирование иероглифами. Но эту программу Сирл не может внести с собой при входе в Комнату. Ее еще нет, она не существует. А при входе в Комнату Сирл несет с собой другую программу (В). Ее он выполнять не будет, но он подаст ей на вход те иероглифы, которые сам получит в качестве вопроса, и от нее получит программу А, которую Сирл выполнит. Только так может быть организовано реальное функционирование «китайской комнаты». Так – или никак. А программа В включает в себе «Базу знаний», и она понимает китайский язык. Ну, и здесь тогда видны все нелепости, какие высказывались вокруг этой ситуации. Программа А действительно не может быть получена заранее и взята Сирлом с собой при входе в Комнату. Из этого Сирл делает вывод, что не может существовать программа В! (Здесь хорошо видно также и то, что именно непонимание самопрограммирования – то есть, создания одной программой другой программы – является главным «камнем преткновения» для них всех: для Сирла, Черчлендов, их оппонентов – раз уж никто из них в течение целых десятилетий на это так и не указал и не разобрал ситуацию надлежащим образом. Не понимают они самопрограммирования – а потому и с интеллектом не могут разобраться. Самопрограммирование – это самое главное звено в интеллекте).

Хорошо. Но теперь представим себе, что надоело, сидя в китайской комнате, тасовать китайские (для меня бессмысленные) символы. Предположим, мне пришло в голову интерпретировать эти символы как обозначения ходов в шахматной игре. Какой семантикой теперь обладает система? Обладает ли она китайской семантикой или шахматной, или она обладает одновременно и той и другой? Предположим, что есть еще некая третья личность, наблюдающая за мной в окошко, и она решает, что мое манипулирование символами можно интерпретировать как предсказание курса акций на бирже.⁷⁸ И так далее. Не существует предела количеству семантических интерпретаций, которое можно приписать символам, поскольку, я повторяю, символы носят чисто формальный характер. Они не содержат в себе внутренней семантики.

Можно ли каким-то образом спасти аналогию Черчлендов? Выше я сказал, что формальные символы не имеют каузальных свойств. Но, конечно, программа всегда выполняется той или иной конкретной аппаратурой, и эта аппаратура обладает своими специфическими физическими, каузальными свойствами. Любой реальный компьютер порождает различные физические явления. Мой компьютер, к примеру, выделяет тепло, производит монотонный шум и т.д. Существует ли здесь какое-либо строгое логическое доказательство, что компьютер не может производить аналогичным образом эффект сознания? Нет. В научном смысле об этом и речи быть не может, однако это совсем не то, что призвано опровергать рассуждение о китайской комнате, и не то, на чем будут настаивать сторонники сильного ИИ, поскольку любой производимый таким образом эффект будет достигнут за счет физических свойств реализующей программу среды. Основное утверждение сильного ИИ заключается в том, что физические свойства реализующей среды не имеют никакого значения. Имеют значение лишь программы, а программы – это чисто формальные объекты.⁷⁹

Таким образом аналогия Черчлендов между синтаксисом и электромагнитным излучением наталкивается на дилемму: либо синтаксис следует понимать чисто формально, через его абстрактные математические свойства, либо нет. Если выбрать первую альтернативу, то аналогия становится несостоятельной, поскольку синтаксис, понимаемый таким образом, не имеет физических свойств. Если же, с другой стороны, рассматривать синтаксис в плоскости физических свойств реализующей среды, тогда аналогия действительно состоятельна, но она не имеет отношения к сильному ИИ.

ПОСКОЛЬКУ сделанные мною утверждения довольно очевидны – синтаксис это не то же самое, что семантика; мозговые процессы порождают психические явления – возникает вопрос, а как вообще возникла эта путаница? Кому могло прийти в голову, что компьютерное моделирование ментального процесса полностью ему идентично⁸⁰? В конце концов весь смысл моделей заключается в том, что они улавливают лишь какую-то часть моделируемого явления и не затрагивают остального. Ведь никто не думает, что мы захотим поплавать в бассейне, наполненном шариками для пинг-понга, моделирующими молекулы воды. Можно ли тогда считать, что компьютерная модель мыслительных процессов на самом деле способна мыслить?

Отчасти эти недоразумения объясняются тем, что люди унаследовали некоторые положения бихевиористских психологических теорий прошлого поколения. Под тестом Тьюринга скрывается соблазн считать, что если нечто ведет себя так, как будто оно обладает ментальными процессами, то оно и на самом деле должно ими обладать. Частью ошибочной бихевиористской концепции было также и то, что психология для того, чтобы оставаться

⁷⁸ В.Э.: Это просто бессмыслица (порожденная представлениями Сирла о том, что символы якобы первичны и «носят чисто формальный характер»). Пусть он сначала напишет такую программу, действия которой можно одновременно интерпретировать как разговор на китайском (о том, какой цвет больше нравится), как игру в шахматы и как расчет биржевых курсов, – пусть сначала напишет, прежде чем говорить нам, программистам, такие вещи! Конечно, символу можно «приписать» какое угодно «значение», да только слаженной и осмысленной игры не будет; будет один «шум»; слаженная игра символов же получается только при одной интерпретации – той, которая соответствует действительности.

⁷⁹ В.Э.: Нет, программы не формальные объекты. Не существует программ вне физического носителя. Сущность программы (ее алгоритм) не зависит от природы носителя, но нет программ без носителя. Программа всегда – физический объект. (И только это обстоятельство и дает ей возможность быть выполненной и повлиять на что-то в материальном мире; будь она «формальным объектом», она ни на что в мире не влияла бы). Как абстрактный, формальный объект программу представляют только те, кто весьма поверхностно знают устройство и работу компьютеров.

⁸⁰ В.Э.: В самом деле – кому это могло придти в голову всё время болтать про моделирование, когда очевидно, что речь идет о том, чтобы два разных устройства делали одну и ту же работу! (Как землекоп и экскаватор в примере {PENRS1} [МОИ № 17, стр.30]).

научной дисциплиной, должна ограничиваться изучением внешне наблюдаемого поведения. Парадоксально, но этот остаточный бихевиоризм связан с остаточным дуализмом. Никто не думает, что компьютерная модель пищеварения способна что-то переварить на самом деле, но там, где речь идет о мышлении, люди охотно верят в такие чудеса, потому что забывают о том, что разум – это такое же биологическое явление, как и пищеварение. По их мнению, разум – это нечто формальное и абстрактное, а вовсе не часть полужидкой субстанции, из которой состоит наш головной мозг. Полемическая литература по искусственному интеллекту обычно содержит нападки на то, что авторы называют дуализмом, но при этом они не замечают, как сами демонстрируют ярко выраженный дуализм, поскольку если не принять точку зрения, что разум совершенно не зависит от мозга или какой-либо другой физически специфической системы, то следует считать невозможным создание разума только за счет написания программ.⁸¹

Исторически в странах Запада научные концепции, в которых люди рассматривались как часть обычного физического или биологического мира, часто встречали противодействие со стороны реакции. Идеям Коперника и Галилея противились, потому что они отрицали, что Земля является центром Вселенной. Против Дарвина выступали потому, что он утверждал, что люди произошли от низших животных. Сильный ИИ правильнее всего было бы рассматривать как одно из последних проявлений этой антинаучной традиции, так как он отрицает, что человеческий разум содержит что-то существенно физическое или биологическое. Согласно утверждениям сильного ИИ, разум не зависит от мозга. Он представляет собой компьютерную программу и по существу не связан ни с какой специфической аппаратурой.⁸²

Многие люди, сомневающиеся относительно физической значимости искусственного интеллекта, полагают, что компьютеры, может быть, и смогут понимать китайский язык или думать о числах, но принципиально не способны на проявления чисто человеческих свойств, а именно (и далее следует их излюбленная человеческая специфика): любовь, чувство юмора, тревога за судьбу постиндустриального общества в эпоху современного капитализма и т.д. Но специалисты по ИИ справедливо настаивают, что эти возражения не корректны, что здесь как бы отодвигаются футбольные ворота. Если искусственное моделирование интеллекта окажется успешным, то психологические вопросы уже не имеют сколь-нибудь важного значения. В этом споре обе стороны не замечают различия между моделированием и воспроизведением. Пока речь идет о моделировании, то не стоит никакого труда запрограммировать мой компьютер, чтобы он напечатал: «Я люблю тебя, Сюзи»; «Ха-ха!» или «Я испытываю тревоги постиндустриального общества». Важно отдавать себе отчет в том, что моделирование – это не то же самое, что воспроизведение; и этот факт имеет такое же отношение к размышлениям об арифметике, как и к чувству тревоги. Дело не в том, что компьютер доходит только до центра поля и не доходит до ворот. Компьютер даже не трогается с места. Он просто не играет в эту игру.

§7. Черчленды. «Искусственный интеллект: Может ли машина мыслить?»

Классический искусственный интеллект едва ли будет воплощен в мыслящих машинах; предел человеческой изобретательности в этой области, по-видимому, ограничится созданием систем, имитирующих работу мозга

ПОЛ М. ЧЕРЧЛЕНД, ПАТРИЦИЯ СМИТ ЧЕРЧЛЕНД

В НАУКЕ об искусственном интеллекте (ИИ) происходит революция. Чтобы объяснить ее причины и смысл и представить рассуждения Джона Р. Сирла в перспективе, мы прежде должны обратиться к истории.

В начале 50-х годов традиционный, несколько расплывчатый вопрос о том, может ли машина мыслить, уступил более доступному вопросу: может ли мыслить машина, манипулирующая физическими символами в соответствии с правилами, учитывающими их структуру.

⁸¹ В.Э.: Ну, тут Сирл демонстрирует такую неосведомленность о взглядах своих оппонентов, что и отвечать ему даже не стоит. (Пусть сначала поймет точку зрения противников, а потом уж пишет статьи о них!). Можно ли построить искусственный котел, который будет пищевые продукты в кислой среде расщеплять на химус – то есть, делать ту же работу, что и человеческий желудок? Можно. Можно ли построить искусственный компьютер, который будет обрабатывать ту же информацию и так же, как человеческий мозг, то есть – делать ту же работу? Можно. Какой тут дуализм? Где он?

⁸² В.Э.: Ох, уж лучше ты, Джон, помолчал бы, чем нести такую ахинею! (Возражать нужно против действительных взглядов противников, а не против тобой же вымышленных).

Этот вопрос сформулирован точнее, потому что за предшествовавшие полвека формальная логика и теория вычислений существенно продвинулись вперед. Теоретики стали высоко оценивать возможности абстрактных систем символов, которые претерпевают преобразования в соответствии с определенными правилами. Казалось, что если эти системы удалось бы автоматизировать, то их абстрактная вычислительная мощь проявилась бы в реальной физической системе. Подобные взгляды способствовали рождению вполне определенной программы исследований на достаточно глубокой теоретической основе.

Может ли машина мыслить? Было много причин для того, чтобы ответить да. Исторически одна из первых и наиболее глубоких причин заключалась в двух важных результатах теории вычислений. Первый результат был тезисом Черча, согласно которому каждая эффективно вычислимая функция является рекурсивно вычислимой. Термин «эффективно вычислимая» означает, что существует некая «механическая» процедура, с помощью которой можно за конечное время вычислить результат при заданных входных данных. «Рекурсивно вычислимая» означает, что существует конечное множество операций, которые можно применить к заданному входу, а затем последовательно и многократно применять к вновь получаемым результатам, чтобы вычислить функцию за конечное время. Понятие механической процедуры не формальное, а скорее интуитивное, и потому тезис Черча не имеет формального доказательства. Однако он проникает в самую суть того, чем является вычисление, и множество различных свидетельств сходятся в его подтверждение.

Второй важный результат был получен Аланом М. Тьюрингом, продемонстрировавшим, что любая рекурсивно вычислимая функция может быть вычислена за конечное время с помощью максимально упрощенной машины, манипулирующей символами, которую позднее стали называть универсальной машиной Тьюринга. Эта машина управляется рекурсивно применимыми правилами, чувствительными к идентичности, порядку и расположению элементарных символов, которые играют роль входных данных.

ИЗ ЭТИХ двух результатов вытекает очень важное следствие, а именно: что стандартный цифровой компьютер, снабженный правильной программой, достаточно большой памятью и располагающий достаточным временем, может вычислить любую управляемую правилами функцию с входом и выходом. Другими словами, он может продемонстрировать любую систематическую совокупность ответов на произвольные воздействия со стороны внешней среды.

Конкретизируем это следующим образом: рассмотренные выше результаты означают, что соответственно запрограммированная машина, манипулирующая символами (в дальнейшем будем называть ее МС-машиной), должна удовлетворять тесту Тьюринга на наличие сознательного разума. Тест Тьюринга – это чисто бихевиористский тест, тем не менее его требования очень сильны. (Насколько состоятелен этот тест, мы рассмотрим ниже, там где встретимся со вторым, принципиально отличным «тестом» на наличие сознательного разума.) Согласно первоначальной версии теста Тьюринга, входом для МС-машины должны быть вопросы и фразы на естественном разговорном языке, которые мы набираем на клавиатуре устройства ввода, а выходом являются ответы МС-машины, напечатанные устройством вывода. Считается, что машина выдержала этот тест на присутствие сознательного разума, если ее ответы невозможно отличить от ответов, напечатанных реальным, разумным человеком. Конечно, в настоящее время никому не известна та функция, с помощью которой можно было бы получить выход, не отличающийся от поведения разумного человека. Но результаты Черча и Тьюринга гарантируют нам, что какова бы ни была эта (предположительно эффективная) функция, МС-машина соответствующей конструкции сможет ее вычислить.⁸³

⁸³ В.Э.: Вообще это довольно неудачная постановка проблемы. Есть поговорка, что «правильно поставленный вопрос – это уже половина ответа». Так вот, этой первой половины ответа здесь нет. Наоборот, вопрос уводит мысль исследователя совсем «не в ту степь». Я активно следил за литературой по этому вопросу только одно десятилетие – примерно с 1966 по 1976 год (время, предшествовавшее созданию Веданской теории). Тогда я не пропускал ни одной книги, которая появлялась на советских книжных полках, и могу сказать, что в советской литературе проблема ТАК не ставилась – ставилась более разумно и правильно. Когда читаешь теперь это у Черчлендов, то начинаешь лучше понимать, как и почему Сирл пришел к той ерунде, которую он говорит. Впрочем, может быть, что Сирл и Черчленды тоже не совсем точно отражают даже американскую точку зрения: ведь они все – не инженеры, не конструкторы, не программисты, а философы – преподаватели философии в университетах Калифорнии, только разных: Сирл в *UC Berkeley*, а Черчленды в Сан-Диего.

Это очень важный вывод, особенно если учесть, что тьюринговское описание взаимодействия с машиной при помощи печатающей машинки представляет собой несущественное ограничение. То же заключение остается в силе, даже если МС-машина взаимодействует с миром более сложными способами: с помощью аппарата непосредственного зрения, естественной речи и т.д. В конце концов более сложная рекурсивная функция всё же остается вычислимой по Тьюрингу. Остается лишь одна проблема: найти ту несомненно сложную функцию, которая управляет ответными реакциями человека на воздействия со стороны внешней среды, а затем написать программу (множество рекурсивно применимых правил), с помощью которой МС-машина вычислит эту функцию. Вот эти цели и легли в основу научной программы классического искусственного интеллекта.⁸⁴

Первые результаты были обнадеживающими. МС-машины с остроумно составленными программами продемонстрировали целый ряд действий, которые как будто относятся к проявлениям разума. Они реагировали на сложные команды, решали трудные арифметические, алгебраические и тактические задачи, играли в шашки и шахматы, доказывали теоремы и поддерживали простой диалог. Результаты продолжали улучшаться с появлением более емких запоминающих устройств, более быстродействующих машин, а также с разработкой более мощных и изощренных программ. Классический, или «построенный на программировании», ИИ представлял собой очень живое и успешное научное направление почти со всех точек зрения. Периодически высказывавшееся отрицание того, что МС-машины в конечном итоге будут способны мыслить, казалось проявлением необъективности и неинформированности. Свидетельства в пользу положительного ответа на вопрос, вынесенный в заголовок статьи, казались более чем убедительными.

Конечно, оставались кое-какие неясности. Прежде всего МС-машины не очень-то напоминали человеческий мозг. Однако и здесь у классического ИИ был наготове убедительный ответ. Во-первых, физический материал, из которого сделана МС-машина, по существу не имеет никакого отношения к вычисляемой ею функции. Последняя зафиксирована в программе. Во-вторых, технические подробности функциональной архитектуры машины также не имеют значения, поскольку совершенно различные архитектуры, рассчитанные на работу с совершенно различными программами, могут тем не менее выполнять одну и ту же функцию по входу-выходу.

Поэтому целью ИИ было найти функцию, по входу и выходу характерную для разума, а также создать наиболее эффективную из многих возможных программ для того, чтобы вычислить эту функцию. При этом говорили, что тот специфичный способ, с помощью которого функция вычисляется человеческим мозгом, не имеет значения. Этим и завершается описание сущности классического ИИ и оснований для положительного ответа на вопрос, поставленный в заголовке статьи.

МОЖЕТ ЛИ машина мыслить? Были также кое-какие доводы и в пользу отрицательного ответа. На протяжении 60-х годов заслуживающие внимания отрицательные аргументы встречались относительно редко. Иногда высказывалось возражение, заключающееся в том, что мышление – это не физический процесс и протекает он в нематериальной душе. Однако подобное дуалистическое воззрение не выглядело достаточно убедительным ни с эволюционной, ни с логической точки зрения. Оно не оказало сдерживающего влияния на исследования в области ИИ.

Гораздо большее внимание специалистов по ИИ привлекли соображения иного характера. В 1972 г. Хьюберт Л. Дрейфус опубликовал книгу, в которой резко критиковались парадные демонстрации проявлений разума у систем ИИ. Он указывал на то, что эти системы не адекватно моделировали подлинное мышление, и вскрыл закономерность, присущую всем этим неудачным попыткам. По его мнению, в моделях отсутствовал тот огромный запас неформализованных общих знаний о мире, которым располагает любой человек, а также способность, присущая здравому рассудку, опираться на те или иные составляющие этих знаний в зависимости от требований изменяющейся обстановки.⁸⁵ Дрейфус не отрицал принципиальной возможности создания искусственной физической системы, способной мыслить, но он весьма критически

⁸⁴ В.Э.: Может, в каком-то смысле так и было, но на той легендарной «Дартмутской конференции» 1956 года, где термин «искусственный интеллект» впервые вышел в свет, такие цели не ставились. (Цели были гораздо скромнее).

⁸⁵ В.Э.: Правильно говорил.

отнесся к идее о том, что это может быть достигнуто только за счет манипулирования символами с помощью рекурсивно применяемых правил.

В кругах специалистов по искусственному интеллекту, а также философов рассуждения Дрейфуса были восприняты главным образом как недальновидные и необъективные, базирующиеся на неизбежных упрощениях, присущих этой еще очень молодой области исследований. Возможно, указанные недостатки действительно имели место, но они, конечно, были временными. Настанет время, когда более мощные машины и более качественные программы позволят избавиться от этих недостатков.⁸⁶ Казалось, что время работает на искусственный интеллект. Таким образом, и эти возражения не оказали сколько-нибудь заметного влияния на дальнейшие исследования в области ИИ.

Однако оказалось, что время работало и на Дрейфуса: в конце 70-х – начале 80-х годов увеличение быстродействия и объема памяти компьютеров повышало их «умственные способности» ненамного. Выяснилось, например, что распознавание образов в системах машинного зрения требует неожиданно большого объема вычислений. Для получения практически достоверных результатов нужно было затрачивать всё больше и больше машинного времени, намного превосходя время, требуемое для выполнения тех же задач в биологической системе зрения. Столь медленный процесс моделирования настораживал: ведь в компьютере сигналы распространяются приблизительно в миллион раз быстрее, чем в мозге, а тактовая частота центрального процессорного устройства компьютера примерно во столько же раз выше частоты любых колебаний, обнаруженных в мозге. И всё же на реалистических задачах черепаха легко обгоняет зайца.

Кроме того, для решения реалистических⁸⁷ задач необходимо, чтобы компьютерная программа обладала доступом к чрезвычайно обширной базе данных. Построение такой базы данных уже само по себе представляет довольно сложную проблему, но она усугубляется еще одним обстоятельством: каким образом обеспечить доступ к конкретным, зависящим от контекста фрагментам этой базы данных в реальном масштабе времени. По мере того, как базы данных становились всё более емкими, проблема доступа осложнялась. Исчерпывающий поиск занимал слишком много времени, а эвристические методы не всегда приводили к успеху.⁸⁸ Опасения, подобные тем, что высказывал Дрейфус, начали разделять даже некоторые специалисты, работающие в области искусственного интеллекта.

Приблизительно в это время (1980 г.) Джон Сирл высказал принципиально новую критическую концепцию, ставившую под сомнение само фундаментальное предположение классической программы исследований по ИИ, а именно – идею о том, что правильное манипулирование структурированными символами путем рекурсивного применения правил, учитывающих их структуру, может составлять сущность сознательного разума.

Основной аргумент Сирла базировался на мысленном эксперименте, в котором он демонстрирует два очень важных обстоятельства. Во-первых, он описывает МС-машину, которая (как мы должны понимать) реализует функцию, по входу и выходу способную выдержать тест Тьюринга в виде беседы, протекающей исключительно на китайском языке.⁸⁹ Во-вторых,

⁸⁶ В.Э.: А это была неправильная установка. (Я ее помню). Ожидали как-то, что разум возникнет спонтанно от одной вычислительной мощности. Но он не возникнет сам; а о том, как его на самом деле надо выращивать – не говорили. (Не знали).

⁸⁷ В.Э.: Плохой перевод; видимо, надо по-русски – «реальных задач». (Или имеются в виду задачи реального времени?)

⁸⁸ В.Э.: Природа тут использует два приема, которых обычно нет в наших компьютерных программах. Первый – это древовидные (а не линейные, как в каком-нибудь *FoxPro*!) структуры данных. И второй – это многократное дублирование. Не одна запись с конкретной информацией, как в *FoxPro*, а сразу сто тысяч! (Или миллион).

⁸⁹ В.Э.: Сирл не описывает такую МС-машину. Его начальные установки чрезвычайно расплывчаты. (На этой неопределенности он и играет). В «китайской комнате», чтобы она могла выдержать тест Тьюринга и выполнялось начальное предположение Сирла (подчеркнутое красным), должны присутствовать три элемента разной природы: 1) база данных, названная выше «Базой знаний» и содержащая кадры чьего-то опыта в использовании китайских иероглифов; 2) программа типа Доллоса (обозначенная выше – Примечание *1 – как *B*), способная обращаться к «Базе знаний» и генерировать программы конкретного поведения; 3) программа (обозначенная выше как *A*), сгенерированная программой *B* при использовании «Базы знаний» и предписывающая конкретные манипуляции с иероглифами. Вот, если бы Сирл выделил бы эти три элемента, то можно было сказать, что он описал какую-то машину. Но у него вместо этой картины – сплошное серое пятно. Чей опыт накоплен в «Базе знаний»? Программиста, писавшего

внутренняя структура машины такова, что независимо от того, какое поведение она демонстрирует, у наблюдателя не возникает сомнений⁹⁰ в том, что ни машина в целом, ни любая ее часть не понимают китайского языка. Всё, что она в себе содержит, – это говорящий только по-английски человек, выполняющий записанные в инструкции правила, с помощью которых следует манипулировать символами, входящими и выходящими через окошко для почтовой корреспонденции в двери. Короче говоря, система положительно удовлетворяет тесту Тьюринга,⁹¹ несмотря на то, что не обладает подлинным пониманием китайского языка и реального семантического содержания сообщений (см. статью Дж. Сирла «Разум мозга – компьютерная программа?»).

Отсюда делается общий вывод, что любая система, просто манипулирующая физическими символами согласно чувствительным к структуре правилам, будет в лучшем случае лишь жалкой пародией настоящего сознательного разума, поскольку невозможно породить «реальную семантику», просто крутя ручку «пустого синтаксиса». Здесь следует заметить, что Сирл выдвигает не бихевиористский (не поведенческий) тест на наличие сознания: элементы сознательного разума должны обладать реальным семантическим содержанием.

Возникает соблазн упрекнуть Сирла в том, что его мысленный эксперимент не адекватен, поскольку предлагаемая им система, действующая по типу «кубика-рубика»⁹², будет работать до абсурда медленно. Однако Сирл настаивает, что быстроедействие в данном случае не играет никакой роли. Думающий медленно всё же думает верно. Всё необходимое для воспроизведения мышления, согласно концепции классического ИИ, по его мнению, присутствует в «китайской комнате».

Статья Сирла вызвала оживленные отклики специалистов по ИИ, психологов и философов. Однако в общем и целом она была встречена еще более враждебно, чем книга Дрейфуса. В своей статье, которая одновременно публикуется в этом номере журнала, Сирл приводит ряд критических доводов, высказываемых против его концепции. По нашему мнению, многие из них правомерны, в особенности те, авторы которых жадно «кидаются на приманку», утверждая, что, хотя система, состоящая из комнаты и ее содержимого, работает ужасно медленно, она всё же понимает китайский язык.

Нам нравятся эти ответы, но не потому, что мы считаем, будто китайская комната понимает китайский язык. Мы согласны с Сирлом, что она его не понимает.⁹³ Привлекательность этих аргументов в том, что они отражают отказ воспринять важнейшую третью аксиому в рассуждении Сирла: «Синтаксис сам по себе не составляет семантику и его недостаточно для существования семантики». Возможно, эта аксиома и справедлива,⁹⁴ но Сирл не может с полным основанием утверждать, что ему это точно известно. Более того, предположить, что она

программу *В*? Самой программы *В*? Китайца, стоящего за дверью и проводящего тестирование по Тьюрингу? У Сирла не то что ответов – у него даже вопросов таких нет. Сирл в «китайской комнате» выступает в роли интерпретатора (так называются программы, выполняющие другие программы). Которую программу будет интерпретировать Сирл в «китайской комнате»: программу *А* или программу *В*? Опять, естественно, у Сирла нет ни ответа, ни – еще хуже! – даже вопроса. «Программу» – и всё! Вся постановка проблемы у Сирла чрезвычайно туманна – и в этом тумане делаются ошибочные выводы. Но и Черчленды принимают этот туман всерьез.

⁹⁰ В.Э.: Это у кого не возникает сомнений?!

⁹¹ В.Э.: То, что она удовлетворяет тесту Тьюринга, был не факт, а предположение (начальное допущение) Сирла. Если он (и Черчленды) полагают, что программу *А* можно получить еще за дверью комнаты ДО тестирующих вопросов, то это начальное допущение не выполняется. Если оно выполняется, то в комнате присутствует программа *В* и «База знания» (владеющие китайским языком).

⁹² В.Э.: Кубика Рубика.

⁹³ В.Э.: Если комната не понимает китайский язык, то не выполняется начальное условие Сирла (подчеркнутое в его статье красным цветом), и тогда Сирл не в состоянии выдать в окошко никаких иероглифов, кроме случайных.

⁹⁴ В.Э.: Эта аксиома (как и вся проблема в целом) сформулирована у Сирла очень туманно. Чтобы о ней что-либо сказать, нужно было бы сперва уточнить, что же такое у него «синтаксис», что такое «семантика». Но в общих чертах: семантика есть тогда, когда программа обращается к базе данных, названной у нас «База знаний», в которой накоплен чей-то жизненный опыт. А причислять ли это обращение к базе данных к «манипулированию символами», или не причислять – это дело определений. (Вообще эти понятия «синтаксис», «семантика» лучше выбросить в мусорник и вместо них пользоваться более точными программистскими терминами).

справедлива, – значит напрашиваться на вопрос⁹⁵ о том, состоятельна ли программа исследований классического ИИ, поскольку эта программа базируется на очень интересном предположении, что если нам только удастся привести в движение соответствующим образом структурированный процесс, своеобразный внутренний танец синтаксических элементов, правильно связанный со входами и выходами, то мы можем получить те же состояния и проявления разума, которые присущи человеку.

То, что третья аксиома Сирла действительно напрашивается на этот вопрос, становится очевидно, когда мы непосредственно сопоставляем ее с его же первым выводом: *«Программы не являются сущностью разума и их наличия не достаточно для наличия разума»*. Нетрудно видеть, что его третья аксиома уже несет в себе 90% почти идентичного ей заключения. Вот почему мысленный эксперимент Сирла специально придуман для того, чтобы подкрепить третью аксиому. В этом вся суть китайской комнаты.

Хотя пример с китайской комнатой делает аксиому 3 привлекательной для непосвященного, мы не думаем, что он доказывает справедливость этой аксиомы, и чтобы продемонстрировать несостоятельность этого примера, мы предлагаем в качестве иллюстрации свой параллельный пример. Часто один удачный пример, опровергающий оспариваемое утверждение, значительно лучше проясняет ситуацию, чем целая книга, полная логического жонглирования.

В истории науки было много примеров скепсиса, подобного тому, который мы видим в рассуждениях Сирла. В XVIII в. ирландский епископ Джордж Беркли считал невозможным, чтобы волны сжатия в воздухе сами по себе могли быть сущностью звуковых явлений или фактором, достаточным для их существования. Английский поэт и художник Уильям Блейк и немецкий поэт-естествоиспытатель Иоганн Гете считали невозможным, чтобы маленькие частички материи сами по себе могли быть сущностью или фактором, достаточным для объективного существования света. Даже в нынешнем столетии находились люди, которые не могли себе вообразить, чтобы неодушевленная материя сама по себе, независимо от того, насколько сложна ее организация, могла быть органической сущностью или достаточным условием жизни. Совершенно очевидно то, что люди могут или не могут себе представить, зачастую никак не связано с тем, что на самом деле существует или не существует в действительности. Это справедливо, даже когда речь идет о людях с очень высоким уровнем интеллекта.

Чтобы увидеть, каким образом эти исторические уроки можно применить к рассуждениям Сирла, применим искусственно придуманную параллель к его логике и подкрепим эту параллель мысленным экспериментом.

Аксиома 1. *Электричество и магнетизм – это физические силы.*

Аксиома 2. *Существенное свойство света – это свечение.*

Аксиома 3. *Силы сами по себе появляются сущностью эффекта свечения и не достаточны для его наличия.*

Заключение 1. *Электричество и магнетизм не являются сущностью света и не достаточны для его наличия.*

Предположим, что это рассуждение было опубликовано вскоре после того, как Джеймс К. Максвелл в 1864 г. высказал предположение, что свет и электромагнитные волны идентичны, но до того, как в мире были полностью осознаны систематические параллели между свойствами света и свойствами электромагнитных волн. Приведенное выше логическое рассуждение могло показаться убедительным возражением против смелой гипотезы Максвелла, в особенности если бы оно сопровождалось следующим комментарием в поддержку аксиомы 3.

«Рассмотрим темную комнату, в которой находится человек, держащий в руках постоянный магнит или заряженный предмет. Если человек начнет перемещать магнит вверх–вниз, то, согласно теории Максвелла об искусственном освещении (ИО), от магнита будет исходить распространяющаяся сфера электромагнитных волн и в комнате станет светлее. Но, как хорошо известно всем, кто пробовал играть с магнитами или заряженными шарами, их силы (а если на то пошло, то и любые другие силы), даже когда эти объекты приходят в движение, не создают никакого свечения. Поэтому представляется невозможным, чтобы мы могли добиться реального эффекта свечения просто за счет манипулирования силами!»

КОЛЕБАНИЯ ЭЛЕКТРОМАГНИТНЫХ СИЛ представляют собой свет, хотя магнит, который перемещает человек, не производит никакого свечения. Аналогично манипулирование символами в соответствии с определенными правилами может представлять собой разум, хотя у

⁹⁵ В.Э.: Скорее всего, здесь неточный перевод; видимо, должно было быть: «покушаться на...».

основанной на применении правил системы, находящейся в «китайской комнате» Дж. Сирла, настоящее понимание как будто отсутствует.

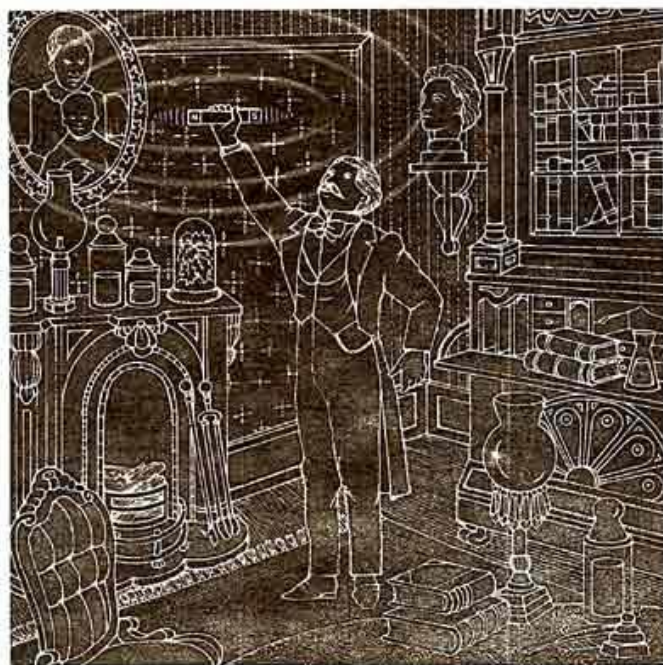
Что же мог ответить Максвелл, если бы ему был брошен этот вызов?

Во-первых, он, возможно, стал бы настаивать на том, что эксперимент со «светящейся комнатой» вводит нас в заблуждение относительно свойств видимого света, потому что частота колебаний магнита крайне мала, меньше, чем нужно, приблизительно в 10^{15} раз. На это может последовать нетерпеливый ответ, что частота здесь не играет никакой роли, что комната с колеблющимся магнитом уже содержит всё необходимое для проявления эффекта свечения в полном соответствии с теорией самого Максвелла.

В свою очередь Максвелл мог бы «проглотить приманку», заявив совершенно обоснованно, что комната уже полна свечения, но природа и сила этого свечения таковы, что человек не способен его видеть. (Из-за низкой частоты, с которой человек двигает магнитом, длина порождаемых электромагнитных волн слишком велика, а интенсивность слишком мала, чтобы глаз человека мог на них среагировать.) Однако, учитывая уровень понимания этих явлений в рассматриваемый период времени (60-е годы прошлого века⁹⁶), такое объяснение, вероятно, вызвало бы смех и издевательские реплики. «Светящаяся комната! Но позвольте, мистер Максвелл, там же совершенно темно!»

Итак, мы видим, что бедному Максвеллу приходится туго. Всё, что он может сделать, это настаивать на следующих трех положениях. Во-первых, аксиома 3 в приведенном выше рассуждении не верна. В самом деле, несмотря на то, что интуитивно она выглядит достаточно правдоподобной, по ее поводу у нас невольно возникает вопрос. Во-вторых, эксперимент со светящейся комнатой не демонстрирует нам ничего интересного относительно физической природы света. И в-третьих, чтобы на самом деле решить проблему света и возможности искусственного свечения, нам необходима программа исследований, которая позволит установить, действительно ли при соответствующих условиях поведение электромагнитных волн полностью идентично поведению света. Такой же ответ должен дать классический искусственный интеллект на рассуждение Сирла. Хотя китайская комната Сирла и может показаться «в семантическом смысле темной», у него нет достаточных оснований настаивать, что совершаемое по определенным правилам манипулирование символами никогда не сможет породить семантических явлений, в особенности если учесть, что люди еще плохо информированы и ограничены лишь пониманием на уровне здравого смысла тех семантических и мыслительных явлений, которые нуждаются в объяснении.⁹⁷ Вместо того, чтобы воспользоваться пониманием этих вещей, Сирл в своих рассуждениях свободно пользуется отсутствием у людей такого понимания.

КИТАЙСКАЯ КОМНАТА	СВЕТАЯЩАЯСЯ КОМНАТА
Аксиома 1. Компьютерные программы — это формальные (синтаксические) объекты.	Аксиома 1. Электричество и магнетизм — это физические силы.
Аксиома 2. Человеческий разум обладает смысловым содержанием (семантикой).	Аксиома 2. Существенное свойство света — это свечение.
Аксиома 3. Синтаксис сам по себе не является сущностью семантики и его не достаточно для семантики.	Аксиома 3. Силы сами по себе не являются сущностью эффекта свечения и не достаточны для его наличия.
Закключение 1. Программы не являются сущностью разума и их наличия не достаточно для существования разума.	Закключение 1. Электричество и магнетизм не являются сущностью света и не достаточны для его наличия.



⁹⁶ В.Э.: Теперь уже позапрошлого: 19-го века.

⁹⁷ В.Э.: Черчленды пытаются защитить «сильный ИИ», но так как у них самих нет представления о том, как ИИ должен работать, то их защита получается совсем слабой. (Она и меня не убедила бы). Не знаю, как со стороны, но то, что Сирлу ответил я, мне представляются гораздо более сильными аргументами.

Высказав свои критические замечания по поводу рассуждений Сирла, вернемся к вопросу о том, имеет ли программа классического ИИ реальный шанс решить проблему сознательного разума и создать мыслящую машину. Мы считаем, что перспективы здесь не блестящие, однако наше мнение основано на причинах, в корне отличающихся от тех аргументов, которыми пользуется Сирл. Мы основываемся на конкретных неудачах исследовательской программы классического ИИ и на ряде уроков, преподанных нам биологическим мозгом на примере нового класса вычислительных моделей, в которых воплощены некоторые свойства его структуры. Мы уже упоминали о неудачах классического ИИ при решении тех задач, которые быстро и эффективно решаются мозгом. Ученые постепенно приходят к общему мнению о том, что эти неудачи объясняются свойствами функциональной архитектуры МС-машин, которая просто непригодна для решения стоящих перед ней сложных задач.

ЧТО НАМ нужно знать, так это каким образом мозг достигает эффекта мышления? Обратное конструирование является широко распространенным приемом в технике. Когда в продажу поступает какое-то новое техническое устройство, конкуренты выясняют, каким образом оно работает, разбирая его на части и пытаясь угадать принцип, на котором оно основано. В случае мозга реализация такого подхода оказывается необычайно трудной, поскольку мозг представляет собой самую сложную вещь на планете. Тем не менее нейрофизиологам удалось раскрыть многие свойства мозга на различных структурных уровнях. Три анатомические особенности принципиально отличают его от архитектуры традиционных электронных компьютеров.

Во-первых, нервная система – это параллельная машина, в том смысле, что сигналы обрабатываются одновременно на миллионах различных путей. Например, сетчатка глаза передает сложный входной сигнал мозгу не порциями по 8, 16 или 32 элемента, как настольный компьютер, а в виде сигнала, состоящего почти из миллиона отдельных элементов, прибывающих одновременно к окончанию зрительного нерва (наружному коленчатому телу), после чего они также одновременно, в один прием, обрабатываются мозгом.

Во-вторых, элементарное «процессорное устройство» мозга, нейрон, отличается относительной простотой. Кроме того, его ответ на входной сигнал – аналоговый, а не цифровой, в том смысле, что частота выходного сигнала изменяется непрерывным образом в зависимости от входных сигналов.

В-третьих, в мозге, кроме аксонов, ведущих от одной группы нейронов к другой, мы часто находим аксоны, ведущие в обратном направлении. Эти возвращающиеся отростки позволяют мозгу модулировать характер обработки сенсорной информации. Еще важнее то обстоятельство, что благодаря их существованию мозг является подлинно динамической системой, у которой непрерывно поддерживаемое поведение отличается как очень высокой сложностью, так и относительной независимостью от периферийных стимулов. Полезную роль в изучении механизмов работы реальных нейронных сетей и вычислительных свойств параллельных архитектур в значительной мере сыграли упрощенные модели сетей. Рассмотрим, например, трехслойную модель, состоящую из нейроноподобных элементов, имеющих аксоноподобные связи с элементами следующего уровня. Входной стимул достигает порога активации данного входного элемента, который посылает сигнал пропорциональной силы по своему «аксону» к многочисленным «синаптическим» окончаниям элементов скрытого слоя. Общий эффект заключается в том, что та или иная конфигурация активирующих сигналов на множестве входных элементов порождает определенную конфигурацию сигналов на множестве скрытых элементов.

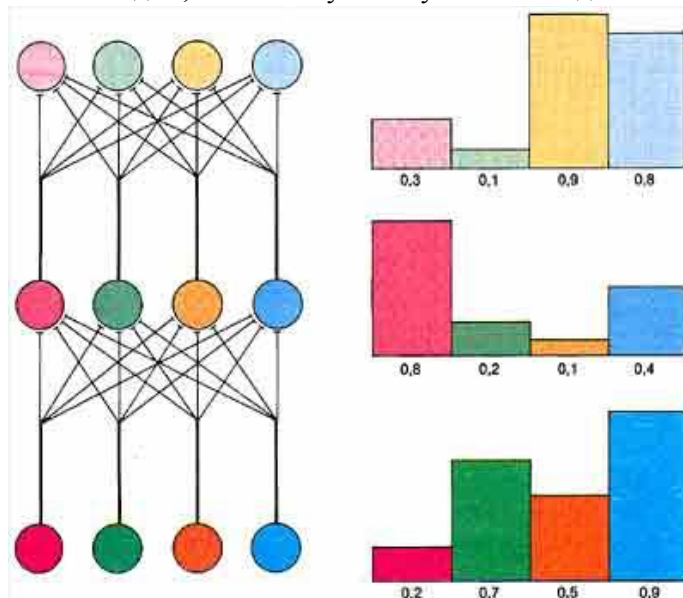
То же самое можно сказать и о выходных элементах. Аналогичным образом конфигурация активирующих сигналов на срезе скрытого слоя приводит к определенной картине активации на срезе выходных элементов. Подводя итог, можно сказать, что рассматриваемая сеть является устройством для преобразования любого большого количества возможных входных векторов (конфигураций активирующих сигналов) в однозначно соответствующий ему выходной вектор. Это устройство предназначено для вычисления специфической функции. То, какую именно функцию оно вычисляет, зависит от глобальной конфигурации синаптической весовой структуры.

НЕЙРОННЫЕ СЕТИ моделируют главное свойство микроструктуры мозга. В этой трехслойной сети входные нейроны (слева внизу) обрабатывают конфигурацию активирующих сигналов (справа внизу) и передают их по взвешенным связям скрытому слою. Элементы скрытого слоя суммируют свои многочисленные входы, образуя новую конфигурацию сигналов. Она передается внешнему слою, выполняющему дальнейшие преобразования. В целом сеть

преобразует любой входной набор сигналов в соответствующий выход, в зависимости от расположения и сравнительной силы связей между нейронами.

Существуют разнообразные процедуры для подбора весов, благодаря которым можно сделать сеть способной вычислить почти любую функцию (т.е. любое преобразование между векторами). Фактически в сети можно реализовать функцию, которую даже нельзя сформулировать, достаточно лишь дать ей набор примеров, показывающих, какие пары входа и выхода мы хотели бы иметь. Этот процесс, называемый «обучением сети», осуществляется путем последовательного подбора весов, присваиваемых связям, который продолжается до тех пор, пока сеть не начнет выполнять желаемые преобразования с входом, чтобы получить нужный выход.

Хотя эта модель сети чрезвычайно упрощает структуру мозга, она всё же иллюстрирует несколько важных аспектов. Во-первых, параллельная архитектура обеспечивает колоссальное преимущество в быстродействии по сравнению с традиционным компьютером, поскольку многочисленные синапсы на каждом уровне выполняют множество мелких вычислительных операций одновременно, вместо того, чтобы действовать в очень трудоемком последовательном режиме. Это преимущество становится всё более значительным по мере того, как возрастает количество нейронов на каждом уровне. Поразительно, но скорость обработки информации, совершенно не зависит ни от числа элементов, участвующих в процессе на каждом уровне, ни от сложности функции, которую они вычисляют. Каждый уровень может иметь четыре элемента или сотню миллионов; конфигурация синаптических весов может вычислять простые одноразрядные суммы или решать дифференциальные уравнения второго порядка. Это не имеет значения. Время вычислений будет абсолютно одним и тем же.



Во-вторых, параллельный характер системы делает ее нечувствительной к мелким ошибкам и придает ей функциональную устойчивость; потеря нескольких связей, даже заметного их количества, оказывает пренебрежимо малое влияние на общий ход преобразования, выполняемого оставшейся частью сети.

В-третьих, параллельная система запоминает большое количество информации в распределенном виде, при этом обеспечивается доступ к любому фрагменту этой информации за время, измеряющееся несколькими миллисекундами. Информация хранится в виде определенных конфигураций весов отдельных синаптических связей, сформировавшихся в процессе предшествовавшего обучения. Нужная информация «высвобождается» по мере того, как входной вектор проходит через эту конфигурацию связей (и преобразуется).

Параллельная обработка данных не является идеальным средством для всех видов вычислений. При решении задач с небольшим входным вектором, но требующих многих миллионов быстро повторяющихся рекурсивных вычислений, мозг оказывается совершенно беспомощным, в то время как классические МС-машины демонстрируют свои самые лучшие возможности. Это очень большой и важный класс вычислений, так что классические машины будут всегда нужны и даже необходимы. Однако существует не менее широкий класс вычислений, для которых архитектура мозга представляет собой наилучшее техническое решение. Это главным образом те вычисления, с которыми обычно сталкиваются живые организмы: распознавание контуров хищника в «шумной» среде; мгновенное вспоминание правильной реакции на его пристальный взгляд, способ бегства при его приближении или защиты при его нападении; проведение различий между съедобными и несъедобными вещами, между половыми партнерами и другими животными; выбор поведения в сложной и постоянно изменяющейся физической или социальной среде; и т.д.

Наконец, очень важно заметить, что описанная параллельная система не манипулирует символами в соответствии со структурными правилами.⁹⁸ Скорее манипулирование символами является лишь одним из многих других «интеллектуальных» навыков, которым сеть может научиться или не научиться. Управляемое правилами манипулирование символами не является основным способом функционирования сети. Рассуждение Сирла направлено против управляемых правилами МС-машин; системы преобразования векторов того типа, который мы описали, таким образом, выпадают из сферы применимости его аргумента с китайской комнатой, даже если бы он был состоятелен, в чем мы имеем другие, независимые причины сомневаться.

Сирлу известно о параллельных процессорах, но, по его мнению, они будут также лишены реального семантического содержания. Чтобы проиллюстрировать их неизбежную неполноценность в этом отношении, он описывает второй мысленный эксперимент, на сей раз с китайским спортивным залом, который наполнен людьми, организованным в параллельную сеть. Дальнейший ход его рассуждений аналогичен рассуждениям в случае с китайской комнатой.

На наш взгляд, этот второй пример не так удачен и убедителен, как первый. Прежде всего то обстоятельство, что ни один элемент в системе не понимает китайского языка, не играет никакой роли, потому что то же самое справедливо и по отношению к нервной системе человека: ни один нейрон моего мозга не понимает английского языка, хотя мозг как целое понимает. Далее Сирл не упоминает о том, что его модель (по одному человеку на каждый нейрон плюс по быстроногую мальчишке на каждую синаптическую связь) потребовала бы по крайней мере 10^{14} человек, так как человеческий мозг содержит 10^{11} нейронов, каждый из которых имеет в среднем 10^3 связей. Таким образом, его система потребует населения 10 000 миров, таких как наша Земля. Очевидно, что спортивный зал далеко не в состоянии вместить более или менее адекватную модель.

С другой стороны, если такую систему всё же удалось бы собрать в соответствующих космических масштабах со всеми точно смоделированными связями, у нас получился бы огромный, медленный, странной конструкции, но всё же функционирующий мозг. В этом случае, конечно, естественно ожидать, что при правильном входе он будет мыслить, а не наоборот, что он на это не способен. Нельзя гарантировать, что работа такой системы будет представлять собой настоящее мышление, поскольку теория векторной обработки может неадекватно отражать работу мозга. Но точно так же у нас нет никакой априорной гарантии, что она не будет мыслить. Сирл еще раз ошибочно отождествляет сегодняшние пределы своего собственного (или читательского) воображения с пределами объективной реальности.

МОЗГ является своеобразным компьютером, хотя большинство его свойств пока остается непознанным. Охарактеризовать мозг как компьютер далеко не просто, и такие попытки не следует считать излишней вольностью. Мозг действительно вычисляет функции, но не такие, как в прикладных задачах, решаемых классическим искусственным интеллектом. Когда мы говорим о машине как о компьютере, мы не имеем в виду последовательный цифровой компьютер, который нужно запрограммировать и которому свойственно четкое разделение на программную часть и аппаратную; мы не имеем также в виду, что этот компьютер манипулирует символами или следует определенным правилам. Мозг – это компьютер принципиально другого вида.

Каким образом мозг улавливает смысловое содержание информации, пока не известно,⁹⁹ однако ясно, что проблема эта выходит далеко за рамки лингвистики и не ограничивается человеком как видом. Маленькая кучка свежей земли означает, как для человека, так и для койота, что где-то поблизости находится суслик; эхо с определенными спектральными характеристиками означает для летучей мыши присутствие мотылька. Чтобы разработать теорию формирования смыслового содержания, мы должны больше знать о том, как нейроны кодируют и преобразуют сенсорные сигналы,¹⁰⁰ о нейронной основе памяти, об обучении и эмоциях и о связи между этими факторами и моторной системой. Основанная на нейрофизиологии теория пони-

⁹⁸ В.Э.: Ну, это опять вопрос терминологии и определения понятий.

⁹⁹ В.Э.: Ну почему не известно? Все предпосылки для понимания этого содержались уже в литературе 1960–1970-х годов, и с середины 1970-х мне уже было ясно, как в принципе это должно работать. Меня на самом деле удивляют эти слова Черчлендов и подобные же в других статьях. (Проблему могут составлять только некоторые технические детали, конкретные алгоритмы – но не сам принцип).

¹⁰⁰ В.Э.: А вот это-то как раз знать необязательно. Нужно идти «сверху», а не «снизу»: по принципу «нисходящего проектирования».

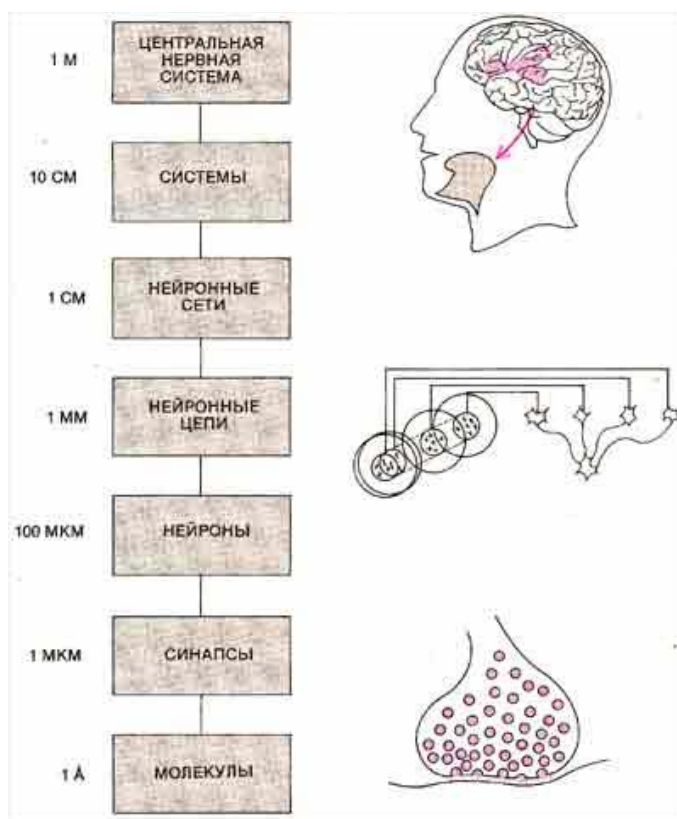
мания смысла может потребовать даже [изменения] наших интуитивных представлений,¹⁰¹ которые сейчас кажутся нам такими неизбежными и которыми так свободно пользуется Сирл в своих рассуждениях. Подобные пересмотры – не редкость в истории науки.

Способна ли наука создать искусственный интеллект, воспользовавшись тем, что известно о нервной системе? Мы не видим на этом пути принципиальных препятствий. Сирл будто бы соглашается, но с оговоркой: *«Любая другая система, способная порождать разум, должна обладать каузальными свойствами (по крайней мере), эквивалентными соответствующим свойствам мозга»*. В завершение статьи мы и рассмотрим это утверждение. Полагаем, что Сирл не утверждает, будто успешная система искусственного интеллекта должна непременно обладать всеми каузальными свойствами мозга, такими как способность чувствовать запах гниющего, способность быть носителем вирусов, способность окрашиваться в желтый цвет под действием пероксидазы хрена обыкновенного и т.д. Требовать полного соответствия будет всё равно, что требовать от искусственного летательного аппарата способности нести яйца.

Вероятно, он имел в виду лишь требование, чтобы искусственный разум обладал всеми каузальными свойствами, относящимися, как он выразился, к сознательному разуму. Однако какими именно? И вот мы опять возвращаемся к спору о том, что относится к сознательному разуму, а что не относится. Здесь как раз самое место поспорить, однако истину в данном случае следует выяснять эмпирическим путем – попробовать и посмотреть, что получится. Поскольку нам так мало известно о том, в чем именно состоит процесс мышления и семантика, то всякая уверенность по поводу того, какие свойства здесь существенны, была бы преждевременной. Сирл несколько раз намекает, что каждый уровень, включая биохимический, должен быть представлен в любой машине, претендующей на искусственный интеллект. Очевидно, это слишком сильное требование. Искусственный мозг может и не пользуясь биохимическими механизмами, достичь того же эффекта.

Эта возможность была продемонстрирована в исследованиях К. Миды в Калифорнийском технологическом институте. Мид и его коллеги воспользовались аналоговыми микроэлектронными устройствами для создания искусственной сетчатки и искусственной улитки уха. (У животных сетчатка и улитка не являются просто преобразователями: в обеих системах происходит сложная параллельная обработка данных.) Эти устройства уже не простые модели в миникомпьютере, над которым посмеивается Сирл; они представляют собой реальные элементы обработки информации, реагирующие в реальное время на реальные сигналы: свет – в случае сетчатки и звук – в случае улитки уха. Схемы устройств основаны на известных анатомических и физиологических свойствах сетчатки кошки и ушной улитки сипухи, и их выход чрезвычайно близок к известным выходам органов, которые они моделируют.

В этих микросхемах не используются никакие нейромедиаторы, следовательно, нейромедиаторы, судя по всему, не являются необходимыми для достижения желаемых результатов. Конечно, мы не можем сказать, что искусственная сетчатка видит что-то, поскольку ее выход не поступает на искусственный таламус или кору мозга и т.д. Возможно ли по программе Миды построить целый искусственный мозг, пока не известно, однако в настоящее время у нас нет



¹⁰¹ В.Э.: Смотря какие у кого эти «интуитивные представления». Но я думаю, что у Черчлендов они такие, что менять не придется. (Нужно просто немножко больше знать о компьютерах и программах).

свидетельств, что отсутствие в системе биохимических механизмов делает этот подход нереалистичным.

НЕРВНАЯ СИСТЕМА охватывает много масштабов организации, от молекул нейромедиаторов (внизу) до всего головного и спинного мозга. На промежуточных уровнях находятся отдельные нейроны и нейронные цепи, подобные тем, что реализуют избирательность восприятия зрительных стимулов (в центре), и системы, состоящие из многих цепей, подобных тем, что обслуживают функции речи (справа вверху). Только путем исследований можно установить, насколько близко искусственная система способна воспроизводить биологические системы, обладающие разумом.

ТАК ЖЕ как и Сирл, мы отвергаем тест Тьюринга как достаточный критерий наличия сознательного разума. На одном уровне основания для этого у нас сходные: мы согласны, что очень важно, каким образом реализуется функция, определенная по входу-выходу; важно, чтобы в машине происходили правильные процессы. На другом уровне мы руководствуемся совершенно иными соображениями. Свою позицию относительно присутствия или отсутствия семантического содержания Сирл основывает на интуитивных представлениях здравого смысла. Наша точка зрения основана на конкретных неудачах классических МС-машин и конкретных достоинствах машин, архитектура которых ближе к устройству мозга. Сопоставление этих различных типов машин показывает, что одни вычислительные стратегии имеют огромное и решающее преимущество над другими в том, что касается типичных задач умственной деятельности. Эти преимущества, установленные эмпирически, не вызывают никаких сомнений. Очевидно, мозг систематически пользуется этими вычислительными преимуществами. Однако он совершенно не обязательно является единственной физической системой, способной ими воспользоваться. Идея создания искусственного интеллекта в небιологической, но существенно параллельной машине остается очень заманчивой и в достаточной мере перспективной.

§8. Продолжение ответа

2011.05.01 17:03 воскресенье

Итак, Сергей, я поместил сюда и прокомментировал обе статьи – Сирла и Черчлендов.

Сирл родился в 1932 году, и 31 июля ему исполняется 79 лет.

Но Черчленды моложе – Пол всего лишь на 4 года старше меня, а Патриция – на 3 (почти что ровесники). В статье 2002 года они указали такие данные о себе: Patricia S. Churchland and Paul M. Churchland are at the Philosophy Department, University of California San Diego, La Jolla. California 92093, USA. Correspondence to P.S.C. e-mail: pschurchland@ucsd.edu doi:10.1038/nrn958.

То есть, они позиционировали себя открытыми для контактов – указали адрес е-почты и т.д.

Может быть, нам попытаться обсудить затронутые в этих двух статьях вопросы с ними (в рамках ПОТИ и в свете того, что я говорил в §4)? Это означало бы, что я тогда написал бы рецензию на эти две статьи (в принципе то же самое, что писал в подстрочных примечаниях, только изложенное более систематически и сдержанно), а ты перевел бы это на английский язык и предложил бы им ответить – мол, для русского интернетовского сайта.

Но это, разумеется, не обязательно.

Спасибо за интересные статьи! Я теперь сам почти ничего по этой тематике не ищущ; то, что попадает ко мне, попадает либо случайно, либо по чьей-то наводке.

Валдис Эгле

1 мая 2011 года

§9. Письма Сергея Марьясова от 2 и 17 мая 2011 г.

no Sergey sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 2. maijs 22:23
temats RE: R-POTI-3
piegādātājs rambler.ru

Валдис, добрый день!

Пока отвечу на твоё письмо кратко: идея с рецензией на эти две статьи – интересная. Можем попробовать подготовить и отправить такое письмо авторам.

В течении нескольких дней напишу более подробный ответ.

С.М.

no Sergey sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 17. maijs 06:13
temats R-POTI-3
piegādātājs rambler.ru

Валдис, добрый день!

Первоначально статью Дж.Р. Сирла «Разум мозга – компьютерная программа?» я прочитал очень поверхностно. Слишком много неопределённостей у автора в тех вопросах, которые нас интересуют (плюс неточности перевода). Особенно если это касается программирования. Так же как и ты, я пытался угадать, какую программу написали программисты для «китайской комнаты», в каком смысле использован термин «моделирование», и в чём суть направления «сильного ИИ», которое, в отличие от «слабого ИИ», и подвергается острой критике в данной статье¹⁰².

Судя по тому, что мне удалось найти в интернете, термин «сильный ИИ» не используется в более поздней литературе. Похоже, что первый мощный натиск, направленный на решение проблемы создания ИИ, в результате затянувшихся «временных трудностей» распался на множество более мелких потоков, направленных на решение частных возможностей человеческого интеллекта. Значительные успехи достигнуты в создании систем распознавания речи, распознавания объектов в пространстве и определения их месторасположения, поиска и принятия решения в условиях недостаточности исходных данных, создания самообучающихся систем и т.д. В результате современные попытки классифицировать направления ИИ скорее отражают наиболее эффективные подходы к решению конкретных задач.

Использование слова «моделирование» в статье (или её переводе?) неудачно. Термин «модель» обычно употребляется в смысле идеализированного или упрощённого описания чего-либо, но может быть применён и для обозначения создания чего-то нового, выполняющего некоторую функцию (или ряд функций) уже существующего объекта¹⁰³. Например: птица – оригинал, самолёт – её модель. Оба объекта летают. В самолёте реализована функция управляемого полёта, а в остальном объекты совсем не похожи друг на друга¹⁰⁴. Аналогично деятельность человеческого мозга характеризуется наличием интеллекта, а некоторая компьютерная система может обладать интеллектом, сравнимым с человеческим в процессах обучения (и самообучения), анализа и решения задач из различных научных областей и т.д. Но, так же как самолёт не является копией птицы (например, ему не нужно высиживать птенцов ☺), так и ИИ не должен быть точной копией человеческого организма¹⁰⁵. В этом смысле термины «моделирование разума» или «моделирование мышления» у Сирла, как создание некоторой концепции ИИ, вполне употребимы (но затрудняют понимание смысла статьи (или «семантики» текста по Сирлу ☺).

¹⁰² После повторного прочтения статьи я вынужден отказаться от своего утверждения в предыдущем письме (прим. 16), что «Дж. Сирл – противник идеи ИИ». Сирл пытается доказать только несостоятельность подхода сторонников «сильного ИИ», суть которого в статье описана в слишком общем виде. Под него подходят и нейронные сети и с некоторыми допущениями и Веданская теория ИИ (твоё примечание 53).

¹⁰³ model: 2) образец; уменьшенная, упрощенная копия; 3) абстрактная схема; концепция. Источник: ABBYY Lingvo Online. <http://lingvo.abbyyonline.com>

¹⁰⁴ Твой пример с землекопом и экскаватором (прим. 77), которые делают одинаковую работу, ближе к обсуждаемой нами теме. Сразу представляю железного робота, сидящего за соседним столом в нашей офисной комнате. Робот смотрит в экран компьютера, где среди открытых рабочих окошек на первом месте открыт *Facebook*, и вместе с коллегами, сидящими за другими столами, активно участвует в обсуждении вопроса, почему нам не выплатили годовую премию ☺.

¹⁰⁵ Наверное, поэтому у меня имя Долли вызвало ассоциацию с головой профессора Доуэля ☺. Довольно ловкие лапы есть и у шимпанзе, но мы собираемся строить интеллект немного посильнее, чем у этих высокоразвитых животных. В результате подсознательно из функций ИИ я исключил функции управления ногами и руками.

Основным элементом (и наиболее плохо описанным) в «китайской комнате» является «программа», которую в виде «книги правил» для Сирла написали «программисты». Но, в общем-то, суть эксперимента от неё мало зависит. Как ты правильно заметил (прим. 34), цель всей китайской комнаты – показать, что с помощью некоторой программы можно создать видимость присутствия в комнате человека, владеющего китайским языком. Сирл так и говорит: *«Предположим далее, что книга правил написана так, что мои «ответы» на «вопросы» не отличаются от ответов человека, свободно владеющего китайским языком»*¹⁰⁶. Всё, дальше спорить не о чем. Такого рода программы, например, игры в шахматы, карты и т.д. уже существуют, и играют иногда лучше людей. В компьютерных играх электронные герои пытаются по ходу игры вести с живым игроком диалог, пусть весьма упрощённый, но, при наличии некоторой фантазии, создающий эффект наличия в игре ещё одного игрока. В интернете я видел упоминания о работах над электронным собеседником, который мог бы поболтать в чате на какие-то общие темы.¹⁰⁷ Но всё это не доказывает, что нельзя создать ИИ.

Теоретически, даже на базе современных компьютеров и известных средств программирования, такую программу (или «книгу правил») написать можно (какой только от неё практический толк?) при условии наличия бесконечно большого времени и неограниченных ресурсов. Но у матушки Природы, когда она создавала Человека Разумного, таких ресурсов не было. Основной строительный материал ограничен рамками планеты Земля (ну, может быть ещё какая мелочёвка из космоса залетит), основной источник энергии ограничен возможностями звезды с названием Солнце. Да и по времени бесконечность не получается, время жизни Солнечной системы порядка 10 миллиардов лет¹⁰⁸. Так что надо торопиться. А то ещё может какой-нибудь метеорит покрупнее на Землю прилетит – и, если не начинать всё сначала, то точно меняй ранее утверждённые планы, снова придумывай, как эту самую Эволюцию дальше продолжать ☺.

Поэтому Природа в ходе эволюции нашла способ создать интеллект немножко по-другому, чем в «китайской комнате». В результате всё необходимое для реализации алгоритмов интеллекта уместается в относительно небольшой голове, и уже через два-три года (после рождения) позволяет вполне резонно отвечать на простые вопросы, а к годам двадцати обеспечивает возможность так ловко говорить по-китайски, и на такие сложные темы, что не всякий китаец и разберёт, о чём идёт речь ☺.

Постулаты Веданской теории тоже не занимают много места (на бумаге ☺):

«1) Человек (вместе с его разумом) был создан в процессе эволюции живой природы путем естественного отбора.

2) Человеческий разум является продуктом деятельности системы обработки информации в организме.»¹⁰⁹

И это обнадёживает: сложное, в конечном итоге, всегда имеет простое объяснение.

К двум основным постулатам мне бы здесь хотелось добавить два тезиса, которые в виде примечаний прозвучали в R-POTI-3 и были более подробно раскрыты в R-POTI-1:

1) Важная роль самопрограммирования в формировании человеческого разума.

«С точки зрения информатики вообще вся «умственная деятельность» человека представляется одним постоянным, непрерывным процессом создания всё новых и новых мозговых программ – постоянно бурлящим котлом самопрограммирования...

...Мышление = самопрограммирование!

¹⁰⁶ Джон Сирл. «Разум мозга – компьютерная программа?»

¹⁰⁷ В.Э.: Насколько я видел эти «чаты», там интеллектуальный уровень такой, что болтающую программу написать не должно составлять никаких трудностей ☺.

¹⁰⁸ Проверая предположительное время жизни Солнечной системы, наткнулся на теорию увеличения массы Земли и Солнца в галактическом цикле (Е.А. Паршаков. Происхождение и развитие Солнечной системы. <http://parshakov.chat.ru/Book1/glava2.html>). Однако из статьи в частности также следует, что время для создания Человека Разумного сокращается периодами между галактическими зимами (в худшем сценарии 120 млн земных лет). И ещё одну статью (<http://www.utro.ru/articles/2005/03/11/416201.shtml>, хотя для более точного понимания надо смотреть первоисточник) о прерывании эволюционной цепочки каждые 62 млн лет.

¹⁰⁹ R-POTI-1, с.11.

Разум = способность создать мозговые программы определенного качества.»
(R-POTI-1, с.22)

2) Материальность программ.

«Во-первых, программа везде и всюду представляет собой материальный объект...

Во-вторых, сущность программы не зависит от материального объекта.» (R-POTI-1, с.23)

«Программа всегда – физический объект. (И только это обстоятельство и дает ей возможность быть выполненной и повлиять на что-то в материальном мире; будь она «формальным объектом», она ни на что в мире не влияла бы).» (R-POTI-3, прим. 76 [здесь 79])

Валдис, в прим. 83 ты коснулся вопроса хранения и поиска информации в памяти человека. Можешь ли ты подробнее раскрыть, как информация хранится в памяти человека, и можно ли с помощью Веданской теории определить, так ли уж верны высказывания в печати о том, что человек использует порядка 5% памяти?

В заключение хочу поместить рисунок, который попался в тексте обсуждения статьи «Google просит пустить беспилотные автомобили на общие дороги»¹¹⁰ и может послужить хорошей иллюстрацией для нашего четвертого вопроса о взаимоотношении человеческого общества с ИИ ☺).

Что касается моего энтузиазма, то он в текущем месяце ограничивался 10–12 часами в неделю, которые я выделял на наш проект. Но желание глубже разобраться в обсуждаемых нами вопросах не пропадает, а наоборот, по мере прочтения книги L-ARTINT, усиливается. И, в перспективе, думаю, шестым вопросом проекта можно будет обозначить привлечение к Веданской теории более широкого внимания.

СМ



§10. Постулат 1 – Естественный отбор

2011.05.17 15:34 вторник

Здравствуй, Сергей!

Спасибо за ценное письмо!

Я рад, что ты правильно понимаешь главные моменты Веданской теории. (После 33 лет тотального непонимания, игнорирования и высмеивания это мне как «бальзам на душу»). Упомянутые тобой четыре момента (два «постулата» и два «тезиса») – это действительно «краеугольные камни» всей той концепции, которая названа Веданской теорией.

Так как ты правильно понял их в первом приближении, то я теперь могу сделать шаг дальше и разобрать эти вещи еще подробнее.

Начнем с Постулата 1: Человек (вместе с его разумом) был создан в процессе эволюции живой природы путем естественного отбора).

Чтобы до конца понять значение и следствия этого постулата, попытаемся его заменить (на время) альтернативными постулатами, например, такими:

Постулат 1б: Человек (вместе с его разумом) был создан Богом.

Постулат 1в: Человек (вместе с его разумом) был создан интеллектом внеземной цивилизации.

Если приняты постулаты (1б) или (1в), то устройство человеческой психики может быть каким угодно – таким, каким это пожелал сделать Творец (Бог или внеземной Конструктор).

¹¹⁰ <http://habrahabr.ru/blogs/robot/119087/>

Захотел он встроить в человека, например, знаменитый фрейдовский «Эдипов комплекс» – и построил (кто может ему запретить?).

Если же у нас в силе Постулат 1 (в противоположность 1б и 1в его можно обозначить и как 1а), то в психику человека не могли встраиваться такие механизмы и аппараты, которые не были ему нужны в борьбе за выживание в процессе естественного отбора. Тогда всё, что в человеке встроено, должно давать ему какие-то преимущества в конкурентной борьбе с другими биологическими видами. И мы, чтобы предположить существование какого-нибудь аппарата, должны указать, какие именно преимущества давались этим аппаратом. А если не можем указать, то и не имеем права предполагать существование этого аппарата.

В конце 19-го и особенно в начале 20-го века Зигмунд Фрейд выдвинул концепцию, согласно которой в основе человеческой психики лежит «комплекс Эдипа» у мужчин и «комплекс Электры» у женщин. Психика делится на «сознание» и «подсознание» (или, другим обозначением «бессознательное»), и в этом подсознании действуют указанные «комплексы»: мужчины бессознательно желают убить отца, чтобы спать с матерью, а женщины – убить мать, чтобы спать с отцом. На основе этой концепции Фрейд и его последователи объясняли различные психические явления, начиная со сновидений и оговорок и кончая психическими заболеваниями, разработали «психоанализ» как метод лечения.

В первой половине 20-го века концепция Фрейда была встречена враждебно классической психиатрией Европы, но в 1930-е годы она стала весьма популярной в США. В СССР в 1920-е годы была сделана попытка внедрить фрейдизм в качестве официальной психологии параллельно марксизму как официальной социологии, но почему-то это не удалось, и потом все остальные годы советской власти Фрейд был почти что под запретом (его книги не издавались, а изданные ранее в царское время и в 1920-е годы находились в спецфондах). С падением советской власти Фрейд захлестнул все книжные полки по психологии в бывшем советском пространстве; одновременно американское влияние подавило сопротивление европейской классической психиатрии, и в настоящее время фрейдизм является фактически «официальной мировой психологией и психиатрией»; практически все современные «школы» психологии и психиатрии так или иначе строятся на базе фрейдизма.

Однако, если мы принимаем Постулат 1а, то фрейдизм очевидно ложен. Невозможно указать, какую пользу в борьбе за существование в процессе естественного отбора приобретают существа, в «подсознание» которых встроено желание убить отца и спать с матерью.

Таким образом, Постулат 1 конфронтирует нас практически со всей современной психологией и психиатрией. Все эти учения – глупости.

Правильные психология и психиатрия должны строиться с учетом Постулата 1 (ну, и далее – вообще Веданской теории).

§11. Постулат 2 – обработка информации

Далее Постулат 2: Человеческий разум является продуктом деятельности системы обработки информации в организме.

Альтернативой этого постулата мог бы быть, например, такой:

Постулат 2б: Человеческая психика несводима к обработке информации.

Это означало бы, что в человеческой психике присутствуют такие компоненты, которые не могут быть воспроизведены компьютерами. Здесь и то (несводимое к материальным процессам) «идеальное» из университетского курса «диалектического материализма», которое ты обсуждал в письме 7 марта после прочтения книги «Путь Идей»; здесь и позиции Джона Сирла и, видимо, также Роджера Пенроуза.

Фактически Постулат 2 декларирует позицию «сильного ИИ» (по терминам Сирла).

А с практической точки зрения этот постулат означает, что для подлинного понимания психики и интеллекта мы должны использовать тот набор понятий и терминов, которым мы пользуемся в информатике (то есть, в компьютерных науках). Основные из этих понятий: информация, программа, алгоритм, база данных (или структура данных, файл), а также – подпрограмма, сопрограмма, функция, блок, интерфейс и т.д.

Постулат 2 согласуется с Постулатом 1: ведь очевидно, какую пользу и выгоду в борьбе за существование приобретут те системы, которые будут обрабатывать информацию об окружающей среде и о своем собственном состоянии.

Но Постулат 2 (и вытекающий из него арсенал средств) находится в полном противоречии с теми постулатами, которые (обычно неявно) приняты во всех науках, касающихся психики и интеллекта.

Математики неспособны увидеть, что основные понятия их науки как-то связаны с Постулатом 2 и всем, что из него вытекает. О сущности числа они предпочитают оставаться на том уровне, до какого мы все доходим в начальной школе. Ни один математик не способен объяснить, что такое число, – во всяком случае объяснить так, чтобы не возникали очевидные нелепости.

Ни в одном учебнике психологии ты не найдешь принципиальных блок-схем программ обработки информации в психике. Психологи даже понятия не имеют, что такие схемы могут быть и что они нужны.

Ну, об этом постулате в общем виде уже много говорилось в книгах РОТІ, поэтому больше не буду; чтобы говорить дальше, надо переходить прямо к программам и схемам.

§12. Самопрограммирование (опять незаконченное)

Самопрограммирование является САМО- программированием только с точки зрения системы в целом (человека, робота – в общем: субъекта). С точки зрения отдельных программ это вовсе не «само»-программирование, а означает, что одна программа (или группа программ) создает другую программу.

В конце §2 я остановился на вопросе: как вообще должны быть устроены самопрограммирующиеся компьютерные системы? Вот, теперь я (в первом приближении) отвечу на этот вопрос.

...

(И опять, занявшись другими книгами, я не дописал этот параграф, как уже пришло новое письмо... Что же, будем фиксировать здесь реальную последовательность событий).

Глава 4. Конспект письма профессору Фрейбергу

§13. Письмо Сергея Марьясова от 24 мая 2011 г. (Хроникер)

no Sergey sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 24. maijs 18:27
temats переписка R-POTI-3
piegādātājs rambler.ru

Валдис, добрый день.

Посылаю тебе некоторые свои наработки по интересующим нас вопросам.

Всего наилучшего,

СМ

Попытавшись охватить (мысленно ☺) некоторый объем информации по поставленным нами вопросам, я пришёл к выводу, что мне для дальнейшей работы нужен русский текст L-ARTINT, по крайней мере, те параграфы, которые связаны с описанием блок-схемы «Модель психической деятельности человека в Веданской теории» (начиная с §28). Недостаточное знание латышского языка не позволяет мне быстро работать с оригинальным текстом книги (обращение к тезисам, поиск детального описания некоторого утверждения или термина).

Поэтому, учитывая, что ты моё предложение о переводе книги не поддержал¹¹¹, я решил ограничиться некоторым конспектом избранных глав L-ARTINT, довольно близким к тексту оригинала, но уже без необходимости сохранять стиль и форму повествования автора.

Я подготовил конспект (с моими комментариями) параграфа «§28. Hronikers (hronikselektors un hronikskriptors.» (прилагаю отдельным файлом). Валдис, включать мой

¹¹¹ СМ. И очень хорошо, что не поддержал. Сейчас я хорошо представляю, сколько усилий и времени, по крайней мере с моей стороны, это потребовало бы.

конспект в R-POTI-3 или не включать – это на твоё усмотрение, но мне бы очень хотелось, чтобы ты в любом случае прокомментировал мои комментарии к §28.

Конечно же, больше всего меня интересует продолжение абзаца .502 (L-ARTINT, §28) об алгоритмах выбора, что писать в Хронос¹¹² и что нет, а если писать, то в каком виде. Размышляя над этими вопросами, я пришёл к выводу, например, что наверное в компьютере Доллии будет класс программ, которые в общем случае не будут привлекать внимание Хроникера (например команды: сжать пальцами книгу, поднять книгу на высоту полки, аккуратно поставить книгу рядом с уже стоящей книгой и т.д.), и эта принадлежность (признак) будет присваиваться по некоторым правилам, и признак этот будет сниматься в исключительных случаях (например, я решил понаблюдать, как мои пальцы сжимают книгу).¹¹³

И даже пофантазировал на тему, почему некоторые, понятные после некоторых усилий идеи, так трудно «пробивают» себе дорогу в моей голове. Например, некоторая информация поступает на вход Хроникера. Хроникер бегло пробегает её и делает вывод: «Н-да, нечто похожее уже поступало к нам, когда ты был в школе. Ты пытался понять, пытался, и ничего не понял (был занят перешептыванием с соседкой по парте ☺). Затем в университете приятель приносил тебе журнал с подробным описанием этой задачи, и ты опять ничего не понял (спешил на студенческую вечеринку ☺). А теперь ты смотришь в эту книгу (дал почитать коллега по работе) и хочешь чего-то там понять... Надо ли тебе это? И вообще ты никогда не любил «философию»... – снимаю с обработки!» На этом месте сознание переключается на чтение следующих страниц...

Валдис, описаны ли где-то эти алгоритмы более подробно в твоих книгах? Планировал ли ты перейти к более подробному описанию алгоритмов работы и структур данных Хроникера¹¹⁴?

С уважением,
СМ

Конспект Сергея Марьясова (довольно подробный) параграфа 28 «Хроникер (X-селектор и X-скриптор)» из книги В. Эгле L-ARTINT

§28. Хроникер (X-селектор и X-скриптор)

.486. Конечно, в этой короткой статье¹¹⁵ мы не сможем рассмотреть операционную систему Доллии полностью. Все вопросы потребуют значительной монографии (и может быть не одной). Мы рассмотрим (на концептуальном уровне), только несколько наиболее фундаментальных (и,

¹¹² В.Э.: В дальнейшем моем тексте я заменяю слово «Хронос» на слово «Хроникер» (использованное в латышском тексте), чтобы не создавать лишнюю путаницу. Слово «Хронос» по-гречески означает «Время». Слово «хроника» означает «описание времён». Слово «хроникер» на нескольких западных языках означает «лицо, которое описывает времена и пишет хронику – летописец». «Хроникерами» также в некоторых операционных системах (я это слово взял из OS/360) назывались специальные программы, ведущие «хронику» деятельности операционной системы. Давай держаться общепринятых установок: хроникер – это тот, кто пишет; хроника – это то, что написано.

¹¹³ В.Э.: Да, примерно так; у меня это ассоциируется с «прожектором», луч которого скользит по информационным полям мозга (человека или Доллии). Я потом нарисую схему и поговорим об этом подробнее.

¹¹⁴ В.Э.: Разумеется, планировал. Но тут есть две проблемы. Первая заключалась в том, что у меня не было читателей. НИ ОДНОГО. (Охотно читали мои политические, исторические и научно-популярные сочинения, но НИКТО никогда не вникал в действительно научные работы). За 33 года существования Веданской теории ты ПЕРВЫЙ, кто серьезно занялся изучением этих идей. Я и так написал (в общем-то громадное) количество текстов – непонятно кому (себе? в воздух?). Очень трудно идти дальше, когда нет ощущения, что хотя бы кем-то понято то, что уже тобою сказано (а вернее: есть убеждение, что оно не понято никем). Вторая проблема заключается в том, что деятельность программ описывать на словах вообще весьма трудно. Представь себе, что мы рассуждаем, например, об операционной системе UNIX. После первых общих идей (которые описать еще несложно) очень скоро наступает момент, когда дальнейший разговор требует уже очень конкретной постановки вопроса о модулях, интерфейсах и т.д. этой операционной системы – вплоть до перехода к операторам языка программирования (например, C). Общими словами уже ничего не скажешь.

¹¹⁵ СМ. Имеется в виду объем одной книги, в данном случае L-ARTINT. В.Э.: ARTINT – это сборник статей. Имелось в виду одно конкретное письмо, отправленное Иммануэлю Фрейбергу, профессору информатики в университете Квебека, мужу Вайры Вике-Фрейберги, тогда президента Латвии. (Таких статей им было послано несколько).

как неспециалист может подумать, – самых «трудных» и «нерешимых») вопросов, в проектировании такого рода систем. Данная статья приобретёт достаточно законченный вид, если мы укажем, как в принципе можно встроить в операционную систему, Долли:

- 1) самопрограммирование;
- 2) осознанное и неосознанное;
- 3) эмоции;
- 4) мечты;¹¹⁶
- 5) абстрактное мышление;
- 6) интуицию;
- 7) творчество.

.487. Прежде всего мы должны признать, что почти вся общепринятая терминология (и принятые в ней такие понятия, как «разум», «мысль», «чувство», «сознание», «подсознание», «Я», и т.д.) – никуда не годится. Ну что специалист по компьютерам, который разрабатывает операционную систему, может сделать с такими понятиями? В компьютере ничего такого нет (может быть если только «память»). Все эти понятия (или, другими словами, эта модель), только блокируют любые дальнейшие действия (и именно последовательное использование этой модели является одной из главных причин того, что то, что мы хотим сделать сейчас, не было сделано раньше).

.488. Поэтому без долгих речей мы отбросим эту модель и поступим следующим образом. Мы не будем философствовать в том духе, что, например, «сознание» – это что-то обладающее какими-то особыми качествами, которые невозможно свести к физиологическим процессам, и так далее. Вместо этого, рассмотрим только объективные факторы: разницу между деятельностью, которую традиционно называют «осознанной», и деятельностью, которую традиционно определяют как «неосознанную».

.489. Например, когда я был в своей комнате, я взял со стола книгу и поставил ее на полку. Я помню, что я сделал это, и поэтому говорю, что сделал это осознанно. Но в другой раз я о чём-то задумался и внезапно обнаружил, что книга, которая 10 минут назад лежала на столе, в настоящее время находится на полке. Я знаю, что в комнате кроме меня никого другого нет и не было, и поэтому мне очевидно, что это я поставил книгу на полку, но только сделал это «неосознанно».

.490. Чем же различаются между собой эти две ситуации? Различие только в том, что, в первом случае, информация о том, как я провел эту операцию (и, возможно, также, как я раньше думал, делать это или нет), была записана в мою память (и, следовательно, я могу сейчас анализировать или «вспомнить» и «подумать»), а во втором случае такой записи в памяти не было.

.491. Как только вопрос сформулирован таким образом, несложно найти соответствующее техническое решение для нашей операционной системы. В операционную систему Доллии мы встраиваем специальный процесс (если мы говорим о программе) или процессор (если мы рассматриваем работу устройства на физическом уровне), который регулярно записывает в память подробную информацию о том, какая программа работала, и что она сделала (назовем этот процесс или процессор «Хроникер»). Теперь если некоторые действия какой-то программы фиксируются в Хроникере, то они становятся «осознанными», а если не фиксируются, то они «неосознанные».

.492. Чисто «априори», как опытные дизайнеры¹¹⁷ программных систем, мы можем сразу же сказать, что Хроникер не сможет регистрировать абсолютно всё, что делают программы Доллии (слишком их много)¹¹⁸. Если мы попытаемся фиксировать абсолютно всё, Хроникер

¹¹⁶ В.Э.: Сновидения. Многим людям сновидения кажутся чем-то совершенно невозможным для компьютеров. А решается дело довольно просто.

¹¹⁷ В.Э.: В наше время их по-русски называли разработчиками. «Дизайнер» у меня ассоциируется только с оформлением внешнего вида.

¹¹⁸ В.Э.: Точный перевод будет таким: «хроникер не сможет фиксировать в своей хронике всё то, что делает (огромное множество) программ Доллии». То есть, объектами «интереса» хроникера являются не столько сами программы, сколько те результаты, которые они производят (и те условия, которые вызвали их запуск) – в общем, те данные, с которыми они работают на входе и выходе. (Отчасти благодаря этому собственно программы остаются для человека «невидимыми», и он обычно не подозревает о их существовании и даже не верит в это, если ему говорят о мозговых программах).

станет гигантской «опухолью», которая «съест» большую часть ресурсов, и ОС Доллии заполнит ее память огромным количеством совершенно ненужной информации.

.493. С другой стороны, мы убеждены, что такие системы, как Хроникер абсолютно необходимы, в противном случае Доллия не сможет ничему научиться из опыта своей предыдущей деятельности, не будет способна проделать последующий анализ или сравнить результаты своей деятельности с запланированными, определить наличие прогресса и так далее.¹¹⁹

.494. Таким образом, мы видим, что для «разума» Доллии разделение на «осознанное» и «неосознанное» является необходимым следствием компромисса между потребностями и возможностями такого типа систем, т.е. по другому такую систему не сделать.

.495. Скажем, Доллия сделала то же самое, что ранее сделал я: взяла книгу со стола и поставила её на полку. Доллия, чтобы сделать это, должна была, как и любой компьютер, иметь программу, которая пошлёт соответствующий приказ моторчикам Доллии. Откуда эта программа появилась в голове Долли, мы рассмотрим немного позже, а сейчас зафиксируем следующее.

.496. Если Долли, как и я, в одном случае осуществила это действие «осознанно» (т.е. действие было зафиксировано в Хронике), в другом «неосознанно» (т.е. без фиксации в Хронике), то что мы можем сказать об этих программах, которые перемещали книги? (Это были две разные программы, потому что в каждом случае они были созданы по-новому,¹²⁰ хотя и из тех же блоков). И как сильно отличаются эти программы?

.497. Зная программирование, мы можем с уверенностью ответить, что обе программы не могут принципиально отличаться, если они выполнили одну и ту же работу и с тем же результатом. Незначительные различия, конечно, могут быть, но они не связаны с режимом выполнения («осознанно» и «неосознанно»), в такой же степени и такого рода различия могут возникнуть и при выполнении двух программ (перемещения книги¹²¹), выполненных в режиме «осознанно».

.498. На основе таких соображений, мы приходим к выводу, что программы, выполняемые в «осознанном» и «неосознанном» режимах не отличаются принципиально, как это считал, например, Зигмунд Фрейд. «Старина Фрейд» (как его назвал Эйнштейн¹²², когда отказался (?)¹²³ поддержать его кандидатуру на Нобелевскую премию) слишком мало знал об операционных системах реального времени, когда строил свои знаменитые конструкции.

.499. Теперь, когда мы представляем суть функционального блока Хроникер, давайте подумаем о принципах, по которым он будет выбирать, что фиксировать, а что нет в своей летописи (имеется в виду память Хроникера¹²⁴). Очевидно, что он не должен фиксировать и записывать все события, особенно события незначительные с точки зрения перспектив процессов самообучения системы (*un par kuriem viņa tātad vispār pati neko nezinās, kamēr kādā skolā viņai par tiem nepastāstīs*).¹²⁵

.500. Понятно, что надо стремиться записать в хронику те процессы, которые «важны» для системы, и по результатам которых можно было бы чему-либо «научиться». Таким образом Хроникеру будет нужен модуль (программа или процессор в составе Хроникера или тесно с ним связанные), который постоянно оценивает порядок «важности» того или иного процесса в голове

¹¹⁹ СМ. В существующих операционных системах всегда есть некоторые процессы, которые что-то пишут в память (набор файлов или специальная база) о том, что происходило в системе. Вопрос только в том, кто, как и почему решил записывать эту информацию. Мой опыт эксплуатации различных операционных систем показал, что пользы от этих системных записей в случае непонятных сбоев в работе системы мало (вот когда причина ошибки известна, тогда другое дело ☺). В данном случае, есть надежда, что мы сумеем достигнуть целей, намеченных данным абзацем.

¹²⁰ В.Э.: заново.

¹²¹ СМ. Добавлено мною.

¹²² Пайс А. «Научная деятельность и жизнь Альберта Эйнштейна». Наука, Москва, 1989., с. 485.

¹²³ В.Э.: Да, отказался – Эйнштейн не поддержал кандидатуру Фрейда на Нобелевскую премию.

¹²⁴ СМ. Добавлено мною.

¹²⁵ СМ. Мне не очень ясен переход от основной мысли абзаца. Имеется в виду обучение (получение новой информации)? В.Э.: Точный перевод абзаца: .499. Теперь, когда нам ясна сущность такого функционального блока Доллии как хроникер, подумаем еще, по какому принципу он будет выбирать, что фиксировать и что не фиксировать в своей хронике. Ясно, что он не может фиксировать и не фиксирует всё, что происходит в системе; ясно, что должны существовать такие (маловажные с точки зрения обучения системы) процессы в «организме» Доллии, которые в хронике не будут фиксированы вообще никогда (и о которых она, значит, вообще сама ничего не будет знать, пока в какой-нибудь школе ей о них не расскажут).

Доллии, чтобы определить, записывать его в хронику или нет, и ограничиться ли только упоминанием о событии или записать информацию о нём более подробно. Назовём этот второй принципиальный блок Хроникселектор (Х-селектор¹²⁶).

.501. Как опытным программистам, нам понятно, что Х-селектор технически может быть реализован как подпрограмма или компонент Хроникера, или как параллельный процесс, который взаимодействует с Хроникером через общее поле (выбор между двумя ситуациями не имеет значения)¹²⁷. Мы выбираем первый вариант и считаем, что Хроникер состоит из двух частей или субблоков. Один (Х-селектор), отбирает, что записывать и в каком виде в память, другой управляет собственно процессом записи (назовем его, скажем, «Хроникскриптор» (или Х-скриптор¹²⁸), а название «Хроникер» оставим для их совместного обозначения).

.502. Далее, проектируя и детализируя работу Х-скриптора, мы должны определить, как записать в память ту или иную информацию и какая нужна для этого структура данных, а, детализируя работу Х-селектора, мы должны подумать об алгоритмах, которые будут определять, что именно фиксировать в Хроникере, а что нет.¹²⁹

.503. Как мы видели в примере с книгой, информация о некотором действии может попасть, а может и не попасть под «всевидящее око»¹³⁰ Хроникера¹³¹ (информации не попала в Хроникер, так как я о чём-то «задумался», и именно то, о чём я задумался в этот момент записывалось в Хроникер, потому что Х-селектор посчитал «мои мысли» более важной информацией, чем информация о перемещении книги).

.504. Образно деятельность Хроникера мы можем представить себе как прожектор над головой Долли, который непрерывно освещает ночные «джунгли». Луч света выхватывает из ночи и освещает тот или иной участок зарослей, но некоторые «медвежьи углы», могут так и остаться в темноте навсегда.¹³²

¹²⁶ СМ. Сокращено мною (и звучит, как мне кажется лучше ☺).

¹²⁷ СМ. Мы говорим об операционной системе реального времени, где множество (сотни и более?) программ работают одновременно в режиме реального времени. Выборка и запись в Хроникер информации о работе программ должна идти параллельно с самими программами, чтобы не задерживать выполнение действий Доллии. В системах обработки транзакций (с которыми я работал) под каждую новую транзакцию запускался отдельный процесс, но поскольку каждый из этих процессов фактически обрабатывался последовательно, я мог наблюдать, как при некотором количестве одновременно открытых процессов система «захлёбывалась» и полностью теряла способность обрабатывать новые входящие транзакции. При обсуждении этого абзаца может конечно сказаться разница в понимании терминологии, но мне кажется, что Хроникер должен запускать под каждую программу новый процесс, который должен работать в параллельном режиме с другими такими же процессами.

¹²⁸ СМ. Сокращено тоже мною.

¹²⁹ СМ. И это самые главные вопросы в более детальном (дальнейшем) обсуждении Хроникера.

¹³⁰ В.Э.: в «поле зрения».

¹³¹ СМ. При использовании выражения «всевидящее око», я несколько задумался, а правильно ли я понял предыдущие двадцать абзацев. Думал довольно долго и пришёл к выводу, что если «всевидящее око» находится в Х-селекторе, то всё нормально (в русском языке «всевидящее око» имеет смысл видящего абсолютно всё, всё без утайки и всё сразу). Особенно, если «веко» этого «ока» находится в Х-селекторе, и похоже это «веко» больше на веко не очень чистоплотного инспектора, у которого оно тем плотнее закрыто, чем лучше инспектор «знает» (в смысле взаимовыгодной дружбы ☺) проверяемого.

¹³² СМ. Сравнение яркое, но мне не очень понятное. Хроникер в целом, я скорее бы сравнил с таможенным пунктом на границе двух очень дружественных стран, который работает по упрощённой таможенной процедуре. Грузовики, подъезжая к границе, перестраиваются в пять, шесть колонн (или больше, по числу открытых таможенных окошек; и чем больше к границе одновременно подходит грузовиков, тем больше таможенных окошек открыто) и на миг приостанавливаются у линии «Стоп», дожидаясь, когда таможенник переключит красный свет своего светофора на зелёный. Как правило, быстрым взглядом окинув грузовик и водителя, таможенник мгновенно зажигает зелёный свет, и грузовик с минимальной потерей времени уносится дальше в пункт конечного назначения. И только иногда, по только одному ему понятным признакам, таможенник может задержать чуть дольше грузовик и скопировать документы водителя, накладные и записать ещё какую-то информацию о грузе и тут же разрешить проезд дальше. В этих странах водители и таможенники очень сознательный народ, и явных нарушений таможенных правил быть не может, но вот ситуации (тип машины, тип груза, конечный путь назначения и т.д.) могут выпадать за рамки, описанные таможенными правилами, и в этом случае обязанность таможенника всю «странную», с точки зрения правил, информацию зафиксировать и передать в головную таможенную контору (Х-скриптор), которая все эти бумаги должным образом рассортирует, упорядочит и подошьет в дело, которое потом м.б. будет отправлено в суд, законодательное собрание и

§14. Проектор Хроникера

В.Э.:

2011.05.26 00:33 ночь на четверг

Ну, ты, Сергей, герой! Такие джунгли латышского языка прошел – и еще находясь в Техасе!

В основном ты понял всё правильно (и твои сокращения я принимаю), однако последний абзац всё же переведу точно:

.504. Образно общую деятельность хроникера мы можем представить себе как движение прожекторного луча по «джунглям» происходящих в голове Доллии процессов. Проектор может осветить ту или иную вещь, но имеется множество таких углов, в которые он не заглядывает никогда.

С этого образа проектора я как раз и собирался начать, когда обдумывал ответ на твой вопрос о памяти из предыдущего письма (§9). (А к начатому 17 мая ответу всё никак не мог вернуться, потому что еще во время твоего ремонта переключился на некоторые «детективы реального времени» – на серию книг {R-CHIKA1}, {R-CHIKA2}, {R-CHIKA3}, {R-CHIKA4} [= МОИ №66, № 67, № 68, №69] – скоро выставлю в Интернет – и не хотелось там всё оставлять в совершенно незавершенном виде).

Итак, принципиальное устройство человеческой памяти (и, соответственно, памяти Доллии) я представляю (или Веданская теория представляет) так:

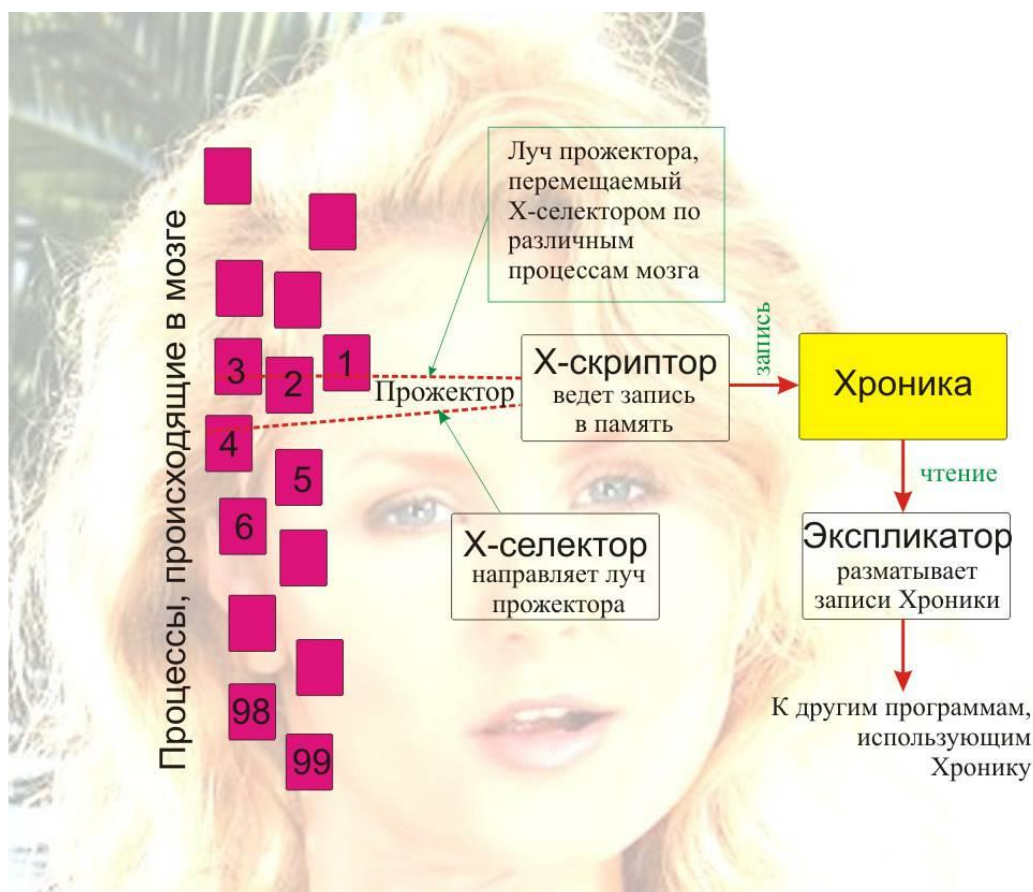


Рис.1. Принципиальное устройство Хроникера в голове Доллии

2011.05.26 15:39 четверг

Красные прямоугольники слева обозначают различные процессы, происходящие параллельно в многопроцессорном мозге Доллии. Х-скриптор ведет запись в память (в Хронику) тех

т.д., чтобы ситуация несоответствия существующих правил с реальным порядком дел больше не повторялась.

данных, на которые в данный момент направлен его «луч прожектора». (В примере рисунка «поле зрения» луча касается процессов 1, 2, 3, 4). В данный момент эти процессы проходят «осознанно» (потому, что сведения о них потом можно будет почерпнуть из Хроники).

Х-селектор может изменить направление луча. Например, он может «потянуть за зеленую стрелку» и направить луч на процессы 4, 5, 6. Тогда процессы 5 и 6, прежде «неосознанные», станут «осознанными», но зато из «сознания» выпадут процессы 1, 2, 3 (при фиксированной «ширине луча», т.е. пропускной способности Х-скриптора).

Но есть и такие процессы мозга (к примеру 98, 99), на которые Х-селектор никогда не в состоянии направить Прожектор. Такие процессы всегда остаются «бессознательными».

Поле зрения Прожектора всегда ограничено. Это обусловлено и вообще технической невозможностью записывать в память абсолютно всё, и нецелесообразностью этого – такая «захлавленная» Хроника не может быть использована столь же быстро и эффективно, как отобранная и хорошо организованная Хроника (Постулат 1: для выживания в естественном отборе нужна эффективная Хроника, а не свалка мусора).

В компьютерных аналогах Хроникер можно себе представлять так. Красные прямоугольники – это различные программы и их входно-выходные данные, находящиеся в различных адресах оперативной памяти. Луч прожектора – это общее (для Х-скриптора и Х-селектора) поле Х, содержащее адреса (указатели). Х-скриптор работает так, что просто хватается куски данных, находящиеся по этим адресам, и запикивает их на диск (в Хронику). А Х-селектор (по каким-то своим соображениям) меняет это общее поле Х и, записав туда (частично или полностью) новые адреса, тем самым направляет Х-скриптор к другим местам оперативной памяти (к другим процессам и их данным).

Поиск данных в Хронике – это уже работа других программ (не Хроникера). На схеме Рис.1 эти программы обозначены как Экспликатор.

§15. О памяти

Теперь я отвечу на твой конкретный вопрос из §9:

Можешь ли ты подробнее раскрыть, как информация хранится в памяти человека, и можно ли с помощью Веданской теории определить, так ли уж верны высказывания в печати о том, что человек использует порядка 5% памяти?

Тут вообще, если отвечать «по полной программе», то можно неделю работать. Поэтому попытаюсь по возможности коротко и конспективно:

1. Веданская теория – это теория программистская. Нас, программистов (во всяком случае, пока мы остаемся в своей области) вообще не интересует, на каких именно физических процессах основана память компьютера. Для нас важно только одно: что информацию можно туда записать и потом оттуда прочесть. Поэтому физическая природа человеческой памяти – это не область Веданской теории: этот вопрос должны решать конкретные исследования мозга (нейрофизиология).

2. Традиционно человеческая память делится на «краткосрочную» (по-программистски – оперативную) и «долговременную». В схеме Рис.1 процессы 1, 2, 3, ... 99... ведь тоже находятся в памяти – но очевидно «краткосрочной»; Хроника же наша находится в долговременной памяти (но Хроника не единственное, что там «сидит»; там находятся также, как минимум, заготовки неработающих в данный момент программ, возможно, еще другие структуры). Поэтому «память» – это довольно расплывчатое понятие.

3. Утверждения, что человек лишь в незначительной степени использует свою память, слышались давно – с дней моей молодости и раньше. Но они требуют уточнения вопроса: что именно имеется в виду.

4. Имеется ли в виду, (в терминах схемы Рис.1) что Х-скриптор пишет Хронику всё дальше и дальше на пустом пространстве памяти, и за жизнь человека он покрывает только 5% имеющегося пространства (это как если бы у нас был диск на 1 TB, а у нас нашлось информации только на 50 GB). Это ли имеется в виду?

5. Или имеется в виду, что у нас записан в Хронике 1 TB, но при чтении и поиске Экспликатор может найти и использовать только 50 GB из них?

6. Больше я встречался с утверждениями второго рода, и они основываются на точке зрения, что человек вообще никогда ничего не забывает по-настоящему – он только теряет доступ к запомненной информации, т.е. не может её больше найти и использовать. (А эта точка зрения, в свою очередь, основывается на том, что под гипнозом он может вспомнить такое, что не может вспомнить в нормальном состоянии).

7. Эта последняя точка зрения, безусловно, в какой-то степени истинна: «забывание» – это в первую очередь потеря доступа к информации. Но вот исчерпывается ли «забывание» только потерей доступа, а физического стирания информации вообще нет – это другой вопрос, и он тесно связан с вопросом пункта (4): пишет ли X-скриптор всегда на «чистое», свободное место памяти, или он пишет все-таки поверх старой информации, частично её стирая?

8. Лично я думаю, что вариант пункта (4) довольно трудно реализуем в природе (какое такое «свободное место» может быть в мозге, намного превышающее реальные его потребности – опять же Постулат 1: какие преимущества наличие такого неиспользуемого пространства давало в борьбе за выживание в процессе естественного отбора, чтобы такой механизм мог быть создан?).

9. В кукле Доллии мы (теоретически) можем такой механизм встроить – и на её интеллект это не повлияет –, но в природе, мне думается, его нет. Мозг пишет, скорее всего, поверх старой информации. Только он пишет хитро – не так, как настольный компьютер. Настольный компьютер имеет файл в одном экземпляре (не считая страховочных копий): записал поверх новый файл – и старого нет. А у мозга «файл», скажем, в 10⁰⁰⁰ экземплярах; новый файл тоже пишется в 10⁰⁰⁰ экземплярах, но не поверх какого-то одного файла, а поверх 10⁰⁰⁰ файлов так, что те не стираются полностью, а просто становятся более «редкими»: имеются уже только в 9⁹⁹⁹ экземплярах (это, разумеется, только объяснение принципа, а не точная картина). Таким образом, при записи новой информации старые файлы становятся всё труднее и труднее найти при (рандомизированном) поиске (т.е. они «забываются»), экземпляров файлов становится всё меньше (воспоминания «блекнут»). В конце концов у некоторых файлов исчезает последний экземпляр – но не у всех.

10. Так, я думаю, это принципиально устроено у человека. А насчет цифры 5% – тут требуется уточнение в постановке вопроса.

О памяти можно говорить еще очень и очень много. Но, я думаю, этих вещей мы коснемся тогда, когда ты поднимешь конкретные вопросы. А пока что отправлю эту книгу такой, какая она есть на данный момент.

§16. Письма Сергея Марьясова от 31 мая и 7 июня 2011 года (Gsvano)

no Sergey sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 31. maijs 16:49
temats RE: переписка R-POTI-3
piegādātājs rambler.ru

Валдис, добрый день.

У меня готов конспект 29 параграфа (§29 Emociators (emoġenerators и emoselektors)). Но мне нужно поработать ещё над своими комментариями.

СМ

no Sergey sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 7. jūnijs 22:24
temats FW: переписка R-POTI-3
piegādātājs rambler.ru

Валдис, добрый день.

Работа над R-POTI-3 интересна тем, что по её ходу возникает много самых разных интересных мыслей (идей). Так и в этот раз: до комментариев к 29 параграфу пока не добрался, но подготовил материал, который ближе к ходу повествования об операционной системе Доллии.

СМ

Валдис, спасибо за поддержку моих усилий по конспектированию отдельных параграфов книги L-ARTINT. Я и сам не думал, что надо уехать так далеко из Латвии, чтобы начать читать книги на латышском языке ☺. Хотя, конечно, дело не в расстоянии, а в содержании книги!

Попытался я законспектировать «§27. Неизвестная операционная система» (L-ARTINT), но после нескольких попыток отказался от этой идеи: текст скорее литературный, чем технический, а кроме того, в параграфе 2 нашей переписки ты вступительную часть уже охватил, хотя формально вступительные параграфы в двух книгах (§27 L-ARTINT и §2 R-POTI-3) отличаются. В §27 (L-ARTINT), в качестве вступления, тобою сравниваются две похожие задачи: задача построения операционной системы, аналогичной операционной системе неизвестного компьютера¹³³, и задача разработки операционной системы биологического компьютера (человека).

И я бы не стал возвращаться к вступлению, если бы не встретился с похожей постановкой вопроса на форуме «Искусственный интеллект: Ваши идеи»¹³⁴. Один из участников форума под псевдонимом *gsvano* писал¹³⁵:

«...Если бы перед тобой положили компьютер, и, не давая его разбирать, попросили сделать такое же устройство, какова вероятность, что ты бы скопировал механизмы его работы (то есть получил все его возможности), только изучая поведение работающих на нем программ? И это аналогия всего лишь с компьютером, мозг неимоверно сложнее, так как эволюционировал миллиарды лет...»

Поэтому я, в значительной мере под впечатлением книги L-ARTINT, поместил на форуме пост¹³⁶, несколько абзацев из которого могут послужить дополнительной иллюстрацией к §2 (R-POTI-3):

*

¹³³ **СМ.** «...например, если к нам в руки попал неизвестный нам компьютер с незнакомой, но работающей, операционной системой (скажем, наши доблестные разведчики добыли эту штуку на территории «вероятного противника» и передали нам, как специалистам по компьютерным системам, для изучения). Какая система команд (инструкций) используется в этом компьютере, мы не знаем, исходного текста программы у нас нет, деассемблировать программу мы не можем... Мы можем только запустить программу на выполнение и посмотреть, как программа обрабатывает в различных ситуациях и подумать, как такие же результаты в таких же ситуациях могут быть достигнуты...» L-ARTINT, «§27. Неизвестная операционная система», абзац 475.

¹³⁴ **СМ.** <http://ai.obrazec.ru/forum/> **В.Э.:** Я посмотрел этот сайт. Но он мне не очень понравился (во всяком случае – пока; впрочем, я его видел и раньше). Когда я захожу в подобные сайты-форумы, то у меня первые два вопроса таковы: 1) «КТО держит этот сайт, кто его владелец?» и 2) «Каковы права гостей, зарегистрированных участников, и что требуется для регистрации?». Когда мне эти два вопроса ясны, то уже можно думать, стоит ли там регистрироваться или нет, и т.д. Так, например, я знаю, что сайт JFKMURDERSOLVED.COM держит Wim Dankbaar, голландский автор книги «JFK Murder Solved»; сайт <http://www.serial-killers.ru/> держит MAzZY, профессиональный разработчик интернетовских форумов; сайт <http://filozofija.lv/> держит Авеллано, наст. имя Рейнис Лазда, был студентом философии в ЛУ, когда открыл сайт, теперь, наверное, уже окончил (на всех этих сайтах я зарегистрирован, хотя в общем-то давно ни в чем не участвую). Так вот, а на сайте <http://ai.obrazec.ru/forum/> я не смог установить, кто его, собственно, держит (думаю, это были какие-то юноши-студенты, когда 31 октября 2002 года открыли сайт; возможно тот самый *gsvano* и есть владелец – уж больно он активный... Впрочем, нет – он зарегистрирован лишь в январе 2005 года (хотя это могла быть и перерегистрация в связи с какой-то реорганизацией)). Ну, это уже нехорошо – не люблю я анонимные сайты. Далее: на порядочных сайтах рядом с окошками «Имя пользователя» и «Пароль» стоит кнопка «Регистрация»: если ты уже зарегистрирован, то вводишь в рамки всё, что надо, а если не зарегистрирован, то регистрируешься. А у этого сайта кнопки «Регистрация» рядом с окошками НЕТ! Черт знает, куда спрятали – ищи, как дурак! (Я так и не посмотрел, что требуется для регистрации). Значит, квалифицируем их сайт как «Плохо организованный». Ты там зарегистрировался, кажется, 2 апреля. Но сегодня, 30 июня, спустя 3 месяца, ты имеешь статус «Кандидат в участники форума». (Да что эти сопляки – как в КПСС принимают: с годовым кандидатским стажем?! Да пошли они...) Посмотрел некоторые посты, помимо твоих, – в общем-то пустая болтовня, как обычно в интернетовских форумах... Статьи, выставленные там, может быть, можно почитать, но они, насколько я видел, не их собственные, они из журналов и чужих сайтов, и статей не очень много, всего около 50. Так что у меня впечатление такое: это мальчишки, не очень умные, но зато очень заносчивые...

¹³⁵ **СМ.** Правила форума позволяют свободно использовать размещённые на нём материалы.

¹³⁶ **СМ.** «Пост», в терминологии форума, – ответ, утверждение или некоторое сообщение.

«...Именно такая постановка вопроса и затрудняет постройку ИИ. А спроси хорошего программиста, напишет ли он редактор типа *MSO Word* при наличии времени (и желания), то он даже спрашивать не будет, на каком компьютере, под какой операционкой, и на каком языке программирования эту задачу решать (у программиста на этот счёт всегда есть свои особые идеи). Я в такого рода обсуждениях, покупать ли готовую функциональность у поставщиков программного обеспечения, или написать своими силами, участвовал неоднократно. Для программистов было важно понять, какая функциональность необходима (что должна делать программа или система), какая информация на входе и выходе (интерфейсы), в каком она виде (т.е. формат представления данных выходной информации), и очень редко спрашивали, на чём это написано у поставщиков программного обеспечения (или конкурентов), и ещё реже, под какой операционной системой это работает и на каких серверах. А что там у поставщиков (или конкурентов) в исходниках, никто и не разбирался, да и, в общем случае, доступа к такого рода информации не было...

То есть главное – правильная постановка задачи и наличие ресурсов для её реализации.

То же самое можно сказать и по отношению к мозгу. Да, хорошо, что мы приблизительно знаем, как работают нейроны (микропроцессор *INTEL*, *AMD* и т.д.), как работает нейросеть (операционная система *Windows*, *UNIX* и т.д.), взаимодействуют разные части головного мозга (конфигурация серверов, накопителей памяти, сетевых устройств и т.д.), но нам важнее выделить функциональность, которая определит то, что мы интуитивно называем интеллектом. И самое главное, что лезть для этого внутрь мозга (платы с микропроцессором, исходные коды операционной системы, соединительные кабели в серверной) не нужно. Поведение человека, особенности его психической деятельности, мышления и т.д., с точки зрения функциональности изучены довольно хорошо. Осталось только эти знания правильно скомпоновать или, вернее, выделить из них то, что позволит программисту эту функциональность запрограммировать...

Приведу пример: постройка самолёта. Со времён Дедала люди пытались построить крылья, похожие на птичьи, изучали их форму, размах, структуру пера, необходимость пуха... А полетели только тогда, когда выделили понятие подъёмной силы. (Птиц природа, кстати, тоже не один миллиард лет создавала)...»¹³⁷

*

Конспект параграфа 28 «Хроникер (X-селектор и X-скриптор)» (*L-ARTINT*) дался несколько легче (чем параграф 27), отчасти из-за того, что, просматривая литературу, которая могла бы быть полезна для рассмотрения наших вопросов (стр. 12, стр. 14 настоящей книги) я наткнулся на книгу К. Шереметьева «Полноприводной мозг. Как управлять подсознанием», 2007 г.

Эта книга, на мой взгляд, довольно хорошо освещает вопросы о том, как работают сознание и подсознание человека¹³⁸ («осознанное» и «неосознанное» в нашей терминологии), взаимосвязь подсознания и сознания и какова их роль в формировании поведения человека. Объяснение, как это работает (твоими словами «устроено» ☺) у человека, и большое количество понятных примеров дано в книге под нужным углом для разработки ИИ и написано простым и понятным языком.¹³⁹ Описание функционирования сознания и подсознания (функциональность биологического компьютера!) дано без упоминания мозговых программ, но это лишь вопрос терминологии. Идеи, утверждения, выводы книги легко проверяются (достаточно обратиться к анализу собственных поступков, опыта и т.д.) и, главное, не выпадают, а логически укладываются в описание работы модуля Доллии, названного Хроникёром.

И ещё, говоря о Хроникере, я хотел бы выделить следующий твой тезис: «...объектами «интереса» хроникера являются не столько сами программы, сколько те результаты, которые они производят (и те условия, которые вызвали их запуск) – в общем, те данные, с которыми они работают на входе и выходе. (Отчасти благодаря этому собственно программы остаются для человека «невидимыми», и он обычно не подозревает о их существовании и даже не верит в

¹³⁷ **В.Э.:** Единственное замечание к твоему посту: птицы всё-таки не эволюционировали «не один миллиард лет». Если считать началом их истории (и истории их крыльев) тот момент, когда некоторые пресмыкающиеся, типа псевдозухий, ползавшие по деревьям, стали планировать по воздуху вниз, то это триасский период около 230–195 млн. лет назад.

¹³⁸ **В.Э.:** Звучит заманчиво. Если дойдут руки, надо будет прочитать, т.е. поместить в Векордию (что одно и то же).

¹³⁹ **СМ.** М.б. сказывается то, что начинал свою трудовую деятельность автор программистом ☺ (<http://www.sheremetev.info/content/blogsection/4/35/>).

это, если ему говорят о программах).»¹⁴⁰ Размышляя о работе Хроникёра, я думал об этом, но в явном виде это важное утверждение не упомянул.

Не думаю, что на данном этапе работы нам следует детализировать работу отдельных модулей Доллии до уровня описания интерфейсов и полей данных; пока я лично ставлю перед собой более скромную задачу – разобраться с описанием блок-схемы «Модель психической деятельности человека в Веданской теории» и некоторыми вопросами, которые возникают, так сказать, «по пути»). Но, как часто это бывает при чтении особенно интересной книги, хочется заглянуть, что там ближе к концу ☺☺.

§17. Письмо Сергея Марьясова от 10 июня 2011 года (Эмоциатор)

no Sergey sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 10. jūnijs 17:10
temats R-POTI-3
piegādātājs rambler.ru

Валдис, добрый день!

Спасибо за развёрнутый ответ на мой вопрос о памяти человека. Действительно, теперь я вижу, что:

- а) мой вопрос был сформулирован нечётко;
- б) несколько направлений, затронутых в твоём ответе, было бы интересно обсудить в

дальнейшем.

Однако, я хотел бы вернуться к этому обсуждению после рассмотрения всех модулей блока памяти (из твоей блок-схемы «Модель психической деятельности человека в Веданской теории», L-ARTINT, с.112). Т.е. сейчас, в первую очередь, я для себя наметил пройти по всем модулям этой блок-схемы в том порядке, в котором они описаны в книге L-ARTINT.

К конспекту параграфа 29 (в приложенном файле) думаю добавить ещё пару абзацев с некоторым обобщением моих комментариев.

Всего наилучшего,

СМ.

Конспект Сергея Марьясова (довольно подробный) параграфа 29 «Эмоциатор (эмогенератор и эмоселектор)» из книги В.Эгле L-ARTINT

§ 29. Эмоциатор (эмогенератор и эмоселектор).

.505. Часто приходится слышать, что компьютер не может чувствовать каких-либо эмоций. И это понятно, что он не может, – но только до того момента, пока мы не встроим в него такую способность. Как и в предыдущем случае (имеется в виду рассмотрение работы Хроникёра)¹⁴¹, давайте рассмотрим, что значит объективно, что люди находятся в том или ином эмоциональном состоянии.¹⁴² Тогда мы увидим, например, что гнев человека – это приготовление

¹⁴⁰ СМ. Прим. 115, R-POTI-3 [здесь 118].

¹⁴¹ СМ. Добавлено мною.

¹⁴² СМ. В Wikipedia (на русском языке) я подчеркнул два полезных факта относительно современного понимания эмоций:

1) В каждой эмоции человека можно (достаточно условно) выделить инстинктивную часть (достались нам от наших далёких предков) и некоторую часть, сформировавшуюся под воздействием общественной жизни человека: «...Эмоции эволюционно развились из простейших врождённых эмоциональных процессов, сводящихся к органическим, двигательным и секреторным изменениям, до значительно более сложных, утративших инстинктивную основу процессов, имеющих отчётливую привязку к ситуации в целом, то есть выражающих личное оценочное отношение к имеющимся или возможным ситуациям, к своему участию в них... Выражение эмоций имеет черты социально формирующегося, изменяющегося с течением истории языка, что можно видеть из различных этнографических описаний.» (Материал из Википедии – свободной энциклопедии. Статья «Эмоциональный процесс». <http://ru.wikipedia.org/wiki/>)

к некоторым агрессивным действиям, страх – приготовление к бегству или другим действиям самосохранения, и т.д.¹⁴³

.506. Например, состояние гнева характеризуется тем, что компьютер Доллии принял к исполнению общую стратегию атаки, активизировал соответствующую группу программ, и Доллия готова в любой момент их исполнить (чтобы, скажем, ругать, наносить удары, кусаться, царапаться и т.д.). Состояние страха, напротив, характеризуется принятием общей стратегии выхода из некоторой угрожающей ситуации и включением необходимых процессов для её реализации.

.507. Нам понятно, что в голове Доллии параллельно будет происходить множество различных процессов. Каждый из них может работать с разной интенсивностью,¹⁴⁴ начиная от нуля (полностью выключен) и заканчивая максимальным. Общая картина распределения

2) «В научном сообществе существует множество различных взглядов на природу эмоциональных процессов. Какой-то одной, общепринятой теории до сих пор не выработано...» (Материал из Википедии – свободной энциклопедии. Статья «Эмоциональный процесс». <http://ru.wikipedia.org/wiki/>) – это утверждение, по сути, является общей оценкой многообразия подходов и теорий к пониманию эмоций, описанных на англоязычной страничке *Wikipedia* (<http://en.wikipedia.org/wiki/Emotion>).

На этой же страничке *Wikipedia* я для себя отметил тезис Дарвина о том, что эмоции были выделены и развиты в результате естественного отбора, и поэтому похожи и узнаваемы для разных народов и культур (кроме того, по этой же причине люди узнают похожие эмоции не только в поведении представителей других народов, но и у многих видов животных). «В начале 70-х (20-го столетия), Paul Ekman с коллегами начал исследовательскую работу, в результате которой удалось доказать, что всем людям присущи, по крайней мере, пять базовых эмоций: страх, печаль, счастье, злость, и отвращение.» (Gaulin, Steven J.C. and Donald H. McBurney. *Evolutionary Psychology*. Prentice Hall. 2003. ISBN 13: 9780131115293, Chapter 6, p 121–142.)

¹⁴³ **СМ.** Для рассматриваемых в этом параграфе вопросов может оказаться полезным определение аффекта, который Л.А. Леонтьев выделил среди эмоциональных процессов (деление эмоциональных процессов по Леонтьеву: аффекты, эмоции, чувства и настроения).

«Аффект, как и любой другой эмоциональный процесс, представляет собой психофизиологический процесс внутренней регуляции деятельности, и отражает бессознательную субъективную оценку текущей ситуации. Его уникальными характеристиками являются кратковременность и высокая интенсивность, в сочетании с выраженными проявлениями в поведении и работе внутренних органов. У животных возникновение аффектов связано с факторами, непосредственно затрагивающими поддержание физического существования, связанными с биологическими потребностями и инстинктами. Содержание и характер аффектов человека претерпевает значительное изменение под влиянием общества, и они могут возникать также в складывающихся социальных отношениях, например, в результате социальных оценок и санкций. Аффект всегда возникает в ответ на уже сложившуюся ситуацию, мобилизуя организм и организуя поведение так, чтобы обеспечить быструю реакцию на неё.» (Леонтьев Алексей Николаевич. *Потребности, мотивы и эмоции* (рус.). Москва. 1971.)

¹⁴⁴ **СМ.** В каком смысле используется термин «интенсивность»? В компьютере есть перераспределение ресурсов, повышение приоритета задач в интересах процессов или программ, которые требуют незамедлительного выполнения, особенно, если речь идёт о работе компьютерных систем в режиме реального времени. Но в компьютере нет ничего похожего на интенсивность работы того же землекопа (который вдруг начинает бросать землю быстрее ☺).

Правда, во времена микропроцессоров *Intel 386* и *Intel 486* среди особо увлечённых компьютерщиков была широко распространена практика увеличения тактовой частоты микропроцессора (так называемый «разгон» центрального процессора), основной целью которой было повышение производительности процессора и, в результате, всего компьютера в целом. Позже количество компонентов компьютера, поддающихся разгону, увеличилось. Причём, если в ранних компьютерах разгон осуществляли энтузиасты, хорошо разбирающиеся в электронике и способные «перепаять» некоторые элементы (например: заменить тактовый генератор на генератор с более высокой частотой), то в настоящее время для этого существуют программные средства, разработанные самими производителями микропроцессоров.

В биологическом компьютере увеличение производительности мозга в моменты, когда требуется срочное принятие некоторой стратегии действий и её незамедлительное выполнение, достигается на биохимическом уровне (аналог повышения тактовой частоты?). Однако такого рода воздействие допускается кратковременно, чтобы не допустить непоправимого ущерба живым элементам головного мозга. С разгоном процессоров происходит нечто похожее: повышение тактовой частоты снижает надёжность работы устройства (увеличение вероятности или количества сбоев), а в некоторых случаях ведёт к выходу из строя внутренних элементов микропроцессора (из-за перегрева, например). Поэтому производители микропроцессоров и выставляют некоторое усреднённое значение тактовой частоты, гарантирующее надёжную (без сбоев) и долговременную работу устройства, но для каждого конкретного микропроцессора эта частота может быть несколько увеличена.

интенсивности процессов будет определять текущее внутреннее состояние системы Доллии, которое традиционно называют эмоциональным состоянием.

.508. Таким образом, чтобы реализовать эмоции Долли, нам понадобятся: 1) возможность изменения интенсивности различных процессов, 2) аппарат, который управляет изменением интенсивности, уменьшением или полной ликвидацией одного процесса и включение или повышение интенсивности других. Блок, в котором будет это реализовано, назовём «эмоциатором».

.509. Как и в предыдущем случае (имеется в виду Хронекер)¹⁴⁵, в эмоциаторе можно выделить две группы функций. Одна из них управляет интенсивностью процессов (уменьшить, выключить, включить, увеличить). При более детальном проектировании, мы должны подумать о том, какими процессами конкретно управлять, как осуществлять связь эмоциатора с ними и как влиять на их интенсивность.

.510. Другая группа функций следит за ситуацией, в которой Долли в настоящее время находится, и, исходя из неё, должна решать, какое именно «эмоциональное состояние» должно быть организовано (т.е. какая именно стратегия поведения системы будет выбрана). При более детальном проектировании мы подумаем о том, по какому алгоритму мы будем определять ситуации, в которых необходимо запускать механизмы конкретных чувств: бегства (страха) или агрессии (гнева), или радости (когда ситуация признается как исключительно благоприятная), *kad atzīt par bezcerīgu, kaut arī tieši dotajā brīdī nekas nedraud (skumjas) utt.*¹⁴⁶

.511. Для реализации этих двух функциональных групп мы выделим в составе эмоциатора два отдельных и концептуально различных по принципам работы субблока.¹⁴⁷ За первую группу функций будет отвечать субблок, который назовём, например, «эмогенератором», а за вторую группу будет отвечать субблок, который назовём, скажем, «эмоселектором». Таким образом, на данном уровне детализации эмоциатор будет состоять из этих двух субблоков.

.512. Чем более подробные алгоритмы мы встроим в эмоселектор, и чем большее количество процессов будет управляться эмогенератором, тем большее количество различных эмоциональных оттенков и нюансов мы можем получить в поведении Доллии.

.513. Таким образом, с построением у Доллии эмоционального аппарата проблем нет (по крайней мере в теории). Однако, может возникнуть вопрос, а нужен ли этот аппарат? Из опыта человеческого общества мы знаем, что эмоции влияют на результаты человеческой деятельности скорее негативно,¹⁴⁸ и, как правило, более эффективно функционируют люди, которым удается

¹⁴⁵ СМ. Добавлено мною.

¹⁴⁶ СМ. Я перевёл для себя как: «или безнадежная, или ситуация в которой ничто не беспокоит (скучать) и так далее». Но уверенности в правильности перевода нет ☹. **В.Э.:** Предложение в целом было таким: «При более детальном проектировании мы будем здесь думать о том, по какому алгоритму устанавливать те ситуации, в которых должны быть включены механизмы бегства (страх), механизмы агрессии (гнев), когда признать ситуацию чрезвычайно хорошей (радость), когда признать ее безнадежной, хотя именно в данный момент ничто и не угрожает (тоска) и т.д.». (Тоска – это когда «всё плохо», хотя непосредственной угрозы или физических страданий нет).

¹⁴⁷ **В.Э.:** Вообще по-русски лучше «подблок»; зачем использовать латинско-английские приставки, когда есть русские?

¹⁴⁸ СМ. Читая в детстве о походах Александра Македонского, я был поражён мужеством и стойкостью 8000 греческих гоплитов-наемников (10–12 тыс. по оценкам других источников), которые сражались на стороне персидского царя Дария в битве при Иссе (333 до н.э.). Сражение было напряжённым, но в некоторый момент армия Александра прорвала фронт почти 100 тысячной армии персов (60 тысячной по другим источникам): отдельные отряды персидской армии побежали, и постепенно растерянность, страх и паническое бегство охватили всё войско Дария. И только греческий отряд наёмников сохранял боевой порядок среди общего разгрома. Греческие наёмники твёрдо удерживали свою позицию даже несмотря на то, что воспользовавшись бегством соседних отрядов персов, македонцы сумели ударить им в открытый фланг.

«Спасайся, кто может!», «Каждый сам за себя!», «А чьи ноги быстрее!?», – вот лозунги (эмоциональное состояние) солдата, бегущего с поля боя; в этом состоянии он подчиняется только древнему чувству самосохранения, которое было «воспитано» естественным отбором в многократных спасительных забегах от мамонтов или саблезубых тигров ☺. «Воевать» со спиной противника гораздо легче, чем сражаться против его меча и щита, в результате войско персов от давки (неуправляемая масса в 50 тыс. человек на небольшом участке местности!) и преследовавших его македонцев пострадало сильнее, чем непосредственно в бою. Потери персов были огромными. Греческий же отряд, подчиняясь «дисциплине строя», чувству коллективной ответственности за стоящего рядом товарища и профессиональным

подавлять свои эмоции,¹⁴⁹ то есть свести к минимуму работу этого аппарата (хотя в отдельных случаях, конечно, работа эмоционального аппарата может оказаться полезной). Мы не пытаемся обсуждать здесь вопрос, хотело бы большинство из нас принести в жертву свой собственный эмоциональный аппарат или нет (или только желало бы, чтобы он генерировал «хорошие», или преимущественно «положительные эмоции»¹⁵⁰). Мы задаём себе вопрос только о том, а нужен ли этот аппарат с точки зрения «чистого проектирования». Или, другими словами, мы собираемся встроить в Доллию эмоциональный аппарат только для того, чтобы полностью имитировать человека, или же это является необходимым условием для построения такого рода систем?

.514. Как мы видели, суть работы эмоционального аппарата такова: в зависимости от принятой [в данной ситуации] основной стратегии [поведения] активизировать, подготовить к немедленному действию определенные процессы.¹⁵¹ Этот принцип мы должны признать действительно необходимым, так как система, которая заранее не подготовится к предстоящим действиям, будет реагировать гораздо медленнее. Если человеку использование такого принципа часто мешает «правильно» реагировать, то это означает, что алгоритмы его эмогенератора не соответствуют целям такого реагирования, которое он (с точки зрения «рационального мышления») признает «правильным»: в повышенную готовность приводятся ненужные, «неправильные» процессы.¹⁵²

.515. Это несоответствие между целями и алгоритмами эмоционального аппарата и «рациональным мышлением» у людей, конечно, развилось исторически: в далёком прошлом, у животных предков человека «рационального аппарата» ещё не существовало, и эмоциональный аппарат был единственным средством генерации ответа.¹⁵³ По сравнению с более поздним «рациональным аппаратом», эмоциональный аппарат действует на сравнительно простых принципах или алгоритмах, в основном эгоистических, однако таких, которые, несомненно, способствовали выживанию индивида в борьбе за существование.

.516. Таким образом, разрабатывая операционную систему Доллии, у нас есть выбор: либо последовать человеческому примеру и встроить в Доллию такой же эгоистически ориентиро-

бойцовским качествам (т.е. «рациональному мышлению»), отступил, сохраняя боевой порядок, и вышел из сражения практически без потерь!

Конечно, если бы Александр Македонский организовал окружение этого отряда, то вряд ли кто-то из греков-наемников остался в живых, но Александр был занят преследованием персидского царя Дария, а его военачальники истреблением более легкой добычи – бегущих персов...

¹⁴⁹ **СМ.** «...Бывший офицер ГРУ, а ныне руководитель школы рукопашного боя Анатолий Тарас рекомендует в драке входить в образ «боевой машины», сметающей все и всех на своем пути. В этом образе исчезают сомнения и страхи. Боец делается необыкновенно решительным и заметно увеличивается скорость реакции...» (К. Шереметьева «Полноприводной мозг. Как управлять подсознанием», 2007 г.) В этом примере происходит не только подчинение ситуации «рациональному мышлению», но и вызов через некоторый образ (осознанная деятельность!) смены эмоционального состояния (в данном случае вызывается состояние решимости), что, в свою очередь, приводит к ускорению программ управления мышцами, зрением и т.д., а в том числе ускоряет и процессы обмена веществ в организме.

¹⁵⁰ **СМ.** Положительные эмоции оказывают благотворное влияние на весь организм человека. Попробую пересказать (кратко) суть открытия М. Норбекова (в соавторстве с Н. Бордюк): через осознанные процессы (желание выздороветь) и неосознанные процессы (осанка, мимика) и эмоциональные центры (общий эмоционально-радостный, счастливый настрой) можно влиять на жизнеспособность биосистемы человека и тем самым добиваться положительных результатов по лечению многих, в том числе и ранее считавшихся неизлечимых, заболеваний. (М. Норбеков. *Опыт Дурака, или ключ к прозрению как избавиться от очков*. Москва. 2008 г. Стр. 104–106).

¹⁵¹ **В.Э.:** То есть, сами по себе эмоции (такие как страх, печаль, счастье, злость и отвращение) еще не есть действия; это есть только готовность к действиям того или иного рода (с легким и быстрым включением самих действий при малейшем дальнейшем импульсе).

¹⁵² **В.Э.:** Ну, например, человек впервые выступает перед широкой аудиторией; надо бы мобилизовать все его умственные способности, а вместо этого его эмогенератор (унаследованный от животных предков) мобилизует кровоснабжение (сердце начинает биться часто), подготавливает охлаждение организма (сильное потоотделение), учащает дыхание – человек волнуется, что отнюдь не помогает ему успешно выступать перед аудиторией.

¹⁵³ **СМ.** Хочу повторно сослаться на определение аффекта в [прим.3](#).

ванный эмоциональный аппарат, либо согласовать его с какими-то целями «рационального мышления».¹⁵⁴

§18. Письмо Сергея Марьясова от 20 июня 2011 года (обобщение)

no Sergey sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 20. jūnijs 05:59
temats R-POTI-3
piegādātājs rambler.ru

Валдис, добрый день!

Посылаю некоторое обобщение к моему пониманию параграфа 29.

Всего наилучшего,

СМ

Обдумывая тезисы 29 параграфа, я для себя построил следующую классификацию эмоциональных процессов человека:

1) инстинктивные эмоциональные процессы, доставшиеся нам от наших далёких предков и представляющие собой совокупность физиологических процессов и некоторого базового набора программ (точнее подпрограмм) поведения. Эти программы восстанавливаются в мозгу человека по записям в ДНК по мере роста мозга.

2) эмоциональные процессы, сформировавшиеся под воздействием общества («рациональное мышление»). Это те же физиологические процессы, но к ним привязаны программы, которые человек приобрёл в результате обучения (обществом и на базе знаний общества) и самообучением, и которые являются либо видоизменением, дополнением программ, доставшихся нам «по наследству», либо, по-сути, замещают «наследственные» программы. «Наследственные» программы при этом никуда не исчезают, они просто не вызываются «операционной системой» человека.

И в том и другом случае в эмоциональном процессе человека участвуют две составляющие: физиологическая и программная. (Выражаясь языком информатики, это аппаратно-программная реализация.) Между этими двумя составляющими существует обратная связь. Положительные эмоции усиливают работу мозговых программ (особенно ту их часть, которая выполняется неосознанно); достижение промежуточных положительных результатов, в свою очередь, приводит к усилению радостных эмоций, которые благотворно воздействуют на всю биосистему человека; и это, в свою очередь, стимулирует дальнейшую работу мозговых программ. Отрицательные эмоции при длительном воздействии наносят ущерб физическому состоянию человека и ведут к усилению отрицательных эмоций (примером такого состояния может быть депрессии).

Нужно ли нам реализовать в случае Доллии аппаратную часть эмоций? Я согласен с К. Шереметьевым, что «...всё, что человек делает, он делает ради одного – чтобы испытать какие-то чувства.»¹⁵⁵ Для человека очень важно в результате некоторой намеченной деятельности, особенно если эту деятельность он определил для себя сам, получить ожидаемые эмоции (например, радость от законченного проекта, похвала друзей и т.д.).

Как будет стимулироваться интеллектуальная деятельность Доллии, если у неё не будет некоторого подобия того, что у человека реализовано на физиологическом уровне?

Рассматривая эмоциональные процессы человека под несколько другим углом, мы можем выделить в них присутствие «наследственных» процессов, вытекающих из индивидуалистических и/или стадных (коллективных) инстинктов (которые тоже реализованы на физиологическом и программном уровнях). Поэтому в случае Доллии возникает вопрос, будет ли Доллия системой, которая по каким-то причинам должна быть частью некоторого сообщества (похожих

¹⁵⁴ СМ. Где-то здесь находятся ключи к вопросам взаимоотношения человека и ИИ. Хотя вопросы соотношения эгоистического (личного) с коллективным (общественным) всегда находились среди вопросов, наиболее дискутируемых в обществе.

¹⁵⁵ СМ. В контексте книги термин «чувства» тождественен нашему термину «эмоции». (К. Шереметьев «Полноприводной мозг. Как управлять подсознанием», 2007 г.).

систем) или, будучи аналогом человека, ей достаточно будет общения с человеческим обществом?

Этот и предыдущий вопрос о стимулах Доллии предлагаю оставить открытыми до окончания рассмотрения блок-схемы «Модель психической деятельности человека в Веданской теории» (L-ARTINT, с.112) (наверное, предварительно также потребуется рассмотреть вопрос типологии людей).

§19. Письмо Сергея Марьясова от 24 июня 2011 года (Программатор)

no Sergey sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 24. jūnijs 16:04
temats R-POTI-3
piegādātājs rambler.ru

Валдис, добрый день!

Поработал с параграфом 30 «Генератор Программ и Анализатор Программ» (L-ARTINT)¹⁵⁶ и, конечно же, захотелось посмотреть на более детальное описание «стартовых блоков» ☺. Не думаю, что это надо делать прямо сейчас, но в некоторой перспективе было бы очень интересно подумать над этой задачей.

В приложенном файле конспект 30-го параграфа с некоторыми моими комментариями.

Всего наилучшего,

СМ

Конспект Сергея Марьясова (довольно подробный) параграфа 30 «Генератор Программ и Анализатор Программ» из книги В. Эгле L-ARTINT

§30. Генератор Программ и Анализатор Программ.¹⁵⁷

.517. Чтобы Доллия могла выполнить какие-то действия, ей, так же, как и любому другому компьютеру, необходима программа этих действий. Под программой, вообще говоря, мы здесь понимаем некоторую компьютерную структуру, которая имеется в наличии до начала выполняемых действий и однозначно определяет результаты этих действий (то есть однозначно определяет, что и каким образом будет сделано).

.518. Очевидно, что мы не можем снабдить Доллию всеми возможными программами наперёд, на все возможные случаи и для всех возможных ситуаций. Поэтому эти программы Доллии придётся подготавливать самой для каждой конкретной ситуации и в каждом конкретном случае и, как правило, непосредственно перед выполнением требуемых действий. Мы называем это самопрограммированием.

.519. Давайте подумаем, как реализовать механизм самопрограммирования в компьютере Доллии. Во-первых, опытному разработчику сразу понятно, что генерация программ должна осуществляться в структурированном виде и на нескольких уровнях. Ни один генератор не справится с задачей создания программы, состоящей, скажем, из ста тысяч отдельных инструкций в одном массиве без какой-либо структуры. С другой стороны, не составляет большого труда создать генератор, который должен собирать текущую программу из нескольких заранее подготовленных блоков простым их объединением.

.520. Мы так и поступим, и на верхнем уровне у нас будет генератор, который, в соответствии с текущей ситуацией, выбирает из находящихся в его распоряжении программных блоков необходимые блоки и соединяет их в нужном порядке. На данном уровне в этом и заключается вся генерация программ.

¹⁵⁶ <http://vekordija.narod.ru/L-ARTINT.PDF>, стр.105.

¹⁵⁷ СМ. В оригинальном тексте «Programgenerators un programanalizators». Но в русском языке такое длинное слово строить не принято (как-то «не звучит»), а сокращать название блоков в ПП и ПА (как это принято в технической литературе) не хочется: потом гадай и ищи, что за этими двумя буквами замаскировано ☺; а так как слово «программа» употребляется в тексте довольно часто, то, чтобы избежать путаницы, название блоков будем писать с большой буквы.

.521. Теперь давайте подумаем, откуда эти блоки появятся. Конечно же, их необходимо предварительно запрограммировать. В общем случае это можно сделать как непосредственно перед генерацией (сборкой) программы верхнего уровня, так и задолго до этого могут быть созданы заготовки, неоднократно уже побывавшие в работе: проверенные, модифицированные и выверенные. Но, – либо сейчас, либо в прошлом, – когда-то и эти [складываемые] блоки генерировались. И к их генерации относится всё то, о чём мы только что говорили в отношении верхнего уровня программы: эти блоки тоже могут быть сгенерированы только из относительно небольшого количества уже ранее существующих блоков более низкого уровня. И так мы продолжаем, спускаясь на всё более и более низкие уровни программ Долли, блоки которых генерировались на всё более и более ранних этапах её существования.

.522. В конце концов, мы приходим к тем блокам, что не были сгенерированы Доллией в процессе самопрограммирования, а были даны ей в готовом виде. Конечно, эти «стартовые блоки» должны дать ей мы – из ничего ничего не получится – но их количество сравнительно невелико и не привязано к конкретным ситуациям; мы и дадим их Доллии (для человека эти «стартовые блоки» являются врожденными, они встроены на генном уровне).¹⁵⁸

.523. Ну и, начиная со «стартового набора», Долли постепенно должна построить для себя блоки или заготовки программ всё более и более высокого уровня таким образом, чтобы, в конце концов, она могла из этих блоков быстро сгенерировать программу, которая требуется непосредственно для конкретной ситуации. Этот процесс сборки заготовок на всё более высоком уровне является центральным и наиболее важным элементом в процессе, который мы называем «обучением» (только малую часть этого процесса составляет простое запоминание информации)¹⁵⁹. Сначала, Доллия научится двигать руками и ногами, затем ходить и держать в руках разные предметы, шевелить языком и губами, издавать звуки, и, наконец, говорить и даже писать...¹⁶⁰

.524. Чтобы Доллия не ленилась и училась по-настоящему, пытаясь делать то одно, то другое, нам потребуется встроить в неё механизм поощрения и наказания, или (скажем более

¹⁵⁸ **СМ.** Задача, сразу скажем, непростая: создать некоторый аналог генома человека в той части, которая определяет механизм работы его мозговых программ, образования эмоций и других процессов, которые мы запланируем встроить в Доллию. **В.Э.:** Интеллект – безусловно сложная система, и решение задачи не может быть простым. Но гены создают как аппаратуру, так и программатуру; наша задача всё-таки проще – хардвер мы считаем уже данным. Стартовые блоки для Доллии должны носить такой характер, как «двинуть рукой», «двинуть ногой», «поменять программу в одном направлении», «поменять программу в другом направлении»... Вот, начиная с таких блоков, пусть она учится, как их комбинировать, чтобы результат получился таким, какой она хочет.

¹⁵⁹ **В.Э.:** То есть, обучение можно разделить на два вида мозговой деятельности: 1) запоминание фактов (таких, например, как «Лондон – столица Англии» или « $7 \times 8 = 56$ » (хотелось ради шутки написать «... = 48», но вдруг найдутся читатели, которые не поймут ☺); и 2) разработка заготовок своих программ (например, программы «Как двигать рукой, чтобы на бумаге получилась буква А». Так вот, основное в обучении, это не первое (как многие думают), а второе.

¹⁶⁰ **СМ.** Определение набора стартовых блоков и их содержание – сложная задача, сравнимая с пониманием результатов эволюции (почему у человека это «устроено так, а не этак») и возможность избежать при построении Доллии атавизмов, которые природа не могла просто «выбросить» в случае эволюционного «изобретения» человека. Этот набор должен, с одной стороны, довольно жёстко определить у Доллии то, что мы называем характером человека, обеспечить формирование того, что мы называем эмоциональными процессами, а, с другой стороны, стимулировать развитие знаний об окружающем мире (как у Человека Разумного, а не законсервировать некоторый уровень знаний, как у Неандертальца). **[В.Э.:** Тут, я думаю, всё обстоит иначе, но об этом поговорим потом.]

Геном человека, кроме того, обеспечивает хрупкий баланс между устойчивым характером системы (индивида) в определённом промежутке времени и возможностью существования (и развития в случае человека) вида (общества), к которому этот индивид относится, в гораздо более длительном периоде времени. Если в случае компьютера мы можем улучшить его работу (не меняя его аппаратную часть), поменяв содержание BIOS, или переустановив на нём новую операционную систему, то в случае человека, эта возможность предусмотрена (эволюцией) только путём рождения нового человека. В интернете имеется информация о том, что геном человека изучен (расшифрован) довольно хорошо (хотя по-прежнему существует необходимость более высокого уровня детализации) (например, <http://www.bibliotekar.ru/lldnk2.htm>). Наверное, было бы интересно посмотреть, насколько изменился генотип Человека Разумного с тех пор, как он из южной части современной Африки начал своё «победное» шествие по планете. **В.Э.:** Если я правильно помню, у человека с шимпанзе совпадают 99 % генов. Это с шимпанзе, а с питекантропом, наверное, не менее 99,9 %.

мягко) порицания¹⁶¹ (состоящий из двух соответствующих генераторов). Генератор наказания должен активно вмешиваться в общую картину протекающих в голове Долли процессов, если она сделала что-то «неправильно», и подавлять те процессы, которые вызвали эту деятельность. Задача генератора поощрения, напротив, стимулировать процессы, которые ответственны за «правильную» деятельность. Механизм поощрения и порицания операционной системы Доллии будет очень важной частью, можно даже сказать, «центральной осью», «вокруг которой всё вращается»; мы должны встроить в Доллию широкий спектр поощрений и порицаний, чтобы добиться от неё «правильного» поведения, но в данном параграфе мы обсуждаем вопросы самопрограммирования и должны к ним вернуться.¹⁶²

525. Для операционной системы Доллии нам понадобится аппарат, который сочетая небольшое количество уже существующих блоков, сможет сгенерировать для Доллии программы, которые необходимы ей в текущий момент. Следуя традиции предыдущих параграфов, назовём этот концептуальный блок «Генератором Программ». Необходимо предусмотреть возможность накапливать в памяти программы, созданные Генератором Программ, модифицировать их в зависимости от результатов их выполнения и использовать для генерации программ более высокого уровня.

526. Если мы теперь посмотрим на поведение Долли в некоторой ситуации, которая требует конкретной реакции, то опытному программисту очевиден основной принцип подготовки этой реакции: создать ряд вариантов программ реагирования, оценить возможные последствия их выполнения, а затем выполнить ту программу, для которой ожидаемый результат является наиболее приемлемым.¹⁶³

¹⁶¹ **СМ.** У автора «механизм наказаний и поощрений», — мне же всегда хотелось, чтобы для стимулирования деятельности людей в первую очередь использовалось поощрение (но реальная жизнь почему-то возможностей для наказания предоставляет гораздо больше ☺). Кроме того, я бы уточнил, что речь идёт о системе самопоощрения или самонаказания. **В.Э.:** Нет, в первую очередь это внешние наказания и поощрения. Когда младенец составляет себе программу, как достать конфету из вазочки, которую мама оставила на полке, и, выполняя эту программу, лезет на стул и падает вниз, то полученные ушибы и боль — это самое настоящее внешнее наказание ему за неправильно составленную программу, а если он конфету достаёт, то это — внешнее поощрение за правильную программу. А «генератор наказаний», работающий в его голове, это тот аппарат, который сигнализирует, что получен ушиб, и это «больно», то есть, «плохо» (идут сигналы в мозговые «центры ада»). «Генератор поощрений» в его голове — это сигналы в «центры рая» мозга в данном случае при получении сладкой еды. (Ведь в принципе в Доллии мы можем сделать наоборот: ударились — и заработал генератор поощрений! Но это будет самоубийственная стратегия Доллоса — ее операционной системы. В природе естественный отбор такие системы истребил).

¹⁶² **СМ.** В этом и двух предыдущих абзацах (522, 523) находятся предпосылки к ответам на мои вопросы, которые я сформулировал (и оставил открытыми), обсуждая материалы параграфа 29 книги В. Эгле L-ARTINT. Однако, если механизм «стартовых блоков», или стартового набора программ (отдалённый аналог BIOS в компьютерах) практически не может быть подвергнут изменению самой Доллией (так же как компьютер не может, в общем случае, подвергнуть изменению свой BIOS из программ, работающих под управлением его операционной системы, а человек практически не может изменить свои персональные особенности, доставшиеся ему вместе с ДНК), то механизм самопоощрения и самонаказания более пластичен в том плане, что его определения того, «что такое хорошо и что такое плохо», могут меняться в результате и под воздействием различных причин (слова воспитателей, пример других людей, новые знания, жизненный опыт, переоценка жизненных ценностей в результате каких-то событий во внешнем мире и так далее). Иначе Коперник (Николай Коперник (*Kopernik, Copernicus*) (1473–1543), польский астроном, создатель гелиоцентрической системы мира) никогда бы открыто не заявил о своём понимании солнечной системы в противопоставление господствующему мнению. Только возможность влиять на настройку «системы самопоощрения или самонаказания» дали ему возможность определить, что Солнце в центре солнечной системы и знание об этом — это «хорошо», а представление о Солнце, вращающемся вокруг Земли — это «плохо».

¹⁶³ **СМ.** Именно так я учил физику в школе: крутил текст параграфа школьного учебника в голове до тех пор, пока он не укладывался в некоторую внутреннюю (внутри меня) согласованную и непротиворечивую картинку, которая позволяла «осознано» решать задачи (кроме того, как только такая внутренняя картинка была сформирована, определения, данные в параграфе, становились понятными с точностью до каждого слова, как если бы кто-то протёр запотевшее стекло и за ним проявился четкий и ясный текст). Точно так же я переводил 528 абзац: крутил в голове возможные значения латышских слов, пока не получил некоторое предложение, имеющее смысл в рамках обсуждаемой темы, а потом записал его уже по правилам русской орфографии.

527. Следовательно, у нас есть необходимость еще в одном функциональном блоке в операционной системе Доллии: в блоке, который будет рассматривать созданные Генератором Программ программы с целью оценить возможные последствия их выполнения, и только после этого принимать решение о выполнении той или иной программы. Назовём этот концептуальный блок, скажем, «Анализатором Программ».

528. Давайте подумаем также о принципах работы Анализатора Программ. Понятно, что мы не можем разрешить ему действительно выполнить анализируемую программу во внешнем мире (так как последствия могут быть катастрофическими для Доллии). Поэтому Анализатор Программ, очевидно, должен работать с некоторым испытательным программным интерпретатором, который генерирует последствия выполнения анализируемой программы (точнее, моделирует выполнение программы в данной конкретной ситуации, строит некоторую итоговую картину, готовит информацию о возможных результатах) в некоторой внутренней компьютерной структуре данных, которые после этого также должны быть проанализированы.¹⁶⁴

529. Назовём такого рода структуры данных «потенциальными продуктами» (программ, изучаемых Анализатором Программ) (эти продукты [представляют собой]: «то, что было бы получено, если бы изучаемая программа была действительно запущена на выполнение»). Позже мы увидим, что потенциальные продукты разных мозговых программ являются главными объектами, которыми оперирует так называемое «абстрактное мышление». Но всё это начинается именно здесь, — возникает из необходимости предварительно, не выполняя собственно программы, оценить результаты тех программ, которые создаются Генератором Программ.

530. Таким образом, в связи с самопрограммированием Доллии, мы на данном уровне детализации можем различить три группы функций: 1) генерация программ, 2) получение потенциальных продуктов (потенциальных результатов) этих программ, 3) оценка этих результатов. Теперь уточним, что Генератор Программ у нас будет осуществлять первую группу функций, Анализатор Программ третью, а для выполнения второй группы функций определим отдельный функциональный блок, который назовём, скажем, «Проектором Программ». А все эти три подблока обозначим одним словом «программатор», который и реализует всё самопрограммирование Доллии.¹⁶⁵

§20. Интенсивность процессов

2011.06.27 15:55 понедельник

В приведенных выше в §§ 16–19 письмах содержался один четко сформулированный вопрос, не отозванный и не отложенный на потом; это вопрос из сноски, носящей сейчас номер [141](#), и относящийся к фразе «Нам понятно, что в голове Доллии параллельно будет происходить множество различных процессов. Каждый из них может работать с разной интенсивностью». Вопрос был таким:

В каком смысле используется термин «интенсивность»? В компьютере есть перераспределение ресурсов, повышение приоритета задач в интересах процессов или программ, которые требуют незамедлительного выполнения, особенно, если речь идёт о работе компьютерных систем в режиме реального времени. Но в компьютере нет ничего похожего на интенсивность работы того же землекопа (который вдруг начинает бросать землю быстрее ☺).

¹⁶⁴ **СМ.** Сохранить стиль автора в этом абзаце мне не удалось, но надеюсь, что его смысл я передал правильно. **В.Э.:** Да, в целом правильно; вот точный перевод абзаца: «Подумаем также и о принципе, по которому Анализатор Программ может работать. Ясно, что он не должен действительно выполнять анализируемую программу во внешнем мире (так как последствия могут быть катастрофическими для Доллии). Поэтому, очевидно, Анализатор Программ должен будет работать в качестве такого интерпретатора исследуемой программы, который создает последствия выполнения этой программы (точнее: модели этих последствий, их образы, информацию о них) в виде каких-то внутренних для компьютера структур данных, которые потом и подвергаются анализу». Этот принцип в Веданской теории называется по-русски «бокоанализом», и о нем была речь, например, также в {R-POTI-2 = МОИ [№42](#)} §9. Там, правда, это давалось уже в контексте математики, но сам принцип бокоанализа создан (естественным отбором) отнюдь не для математики. Ниже я дам принципиальную схему бокоанализа (Рис.3).

¹⁶⁵ **В.Э.:** Последние два абзаца, ввиду их исключительной важности, я подкорректировал; это уже практически перевод (выше тоже местами корректировал, где смысл не передавался правильно).

Результаты работы программ Доллии (и человека тоже) часто заключаются в отправке сигналов (мышцам, железам внутренней секреции, другим процессорам мозга). По данной «линии связи» сигналы могут не отправляться вовсе (интенсивность 0), могут отправляться редкие сигналы или в небольшом количестве, если по параллельным каналам (невысокая интенсивность), и может посылаться максимальное возможное или близкое к максимальному количество сигналов (высокая интенсивность).

От интенсивности посылаемых мозгом сигналов будет зависеть и интенсивность работы тех аппаратов организма, которые этим мозгом управляются. Пошлет Доллос шквал сигналов в надпочечники, и те выбросят в кровь пресловутый адреналин; будет посылать мышцам землекопа лихорадочные серии интенсивных сигналов – и землекоп будет «копать поглубже и бросать подальше». И под соответствующими шквалами сигналов сердце у него будет биться чаще, легкие расширяться и сокращаться с большей частотой и амплитудой...

Кроме интенсивности сигналов, посылаемых по «линиям связи», есть еще один показатель, характеризующий интенсивность мозговых процессов. Это т.н. «пороги». Обычно, если взаимодействуют два мозговых процесса (процессоры) A и B , то B реагирует не на отдельный сигнал от A , а только тогда, когда количество сигналов превысит некоторый «порог». Но этот порог – динамический, он может изменяться (опять же – под управлением Доллоса). В одном состоянии системы порог высокий, и B «не чешется», даже если A посылает ему сравнительно много сигналов. А в другом состоянии системы порог может быть намного снижен, и малейшего сигнала от A уже достаточно, чтобы B заработал.

Существенной частью эмогенератора является именно управление порогами между различными процессорами системы. А эмоциональные состояния, в свою очередь, характеризуются состоянием порогов: в каждом эмоциональном состоянии определенный набор порогов завышены, а другие пороги понижены. Хорошо известно, что если человек, например, и так зол, то уже любой мелочи достаточно, чтобы он «взорвался», а в «нормальном» состоянии он на эту мелочь вообще внимания не обратил бы. Но эти (отличающиеся) реакции определяются в первую очередь именно высотой порогов у разных процессоров мозга (во вторую очередь действуют и другие факторы).

Итак, мы выделили в Доллосе (и, соответственно, в человеке) функциональный блок, назвав его «эмоциатор», состоящий из двух подблоков под названиями «эмогенератор» и «эмоселектор».

Важно понимать, что эмоциатор не обязательно должен быть анатомически единым, то есть, находящимся в одном месте. Он может быть и распределенным – разбросанным по всей системе. Он именно функциональный блок, то есть: система может переходить из одного «эмоционального» состояния в другое – и всё, что осуществляет эти переходы, мы относим к «эмоциатору».

А отличаются «эмоциональные» состояния системы один от другого разными порогами у различных процессоров, разной частотой и количеством испускаемых процессорами сигналов.

Эмогенератор «создает эмоции», то есть, создает ту или иную комбинацию порогов у различных процессоров, устанавливает для них частоту испускаемых сигналов и их количество (интенсивность). А эмоселектор следит за окружающей ситуацией и деятельностью самой системы и решает, какая комбинация всех этих параметров должна быть сейчас установлена (то есть, какая эмоция должна быть в данный момент «включена»).

Ты говорил об увеличении тактовой частоты процессоров *Intel 386* и *486*. Но все эмоции могут быть реализованы «чисто программными» средствами. Пусть Доллия управляется процессорами $A, B, \dots, Z$. Определим для каждого из этих процессоров три константы¹⁶⁶:

a – порог: число импульсов в секунду, которые необходимы, чтобы процессор заработал;

b – частота посылаемых процессором пучков сигналов управляемому им устройству: сколько пучков в секунду;

c – объем пучка: сколько сигналов в пучке.

Поместим эти константы в общее поле с эмогенератором. Теперь эмогенератор, меняя константы a, b и c для разных процессоров $A, B, \dots, Z$, управляет всем эмоциональным состоянием системы (Доллии): у одних процессоров снижает пороги, увеличивает частоту пучков и объемы

¹⁶⁶ Они константы для процессоров $A, B, \dots, Z$, так как те их не меняют. Для всей системы в целом это – переменные.

пучков, у других, наоборот, уменьшает частоту и объёмы или так завышает пороги, что те процессоры оказываются вообще заблокированными.

В биологическом мозге, разумеется, нет цифровых констант; всё это осуществляется аналоговыми процессами: увеличивается или уменьшается пропускная способность какого-то нерва, какой-то связки нейронов и т.п. Но с точки зрения информатики – результат тот же самый.

Так что эмоции в компьютерах реализуются элементарно. Те люди, которые говорят, будто компьютеры не могут испытывать эмоции, просто не знают, что такое эмоции из себя представляют, и как надо строить большие и гибкие компьютерные системы. (А сколько я таких слов слышал за свою жизнь! Не сосчитать... Над такими словами надо просто смеяться, им и отвечать в общем-то не стоит).

В заключение я для наглядности дам принципиальную схему эмоциатора (Рис. 2).

§21. Краткие письма

no Valdis Egle valdis.egle@gmail.com
kam Sergey <sevm@rambler.ru>
datums 2011. gada 26. jūnijs 12:17
temats Re: R-POTI-3
piegādātājs gmail.com

Здравствуй, Сергей!

Опять накопилось много твоих писем без моего ответа. Как ты, наверное, читал, новые архивы Векордии выставляются на сайт <http://vekordija.narod.ru/INDEX.HTM> 4 раза в год в начале каждого квартала. Сейчас приближается одна из этих дат – 1 июля – и идет интенсивная подготовка очередного выпуска. Когда архив появится, ты его скачай: там будут и новые книги, и много переизданных старых с исправленным множеством ошибок, с новым форматом PDF, у которого нормальное оглавление (раньше мне такое не удавалось получить). И тогда будет выставлена в Интернет также и книга POTI-3 с твоими текстами и моей реакцией на них.

Так что – до 1 июля!

В.Э.

no Sergey sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 29. jūnijs 00:20
temats RE: R-POTI-3
piegādātājs rambler.ru

Валдис, добрый день!

О квартальном обновлении книг, статей и других материалов Векордии я, конечно же, помню, и рад, что до предстоящего выпуска мне удалось немного продвинуться вперёд на пути понимания/изучения идей Веданской теории в отношении Искусственного Интеллекта, и, наверное, теперь уже и в более широком понимании положений этой теории.

Несмотря на нехватку времени на все мои проекты (наш проект переписки о теории Интеллекта – один из них), которые мне хотелось бы (а иногда и не хотелось бы ☺) делать, я всё-таки не удержался и заглянул в твою книгу R-DVESA [= МОИ №50]. Пока я только очень бегло пробежался по её страницам, но уже вижу, что она в некотором смысле является подготовительной книгой для понимания теории Интеллекта, а для меня ещё и сборником ответов к тем занимательным задачкам, которые мне подкидывает прочтение L-ARTINT ☺).¹⁶⁷ Так, например, читая об эмоциональных процессах человека, я вспомнил о мужестве и стойкости греческих наёмников в битве при Иссе (333 до н.э.), а в книге нашёл сравнение поведения «труса» и «героя» (абзац 1075, Valdis Egle. «R-DVESA». 2010) [= МОИ №50, стр.61].

СМ

P.S. Прилагаю файл с печатками.

¹⁶⁷ В.Э.: В твоих текстах вместо закрывающей скобки очень часто стоит :), что по общепринятым соглашениям означает «смайлик». Но что-то я засомневался: может быть у твоего Word-а установлена такая автозамена, всегда меняющая) на :)? Или какая-то другая неполадка (ведь если это ☺, то закрывающей скобки нет, а если это скобка, то зачем «:»)? В общем, я не знаю, просто скобка это, или ☺?

2011.07.01 03:08 ночь на пятницу

В.Э.: Спасибо за файл опечаток.

Вот, пришло 1 июля, но я, к сожалению, не успел дописать всё, что хотел здесь сказать. (Уж больно канительная это работа – переиздавать книги; за день можно сделать в среднем 3–4 тома, а у меня их 140). Это и затянуло меня как в «асфальтовую топь» Брукса.¹⁶⁸

Поэтому я временно удалил отсюда незаконченные части (они последуют через некоторое время; два параграфа и три схемы за мной) и помещаю в Интернет эту книгу такой, какая она есть до этого места.

§22. Эмоциатор

2011.07.03 00:49 ночь на воскресенье

Итак, я продолжу то, что не успел написать к 1 июля. В первую очередь: Рис. 2 – Принципиальная схема Эмоциатора:

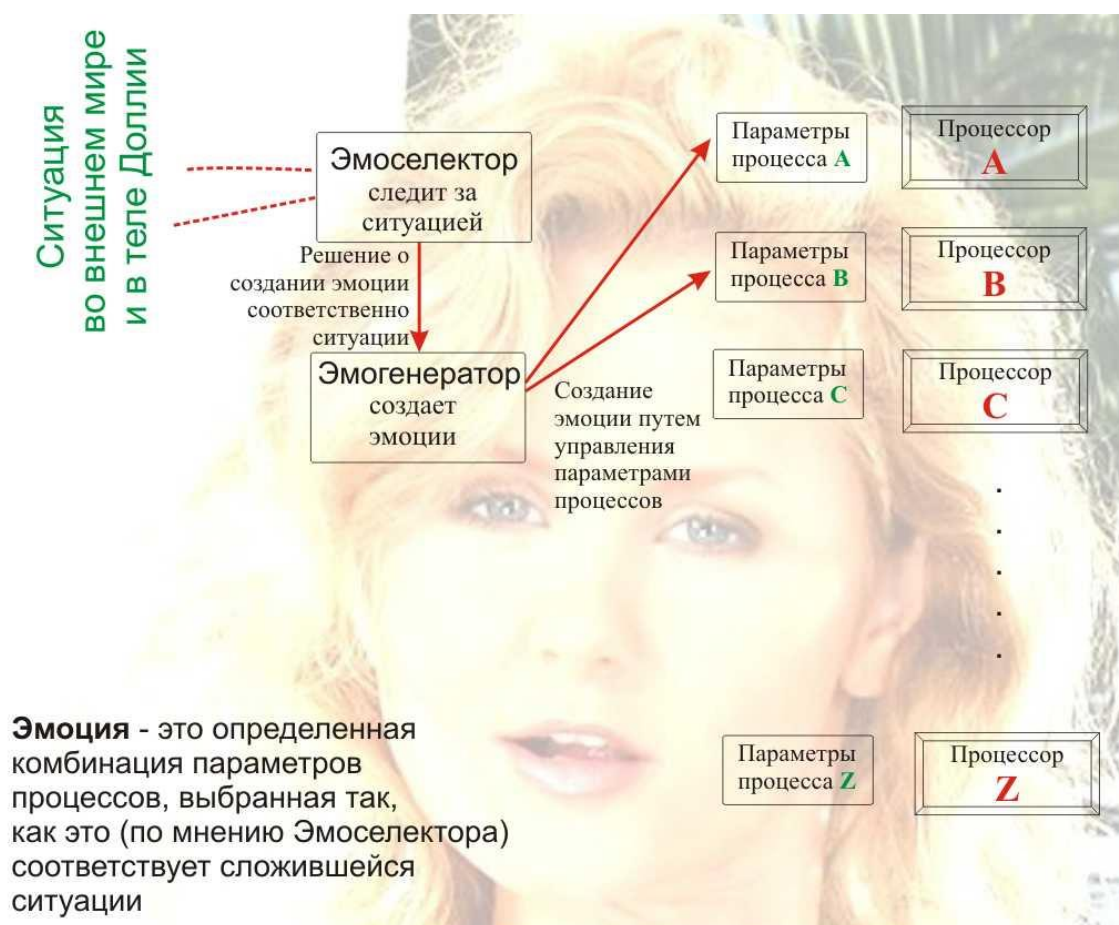


Рис.2. Принципиальное устройство Эмоциатора в голове Доллии

Многое по этой схеме уже говорилось выше. Теперь сделаем просто некоторый обобщающий повтор.

Биологические системы вообще и человек в частности – это системы гибкие. Их гибкость заключается в первую очередь в том, что они (и их части) могут работать с различной интенсивностью, переходя от режимов почти-покоя до развития максимальной возможной мощности, причем градация этих состояний не носит скачкообразный характер, а меняется

¹⁶⁸ Брукс, руководитель разработки OS/360 в фирме IBM, потом написал книгу «Мифический человек-месяц», которую начал с описания того, как динозавр тонет в «асфальтовой топи» (так это перевели на русский; вообще речь идет, конечно, об илистом болоте).

плавно. Такая гибкость позволяет экономить ресурсы и не изнашивать оборудование в спокойных ситуациях, но выдавать всю возможную мощь в ситуациях экстремальных.

Но такие переходы требуют определенных изменений, переключений в системе. Эти переключения (означающие подготовку системы к деятельности того или иного рода в том или ином режиме) и называются эмоциями.

Это очень общее определение эмоций: в нем не говорится, о какой именно эмоции (из традиционно выделяемых) идет речь. Пол Экман выделил пять «базовых эмоций»: «страх, печаль, счастье, злость, и отвращение»? Хорошо, пусть будет так, но на самом деле состояния системы гораздо разнообразнее: параметры у разных процессоров могут по-всякому комбинироваться, и к тому же меняются плавно, а не скачкообразно. Поэтому люди, пытающиеся описать эти состояния в категориях «пяти базовых эмоций», вынуждены прибегать к таким рассказам, как, например: «Его охватило сильное чувство страха вперемешку с отвращением и злостью...». Так что любое деление состояний системы на отдельные эмоции всегда будет условным.

Данное выше определение эмоций очень «конструктивно»: это взгляд конструктора, который сам строит эмоциональные системы, а не просто их наблюдает. Для сравнения вспомним хотя бы подобранные тобою же выше слова из Википедии:

В каждой эмоции человека можно (достаточно условно) выделить инстинктивную часть (достались нам от наших далёких предков) и некоторую часть, сформировавшуюся под воздействием общественной жизни человека: «...Эмоции эволюционно развились из простейших врождённых эмоциональных процессов, сводящихся к органическим, двигательным и секреторным изменениям, до значительно более сложных, утративших инстинктивную основу процессов, имеющих отчётливую привязку к ситуации в целом, то есть выражающих личное оценочное отношение к имеющимся или возможным ситуациям, к своему участию в них... Выражение эмоций имеет черты социально формирующегося, изменяющегося с течением истории языка, что можно видеть из различных этнографических описаний.»

«Инстинктивный», «врожденный»... Всё это не разговор конструктора. Хотя в общих чертах это правильно, но это не дает ключа к построению самого всего этого. (А ключ там – в той Принципиальной схеме Рисунка 2).

И когда мы Основной принцип построения эмоций в системе уже понимаем, тогда мы можем добавить: «Да, процессы А, В, С, ... были уже у ящериц, а вот процессы W, X, Y, Z – это специфически человеческие, они определяют *«личное оценочное отношение к имеющимся или возможным ситуациям»* и они *«социально формирующиеся»* и т.д. Тогда мы окажемся сказавшими то же самое, что в Википедии (и в учебниках психологии), но только на гораздо более глубоком уровне понимания. (Наш уровень понимания следует считать более глубоким потому, что мы знаем, как это встроить в куклу Доллию, а профессора психологии не знают).

Пожалуй, пока хватит об Эмоциаторе. В заключение только приведу еще раз соответствие между традиционно выделенными эмоциями и той основной стратегией Доллоса, приготовление к осуществлению которой данная эмоция означает. (И расположу эти пять «базовых эмоций» в таком порядке, в каком они должны стоять по филогенезу; в оригинальной цитате они перечислены бессистемно):

- страх – стратегия ухода от опасности, как правило, бегства;
- злость – стратегия нападения (на врага, не на добычу);
- отвращение – стратегия отмежевания, ухода;
- счастье – стратегия безопасности при достигнутых целях;
- печаль – стратегия бездействия при недостижимых целях.

§23. Программатор

2011.07.03 13:54 воскресенье

На Рис.3 приведена принципиальная схема устройства Программатора (аппарата самопрограммирования) в голове Доллии (и в человеческой голове тоже).

В этом функциональном блоке мы выделили три подблока: Генератор Программ, Проектор Программ и Анализатор Программ.

Генератор Программ создает новую программу для Доллии. При этом он, разумеется, пользуется какими-то исходными данными (в которых можно выделить две главные части: 1)

информация, определяющая, какая, собственно, программа должна быть создана; и 2) библиотека уже ранее созданных программ, которые могут быть включены во вновь создаваемую как её подпрограммы).

Вновь созданная программа (в схеме: Программа А) существует в мозге Доллии как некоторая структура данных. Если её запустить на выполнение, то она отработает как программа, но если не запускать, то она – данные, которые могут стать исходными для других программ.

И вот, Проектор Программ и есть такая программа, которая анализирует вновь созданную программу А «сбоку», как данные, на самом деле её не выполняя, а лишь интерпретируя, эмулируя, симулируя её (обычно частичное, не во всех деталях) выполнение. Такой процесс в Веданской теории называется бокоанализом (от той ассоциации, что Проектор и эмулируемая им программа находятся рядом – они стоят «бок о бок»).

Цель эмуляции: получить возможные результаты выполнения Программы А, но получить не в реальном мире, а лишь в мозге Доллии – как некоторую структуру данных, построенную Проектором. Такая структура называется в Веданской теории потенциальным продуктом (программы А). Если бы Программа А действительно отработала, то она создала бы реальный продукт (или фактический продукт), – но уже в реальном мире. Этот фактический продукт Доллия могла бы увидеть. А «увидеть» – это означает: построить у себя в голове некоторую модель увиденного объекта (такая модель в Веданской теории называется номиналией).

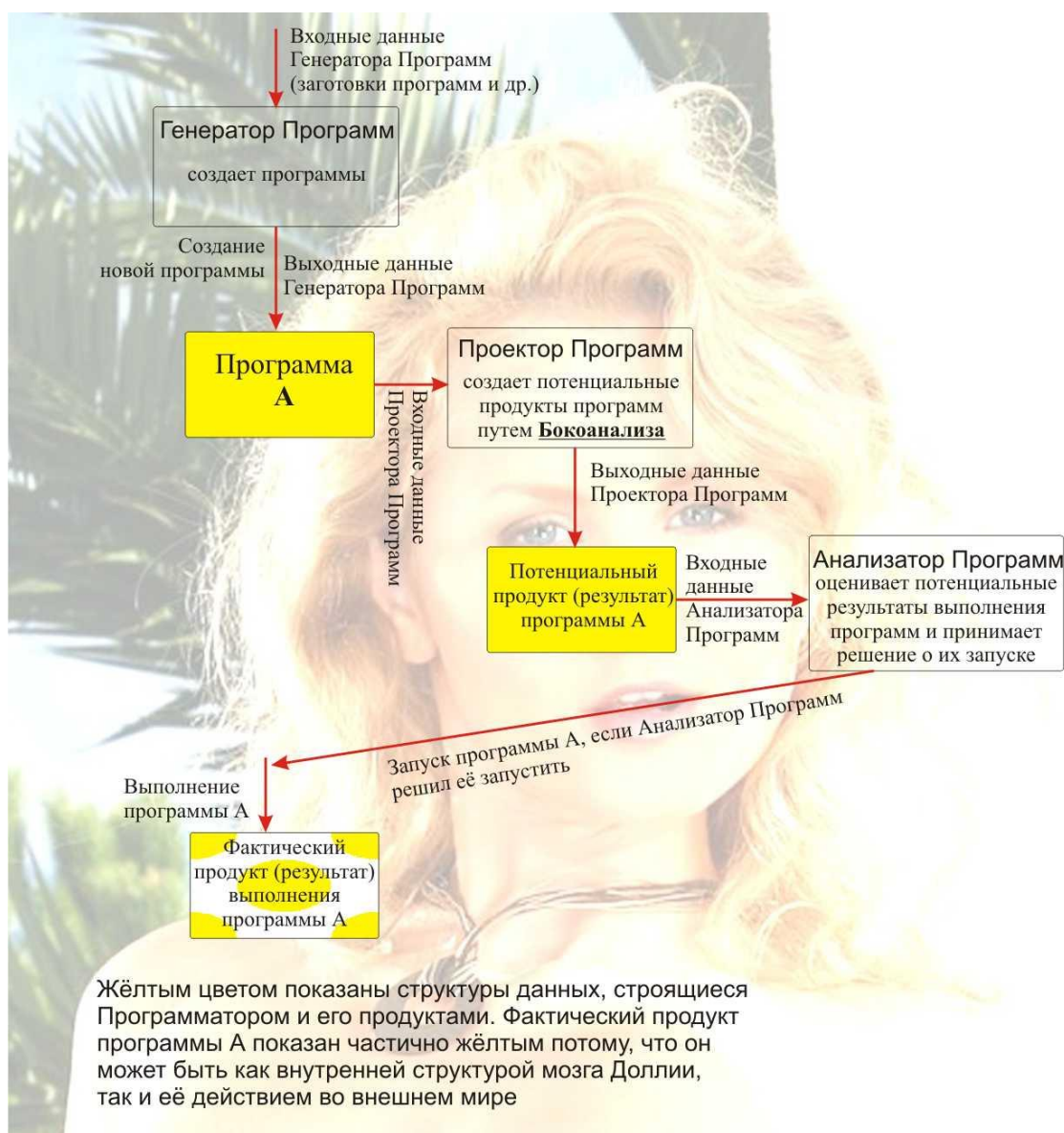


Рис.3. Принципиальное устройство Программатора в голове Доллии

Так вот: потенциальный продукт (построенный Проектором) и номиналия фактического продукта (возникшая, когда Доллия на самом деле видит фактический продукт) – это объекты (структуры данных) очень похожие: построение Проектором потенциального продукта (Программы А) означает, что Доллия как будто «предвидит» (предварительно видит) результаты выполнения Программы А.

Теперь этот потенциальный продукт (как структура данных) подвергается изучению со стороны Анализатора Программ. Анализатор изучает потенциальный продукт (Программы А) и решает: устраивает он Доллию или нет. Если устраивает, значит Программа А запускается на (уже реальное) выполнение. А если не устраивает, то Программа А так и остается лишь сгенерированной, но невыполненной (в бытовой речи скажут: «было намерение, но...»), а через некоторое время разрушается («забывается»).

По-моему, совершенно ясно, зачем Естественный отбор встроил такой аппарат в живых системах. Опасно сразу выполнять любую программу (намерение), которая «взбредет в голову» (т.е. будет сгенерирована Генератором Программ). Полезно бывает сначала «подумать», оценить последствия, прежде чем кидаться осуществлять любую задумку. А это «подумать» как раз и заключается в отработке Проектора и Анализатора Программ.

Создано это было в процессе эволюции с очевидно полезной целью, но как «боковой эффект» это порождает «платоновский мир идей» «абстрактных понятий». Ведь после отработки Проектора Программ потенциальный продукт (Программы А) УЖЕ (!) существует в голове Доллии как некоторый реальный объект, хотя сама Программа А еще не работала и, может быть, вообще никогда не будет работать.

Этот «платоновский мир идей», создаваемый Проектором Программ, помимо многих других вещей, содержит и весь сонм математических понятий, начиная уже с простого натурального числа. (Это очень просто и очевидно, но профессора и академики от математики уже, вот, 33 года не в состоянии это понять и продолжают нести архаичную и бессильную чепуху. Но это уже другая тема – хотя и основная для ПОТИ).

§24. Дерево программы

Как уже говорилось в твоём Конспекте, ни Генератор Программ, ни Проектор Программ не могли бы успешно работать, если Программа А была бы сплошным массивом простых инструкций (элементарных «машинных команд»). Их работа возможна только тогда, если программа структурирована.

Структура программы кодируется в виде дерева, состоящего из блоков (узлов) различных уровней. Вообще вся память человека (и других живых существ тоже) построена на древовидных структурах – это стандартный «почерк Природы». Этому есть прямые экспериментальные подтверждения. Например, давно установлено, что время поиска элемента среди n элементов в человеческой памяти возрастает не пропорционально n , а пропорционально логарифму от n . А такое возможно только в древовидных структурах.

Ну, и самопрограммирование тоже не является исключением. Программа, создаваемая Генератором Программ, строится как дерево.

Чтобы не говорить абстрактно-непонятно, рассмотрим это на конкретном примере. Предположим, наша Доллия решила ограбить банк. (Это вообще-то обычно мужское занятие, хотя богатая криминальная история человечества знает и женщин – грабительниц банков. Допустим, мы справились с задачей создания самопрограммирующейся операционной системы Доллии настолько успешно, что она выросла у нас куклой чрезвычайно самостоятельной, энергичной и бой-



Банк, который Доллия решила ограбить

кой, ни в чем не отстающей от лучших человеческих образцов).

Итак, Доллия решила ограбить банк (допустим, она живет в Риге). Это её намерение – уже первый, самый верхний узел в строящейся программе ограбления (см. Рис.4).

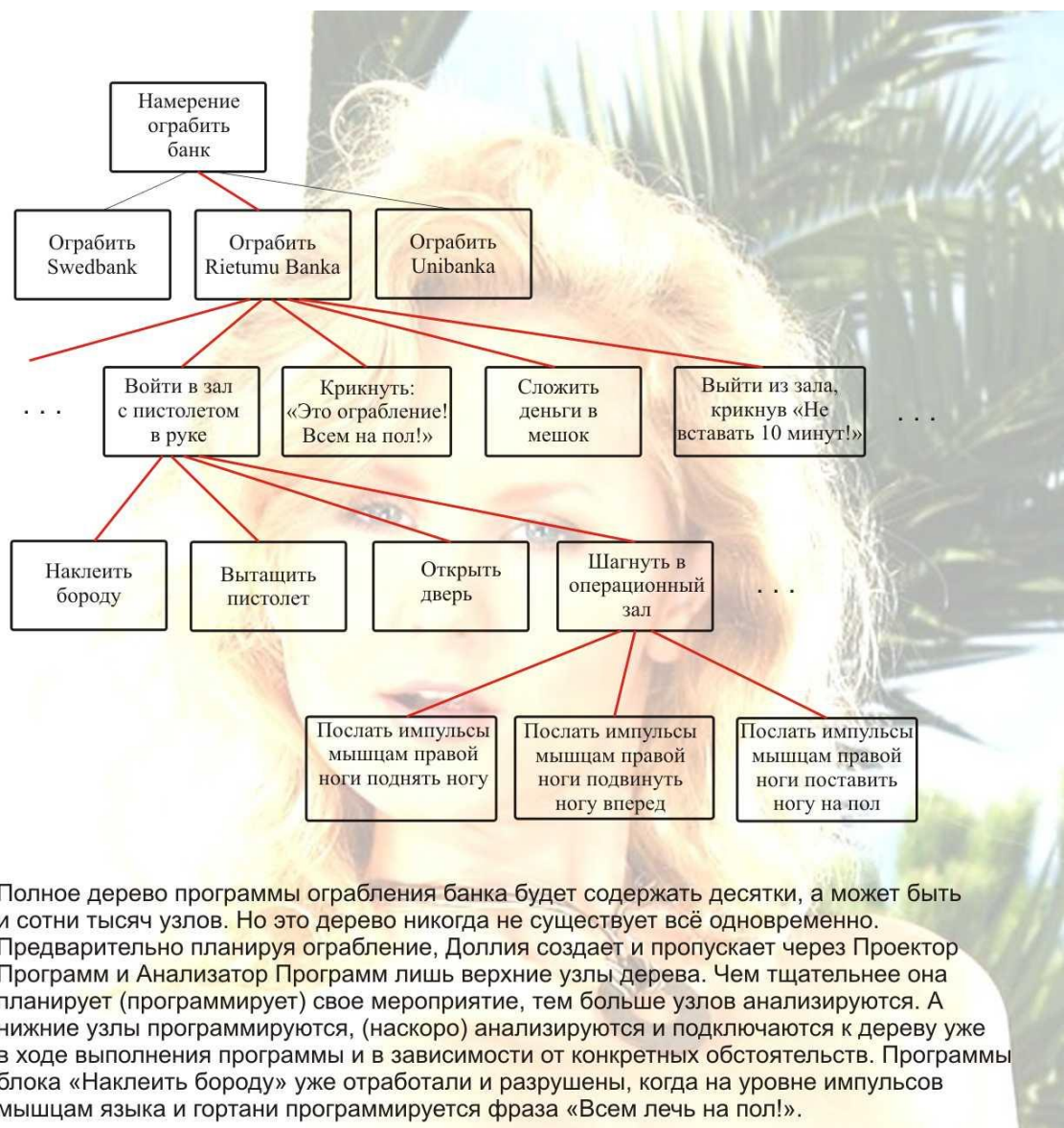


Рис.4. Часть дерева программы ограбления банка «Rietumu banka» в голове Доллии

Как правило, уже этот первый узел будет подвергнут обработке Проектором Программ и Анализатором Программ. Проектор построит картину, в которой Доллия имеет очень очень много денег и может купить очень очень очень много конфет и губных карандашей. Правда, Проектор построит и картину, в которой Доллию сажают в тюрьму. Но Анализатор, оценив обе эти картины, решает, что вторая картина, во-первых, маловероятна (Доллия весьма высокого мнения о своих умственных способностях), и, во-вторых, если даже её и поймают, то там откроется широчайшее поле для деятельности адвокатов: ведь она же кукла! Как можно судить куклу – такое не предусмотрено никакими Уголовными и Процессуальными законами; она поставит в тупик всю юридическую систему Латвийской Республики!

Словом, Анализатор признает программу ограбления пригодной для выполнения, и само-программирование продолжается. Теперь Генератор Программ Доллии создает блоки второго уровня. Какой банк грабить? Допустим, Генератор построил три узла: грабить *Swedbank*, грабить *Rietumu Banka*, грабить *Unibank*. В этот момент в программе ограбления на втором уровне три ветки. Опять прогоняются Проектор и Анализатор, и Доллия решает грабить *Rietumu Banka*: он не в центре города, а на отшибе, на Весетас 7... Теперь две остальные ветки отключаются от

дерева (но они остаются в запасе в памяти: в случае чего могут быть опять подключены вместо основной ветки (отмеченной в схеме красной линией)).

Далее идет программирование третьего уровня: как подъехать к банку, как наклеить бороду (чтобы все думали, что грабит не кукла, а человек мужского пола), как войти в операционный зал, как там действовать, что кричать, как уехать...

Самопрограммирование всегда динамический процесс: одни ветки разрабатываются, прогоняются через Проектор и Анализатор и отбрасываются, другие остаются подключенными к дереву. Подключается масса давно заготовленных программ: как водить машину, как ходить по земле, как заряжать пистолет, как приклеивать бороду, как произносить фразы «Это ограбление! Всем лечь на пол!» и «Не вставать десять минут!»...

Полностью вся программа не составляется до конца перед началом её выполнения. Многие низшие ветки допрограммируются по ходу действия: ведь невозможно предусмотреть, какие препятствия придется обойти, какие действия предпримет персонал банка... Уже исполненные нижние ветки дерева разрушаются (их уже нет в памяти), а новые создаются на ходу...

Ну вот, так выглядит работа Программатора согласно Веданской теории.

§25. Письмо Сергея Марьясова от 18 июля 2011 года (еще раз Эмоциатор)

no Sergey sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 18. jūlijs 16:23
temats R-POTI-3
piegādātājs rambler.ru

Валдис, добрый день!
Текст письма в приложенном файле.
Всего наилучшего,
СМ

Посмотрел твои поправки в текстах конспектов: перевод некоторых абзацев давался мне довольно тяжело, и увидеть твой (авторский) перевод было очень интересно (и полезно для дальнейшего конспектирования ☺).

Попробую кратко ответить на некоторые твои примечания и добавить ещё несколько слов по теме Эмоциатора.

Слово «субблок» можно поправить на «подблок», как любую другую опечатку.

Размышление об использовании слов других языков в русском языке привлекло моё внимание к слову «программатура», которое является русской записью латышского слова¹⁶⁹, означающего «программное обеспечение». Но в отличие от заимствованного слова «софтвер», которое довольно часто используется в русскоязычном интернете, «программатура» используется только в Латвии и латвийском интернете. То же самое можно сказать о распространении слов «аппаратура» и «хардвер». Может быть добавить это пояснение к слову «программатура» на странице 12, где оно встречается первый раз?

Со смайликом у меня действительно образовалась некоторая путаница. В тексте письма он всегда оставался как «:)), и если он совпадал с закрывающей скобкой, то я вторую скобку опускал, чтобы не писать «:)))». А при подготовке файла текстовый редактор почти всегда

¹⁶⁹ В.Э.: Ну, слово-то не латышское, а состоящее сплошь из «иностранных» компонентов. Но оно принято латвийской Комиссией по терминологии (есть такая при Академии наук; официально работает с начала 1920-х годов, а неофициально действовала еще в царское время при Рижском латышском обществе) – принято в качестве «официального» (т.е. рекомендованного в латышском языке) обозначения того, что по-английски называется «software». В общем случае я не считаю для себя обязательными рекомендации этой Комиссии даже когда пишу по-латышски, но если они приняли термин удачный, то почему его не использовать? – и если он «международный» (по образующим корням и суффиксам), то почему его не использовать и по-русски? Термин «аппаратура» в русском языке есть, и в одном предложении с ним также и термин «программатура» будет понятен читателю даже без специальных пояснений. (Тандем «хардвер–софтвер» звучит, в общем-то, уродливо, а тандем «аппаратура–программатура» – можно сказать, просто поэзия! ☺).

автоматически заменял двоеточие со скобкой на «весёлого колобка»¹⁷⁰. Поэтому в дальнейшем в текстах писем буду ставить «:))», в надежде на то, что при переносе в R-PTI-3 ты сделаешь преобразование в «☺»). А в текстовом редакторе буду следить, чтобы смайлик при необходимости был корректно закрыт скобкой.

Вопрос об определении набора стартовых блоков и их содержании (прим. 156) заново в список вопросов, которые мы могли бы обсудить после рассмотрения блок-схемы «Модель психической деятельности человека в Веданской теории» (L-ARTINT, с.112).

В примечании [157](#), корректируя мой тезис о «системе самопоощрения или самонаказания», ты написал, что «...в первую очередь это внешние наказания и поощрения...». Понятно, что источник наказания и поощрения всегда лежит вне Доллии, но мне хотелось подчеркнуть субъективность восприятия реального внешнего мира (и его ответной реакции на действие или бездействие человека) конкретным человеком. В похожих ситуациях, или даже в одной и той же ситуации, разные люди построят отличающиеся номиналии. Та же боль от падения в сочетании с другими обстоятельствами, например, прибежали испуганные родители, потащили ребёнка к врачу и т.д., может привести к тому, что у человека навсегда сохранится боязнь высоты. А в другом случае, когда папа или мама, игнорируя испуг и боль ребёнка, успокоили его и помогли ему добраться до заветной вазочки (и может быть помогли не один раз), а затем позволили ему проделать это путешествие самостоятельно, из него получится если не монтажник-высотник, то заядлый любитель побродить по горам ☺. Даже само восприятие боли у разных людей несколько отличается, и есть люди со сниженным болевым порогом. Но согласен, что с названием «система самопоощрения или самонаказания» я в данном случае промахнулся.

Валдис, отвечая на мой вопрос об интенсивности процессов, ты для описания работы Эмоциатора использовал три константы: *a*, *b*, *c*, которые отвечают за пороговое срабатывание процессора, частоту посылаемых процессором пучков сигналов, и объем пучка, соответственно. Мог бы ты указать на принципиальной схеме Эмоциатора (рис. 2), где действуют сигналы, определяемые этими константами?

Твой тезис о том, что : «... процессы *A*, *B*, *C*, ... были уже у ящериц, а вот процессы *W*, *X*, *Y*, *Z* – это специфически человеческие, они определяют «личное оценочное отношение к имеющимся или возможным ситуациям» и они «социально формирующиеся» и т.д.», настолько увлёк меня, что я начал размышлять о том, откуда берутся сигналы (информация) на входе Эмоциатора и Эмогенератора, как можно реализовать в Эмоциаторе эмоциональные процессы разной продолжительности (то, что в Википедии условно поделили на аффекты, чувства и настроения), как реализуется обратная связь и т.д.. Но, также как и ты, решил оставить эту детализацию «на потом».

Хочется ещё раз коснуться пяти «базовых эмоций». Если со страхом и злостью более или менее понятно, то с остальными тремя эмоциями не всё так просто (по-крайней мере для меня). Работая с твоей статьёй об эмоциях, я сформулировал для себя следующее определение счастья: счастье – это стратегия восстановления¹⁷¹, отдыха, расслабления (и, наверное, отчасти поэтому, беспечности), и, учитывая программный подход Веданской теории, формирование новых целей (старые-то в основном выполнены!). То есть, бездействие в этом состоянии только внешнее, а внутри человека кипит восстановительная энергия и создаются новые программы.¹⁷²

¹⁷⁰ В.Э.: Эта замена зависит от режимов, которые установлены в твоём Word-е, и это может быть по-разному на разных компьютерах. Чтобы текстовый редактор это менял, нужно: 1) чтобы была разрешена автозамена вообще; и 2) чтобы в таблице замен была пара «:» ☺». Я обычно, когда запускаю новый компьютер, разрешаю автозамену, но из таблицы замен удаляю все те глупости, которые там записаны при инсталляции и которые тайно портят мне текст, и оставляю только те замены, которые мне действительно нужны.

¹⁷¹ СМ. В первую очередь такой вывод я делаю на основании теории М. Норбекова (М. Норбеков. *Опыт Дурака, или ключ к прозрению как избавиться от очков*. Москва. 2008 г.).

¹⁷² В.Э.: Ну разумеется!, человеческий мозг, как и компьютер, вообще не может абсолютно ничего не делать (разве что в коме). Допустим, гимназист счастлив от того, что девушка (положительно) ответила на его любовное письмо. И вот, теперь его мозг генерирует всякие программы (бессмысленных) прыжков, восклицаний «Yes!!!» и подобную ерунду (причем одному его процессору еще приходится специально, «сознательно» тормозить остальные, у которых Эмоциатором установлены очень низкие пороги и широкий выход: «счастье льется через край», и гимназисту «хотелось бы еще что-нибудь такое натворить»). Имеется в виду бездействие в области целеустремленных, осмысленных предприятий. Но через некоторое время «уровень счастья» спадает, и он начинает действовать уже более разумно.

Печаль – стратегия выхода из состояния неудач, потерь, когда либо цели не достигнуты, либо стало ясно, что достичь их невозможно, либо, как и в случае недостижимых целей, во внутреннем мире человека, а точнее, согласно определениям Веданской теории, в системе его потенциальных продуктов рушится какая-то важная связка, какой-то важный для всей системы в целом элемент (например, крушение жизне-определяющей идеи или смерть очень близкого человека). Бездействие в этой стратегии тоже только внешнее. В области неосознанных программ происходит перестройка модели внешнего мира (или её важных компонентов), а на физиологическом уровне включаются механизмы забывания. Но у человека, в отличие от животных, физиологические механизмы забывания «проигрывают» мозговым программам, а перестройка мозговых программ происходит очень медленно. Согласно К. Шереметьеву человек не может забыть однажды «установленные» в его подсознании программы.¹⁷³ Изменить поведение человека можно только «снабдив» его новыми программами, которые должны заменить действие старых «вредных» в новых условиях (или обстоятельствах) программ, «повысив» их приоритет в подсознании человека. Термины «установить», «снабдить», «повысить» я беру в кавычки, потому что речь идёт о подсознании, или неосознанных программах, к которым добраться можно только через сознание, или осознанные программы согласно Веданской теории, а связь между ними у человека весьма и весьма непростая.

§26. Поток сигналов

В.Э.:

2011.07.19 15:50 вторник

Ты задал вопрос:

Валдис, отвечая на мой вопрос об интенсивности процессов, ты для описания работы Эмоциатора использовал три константы: a , b , c , которые отвечают за пороговое срабатывание процессора, частоту посылаемых процессором пучков сигналов, и объем пучка, соответственно. Мог бы ты указать на принципиальной схеме Эмоциатора (рис. 2), где действуют сигналы, определяемые этими константами?

Ну, если именно на схеме Рисунка 2, то так:



Рис.5. Один из потоков сигналов и его параметры

¹⁷³ СМ. К. Шереметьев для описания данного механизма использует термин «шаблон» (К. Шереметьев «Полноприводной мозг. Как управлять подсознанием», 2007 г.).

Здесь синими стрелками показан один из потоков сигналов, управляемый Эмогенератором. Предположим, что процессор A у нас такой, что сам он мышцам (и железам внутренней секреции) сигналы не посылает, а только другим процессорам (например, процессору B). Константа a_A – это порог процессора A ; константы b_A и c_A – его выходные характеристики; константа a_B – это порог процессора B ; константы b_B и c_B – его выходные характеристики. Процессор B уже управляет непосредственно (какими-то конкретными) мышцами, посылая им сигналы.

Таким образом, потоки сигналов – это между процессорами внутри мозга и от процессоров к мышцам и к железам внутренней секреции.

Разумеется, всякая схема есть упрощение. Мы фактически рассматриваем только элементы картин, а не законченные картины (которые, разумеется, почти всегда весьма сложны).

§27. От вещи к слову!

Авторов, на работы которых ты ссылаешься (Норбекова и Шереметьева), я, к сожалению не читал. В принципе мне хотелось бы прочитать, включить в Векордию и разобрать их, но это не может быть сделано быстро и не может быть условием продолжения этой Переписки. Поэтому пока приходится продолжать, ориентируясь только по твоим словам, без (с моей стороны) точных ссылок на этих авторов.

Обозревая два последних абзаца твоего письма, мне хочется подробнее вернуться к одному вопросу, о котором уже была речь в переписке ПОТИ. 15 марта 2011 г. в 18:24 ты писал: «С большим интересом прочитал книгу R-ПОТИ-1...» В той книге, в свою очередь, содержалась моя полемика с Грачёвым о том, как идти: «от слова к вещи» или «от вещи к слову» (см. мой пост 31.05.2010 15:33 и примечание к нему в {ПОТИ-1 = МОИ №41} §39). Там имелась также ссылка на латышский текст {VITA2.748}¹⁷⁴, а здесь я этот текст дам в переводе (он был адресован Майе Куле, профессору философии в Латвийском Университете, действительному члену Латвийской Академии наук и – одному из моих врагов):

.747. Инфантильных представлений в человеческом мышлении очень много; мы с ними уже встречались в этих книгах (напр. {199}) и будем встречаться и впредь. Но сейчас нас интересуют инфантильные представления одного специфического вида, а именно – распространенное мнение, что слова якобы имеют какой-то «объективный смысл».

.748. Когда маленький человечек [ребенок] учится говорить, он спрашивает у мамы или папы, что означает слово «пшеница», «самолет» и т.д.; здесь для него слово языка и обозначаемые им вещи выступают (на уровне его понимания) как одинаково объективные, и связь между ними как объективно данная. Алгоритм работы его мозга таков: «Вот, существует слово; надо выяснить, узнать, какой это слово имеет смысл».

.749. На самом деле ни одно слово никакого объективного значения не имеет: любую вещь можно назвать любым словом и (что то же самое), любым словом можно обозначить всё, что угодно. Связь между словом и его «значением» представляет собой (в лучшем случае) только соглашение между людьми (а часто нет даже и соглашения, а имеется одно только произвольное решение автора).

.750. То, что маленький человечек ощущает как объективную связь между словом и его значением, на самом деле представляет собой это соглашение, договоренность между людьми – только заключенное без участия маленького человечка и единственно поэтому ему как бы внешним образом, объективно данное.

.751. Подрастая маленький человечек в нормальном случае начинает понимать, что связь между словом и вещью весьма условна. Однако следы от его первоначального, детского восприятия часто (даже удивительно часто!) остаются в его мышлении. Ему всё еще кажется – если не всегда, то по крайней мере иногда, – что, вот, «имеется слово, понятие», и надо «выяснить, что же оно означает, что это такое». Термин, понятие у него получается первичным (по крайней мере в том смысле, что именно с него начинаются его поиски, его действия в этой области – а не начинаются «с другого конца», с Внешнего мира).

.752. Я мог бы собрать сотни примеров (если бы старался их регистрировать), где даже в работах вполне серьезных ученых имеются явные следы таких инфантильных представлений. Я видел даже диссертации и научные работы на темы «анализ такого-то понятия»; языковеды постоянно ищут «правильное значение слова» и т.д.

.753. В данной лекции я об этом явлении упоминаю потому, что и о слове «философия» те, у кого в мозге работают эти инфантильные представления, обязательно начнут спорить. Они скажут,

¹⁷⁴ <http://vekordija.narod.ru/L-VITA2.PDF> , <https://yadi.sk/i/dvitktHr3KEMxh>.

что данное мною в предыдущем параграфе определение философии «не верно», что следует «искать действительный предмет философии» (как будто слово первично, а предмет вторичен), что надо «исследовать границы философии» (как будто мы не можем проводить эти границы как желаем)...

.754. В моем определении философии ход мысли был такой: есть Внешний мир – есть его модели в головах – модели группируются – одна из таких групп, это Философия. Это был путь от Внешнего мира к понятию. Но они захотят, чтобы путь был противоположным, от понятия к Внешнему миру: есть Философия – эта наука существует и имеет какой-то предмет – и теперь нужно искать, каков же этот предмет и где границы Философии?

.755. Вот, в требовании такого, обратного, пути и будут проявляться их инфантильные представления.

.756. Следы инфантильных представлений хорошо видны и, например, в книге Майи Куле и Рихарда Кулиса «Философия»¹⁷⁵. Там уже на 7-ой странице написано «Предмет философии определяется кругом вопросов, которыми она интересуется. Следует принимать во внимание, что вопросы в течении веков существенно изменялись».

.757. Итак, для авторов сперва существует Философия, потом она о чем-то интересуется, и вот мы все вместе будем выяснять, о чем же, собственно, она в свое время интересовалась...

Я здесь привел всё это для того, чтобы еще раз подчеркнуть то, что говорю всем: зрелый ход мыслей идет не от слова к вещи, а от вещи к слову.

Сперва существуют какие-то состояния мозга (эмоциональные). Потом мы можем их назвать «счастьем», «печалью», «отвращением» – а можем и не называть. И границы этих понятий можем проводить, как нам вздумается.

Нельзя сначала полагать, что существуют «счастье», «печаль», «отвращение» – и потом думать и гадать: а что же это такое? Это детский подход – и, как правило, НЕ результативный. На самом деле эти понятия расплывчаты, под этими словами намешано много всякого и разного, и оттого и возникает путаница и непонимание.

А надо действовать «с другого конца»: вот, есть конкретная ситуация, вот, есть конкретный человек; произошло то-то и то-то, следовательно, его мозг (управляющий компьютер) отработал так-то и так-то. (А было ли там «счастье», «печаль» и «отвращение» – об этом в принципе можно вообще не говорить, разве что в качестве сопровождающей беллетристики).

§28. Письмо Сергея Марьясова от 2 августа 2011 года (Вариатор)

no Sergey sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 2. augusts 00:11
temats R-POTI-3
piegādātājs rambler.ru

Валдис, добрый день!

Неделю был занят так, что практически не оставалось времени на наш проект, а неделю был в отъезде. Да и пришлось немного подумать: вопросы творчества не хотели просто так раскрывать свои секреты ☺. В результате прилагаю некоторые свои заметки по §§23,24 R-POTI-3 и конспект §31 книги L-ARTINT.

Всего наилучшего,
СМ

Понятия Веданской теории «потенциального продукта» и «номиналии» в серии книг R-POTI впервые объясняются в R-POTI-2 (стр.8) и теперь после § 23 (R-POTI-3) подводят окончательную черту в моём застарелом споре с преподавателем диалектического материализма (R-POTI-3, стр.3). Хотя должен признать, что представители диалектического материализма были отчасти правы, когда критиковали представление о том, что «мозг вырабатывает мысль так же, как печень вырабатывает желочь»¹⁷⁶: ведь чтобы электрические сигналы мозга превратились в

¹⁷⁵ Kūle Maija, Kūlis Rihards. «Filosofija». Apgāds Zvaigzne ABC, b.g., b.v. © Apgāds Zvaigzne ABC, 1998.

¹⁷⁶ В.Э.: Врач и профессор физиологии Якоб Молешотт, говоривший на лекциях и писавший в книгах это изречение в середине 19-го века, конечно, не мог иметь представления о работе компьютеров. Но вся эта классическая троица (Фохт, Бюхнер и Молешотт), обруганная Энгельсом «вульгарными

мысли, необходимо, чтобы структура мозга создала некоторую информационную систему, результатом работы которой и будут мысли, идеи и то, что мы называем сознанием человека.

Кстати, думаю, что и закон (диалектического материализма) о переходе количества в качество, тоже не совсем корректное утверждение (подхваченное марксистами, так как оправдывало жестокость классовой войны). Эволюционный процесс показывает нам, что из имеющегося набора элементов путём случайного перебора можно создать новые устойчивые системы (с новыми качествами, отличающимися от качеств исходных элементов), которые полностью или частично можно снова использовать для создания новых систем, которые снова будут проверены на устойчивость. Но ведь некоторое вновь заявленное «количество» может ни к чему и не привести!

Возвращаясь к вопросу понимания природы мысли и мира идей, хочу привести цитату из книги «Сандра Блейкли, Джефф Хокинс. Об интеллекте. Москва–Санкт-Петербург–Киев. 2007», которая привлекла моё внимание и в которой делается попытка объяснения этого вопроса с физиологической точки зрения:

«...головной мозг – это ящик, в котором царят кромешная тьма и тишина. Мозг познает окружающую среду только посредством сигналов, поступающих по сенсорным нервным волокнам. Мозг как сигнальное устройство не воспринимает тело своего хозяина как нечто особенное, отличное от остального внешнего мира. Иными словами, границы между вашим телом и оставшейся частью внешнего мира для мозга не существует. В то же время кора мозга не может построить модель собственно головного мозга, потому что у мозга нет ощущений. Теперь становится понятным, почему наши мысли появляются независимо от наших тел, почему нам кажется, что душа или разум независимы. Ваши мысли существуют в мозге, они физически отделены от тела и остальной части внешнего мира. Разум независим от тела, но является продуктом деятельности мозга.»

Авторы книги описывают природу и работу мозга с точки зрения нейробиологии (и получается у них это довольно интересно), говорят о роли неокортекса в формировании сознания Человека Разумного, о самообучении, творчестве и других связанных с сознанием человека вопросах. И хотя один из авторов в прошлом «разработчик архитектуры компьютеров»¹⁷⁷, в книге нет ни одной попытки привлечь к объяснению природы человеческого сознания и разума понятие мозговых программ. А жаль: выводы авторов о возможностях и роли искусственного интеллекта в будущем могли бы быть точнее и доказательнее.

*

С деревом программы практически всё понятно, хотя хочется заметить, что вопрос о разрушении отработавших, или вообще не нашедших применения в реальной жизни отдельных ветвей программ может оказаться несколько сложнее. Тут у человека, наверное, действуют общие принципы запоминания и забывания¹⁷⁸: те ветви, которые были проработаны более тщательно (в уме, на плане, в групповом обсуждении или путём многократных тренировок в

материалистами», была людьми весьма умными, и вовсе не такими дурачками, какими они изображались в наших университетских курсах «диалектического материализма». Я не думаю, чтобы Мошотт полагал, что мысль представляет собой какое-то вещество. Видимо, его утверждение означало, что мыслительная деятельность может быть выведена из функционирования мозга, сведена к нему (и в таком случае он был прав). А это как раз и не нравилось Энгельсу, который, будучи дуалистом (а не материалистом! – см. {VIEWS.107} и {VIEWS.205 = МОИ № 100}) отстаивал несводимость «идеального к материальному».

¹⁷⁷ СМ. Сандра Блейкли, Джефф Хокинс. Об интеллекте. Москва–Санкт-Петербург–Киев. 2007, стр. 5. Далее Джефф Хокинс пишет о том, что в 1982 году в компании *Grid Systems* он «разработал новый язык программирования под названием *GridTask*, принесший компании большой успех...» (стр. 14).

¹⁷⁸ В.Э.: Разумеется, многие узлы верхнего и среднего уровней программного дерева сохраняются в памяти и пополняют «библиотеку заготовок» для будущего самопрограммирования. Однако ведь программа была нужна и для того, например, чтобы просто протянуть руку и открыть дверцу машины. Эта программа самого низкого уровня была составлена за полсекунды до того, как дверца открылась, а спустя секунду её уже не было в голове Доллии. Когда следующий раз надо будет открывать дверцу, будет составляться новая программа – из тех же заготовок, возможно, учитывая опыт выполнения предыдущей программы (если, например, в дверце какая-то особенность, как-то особо надо ручку повернуть и т.п.) – но это будет новая программа. Для каждого движения, каждого шага, каждого слова на низшем уровне непосредственно перед действием составляется программа этого действия, выполняется и тут же гасится (перекрывается новыми программами для этих же мышц). Нет ни возможности, ни необходимости все эти миллионы, миллиарды мелких программ хранить в памяти.

точной копии здания, как это в течении месяца делала группа захвата «террориста №1») будут сохраняться гораздо дольше и, при необходимости, будут мгновенно вызваны подсознанием человека и по прошествии длительного времени. Вообще, для человека предварительная подготовка дерева программы, аналогичному на рис. 4, кроме нижних уровней будет осуществляться в сознании (на осознанном уровне)¹⁷⁹, и, по мере тренировок, сама структура дерева (структура его узлов) будет передаваться в подсознание человека (на неосознанный уровень). Чем лучше пройдут тренировки, чем больше в них будет учтено возможных ветвлений, тем легче будет выполнение программы в реальном мире – переключение в узлах на новую ветвь программы будет происходить на подсознательном уровне, без привлечения сознания к оценке ситуации, размышлениям и построения новых ветвей программ (действий).

Сама структура дерева программы, а также и узлы, в которых оценивается переключение (или подключение) ветвей программ, скорее всего, будут подвержены тем же механизмам забывания, что и любая другая информация, которую обрабатывает операционная система человека.

Причём из опыта изучения иностранных языков известно, что то, что было связано с какими-то эмоциями, будет помниться гораздо дольше. Так, те ветки, которые будут задействованы в реальном мире, с момента, когда Доллия выйдет из машины со спрятанным под одеждой пистолетом (а может быть эта точка отсчёта включится немного раньше), – запомнятся Доллии надолго.¹⁸⁰ Потом, спустя много дней, она сможет неоднократно перебирать в своей памяти порядок произошедших в банке событий, оценивать упущенные или реализованные возможности, свои и чужие ошибки, и в причудливом виде какие-то сцены будут являться ей в её снах... (но об этом далее, в конспекте 31 параграфа книги L-ARTINT).¹⁸¹

Конспект Сергея Марьясова (довольно подробный) параграфа 31 «Генератор Разнообразия и Сонатор» из книги В. Эгле L-ARTINT

§ 31 Генератор Разнообразия¹⁸² и Сонатор

¹⁷⁹ **СМ.** Конечно, можно привести экспромт, как некоторый крайний пример, когда почти всё планирование предстоящих действий осуществляется в подсознании. Но при внимательном рассмотрении хорошего (удачного) экспромта, окажется, что почти все ветки, составляющих его программ, были разработаны (и может быть не раз использованы) в прошлом.

¹⁸⁰ **В.Э.:** Тут только надо различать, ЧТО запомнится: сама программа – или те обстоятельства, в которых она выполнялась, и её результаты. Сказанное тобою относится в первую очередь ко второму.

¹⁸¹ **В.Э.:** Да, разумеется, со всем этим я согласен. Единственное, что не совсем соответствует моим представлениям, – это та роль, которую ты отводишь «сознательному» и «подсознательному»; мне кажется, что тут на тебя еще «давит груз традиции». С точки зрения Веданской теории в момент, когда программа (мозговая) работает, вообще не имеет никакого значения то, происходит ли это «сознательно» или «подсознательно» – она просто работает. А работает она «сознательно» или «подсознательно» – это зависит только от того, направлен или не направлен на нее тот «луч прожектора», о котором мы говорили вокруг Рис.1. Это обстоятельство (направление «луча») не влияет на собственно работу программы, но оно влияет на то, сможет ли субъект или не сможет потом вспомнить, что программа работала, и, соответственно, проанализировать ее работу, результаты и т.д. «Подсознательно» отработавшая программа не будет фигурировать в дальнейших рассуждениях субъекта, а «сознательно» отработавшая – будет. Только в этом разница.

¹⁸² **СМ.** В оригинале «Randomgenerators», наверное, производное от английского «Random Generator». Однако в русскоязычном интернете это название имеет довольно узкий смысл, связанный с проблемой генерации случайных чисел, поэтому я предлагаю использовать название «Генератор разнообразия», что ближе к содержанию данного параграфа, и освобождает от вопроса менять ли английское «а» (в слове random) на русское «е» или «э» ☺. Можно было бы созвучно названию Сонатор использовать слово Вариатор (от английского слова *variety* – разнообразие, многообразие), но у него в интернете сильная привязка к бесступенчатой трансмиссии. Тем не менее, название Вариатор мне тоже нравится, поэтому окончательное решение оставляю на твоё усмотрение. **В.Э.:** Слова с корнем «рандом» используются (и давно) отнюдь не только для обозначения генераторов случайных чисел, а для указания присутствия случайного фактора вообще. Так, еще в 1960-х годах у нас в Университете на лекциях по статистике говорилось о рандомизации как приеме образования случайных выборок и т.п. Слово в общий обиход вошло, конечно, из английского языка, но в самом английском оно появилось из старофранцузского, где «randir» означало «быстро бегать» (таким образом, это слово родственно английскому «run»); это сами англичане запутались, где надо писать «и», где «а», где произносить «а», где «е», а мы не обязаны

.531. Теперь давайте подумаем, каким образом реализовать в операционной системе Доллии творчество. Если какая-то структура (программа или данные) в её компьютере уже существует и имеется (в той или иной форме) алгоритм, по которому эту структуру оценивать, тогда нам понятно, как реализовать работу соответствующего функционального блока.

.532. Из своего человеческого опыта, мы знаем, что творчество, пожалуй, самый сложный элемент человеческой деятельности. Например, прочитав книгу или посмотрев фильм, все мы можем относительно легко оценить произведение как «хорошее» или «плохое». Но каждый, кто пытался сам придумать новый сюжет для книги или фильма, знает, как это трудно. Также, как специалисты, мы знаем, насколько велика проблема, например, генерации случайных чисел в компьютерах.

.533. В таком случае, как компьютер Доллии может создать нечто совершенно новое, до того не существовавшее? Очевидно, что единственный алгоритм, который может быть использован в этом случае, это образование новых соединений путём комбинирования существующих элементов. Этот принцип мы будем использовать в работе многих блоков операционной системы Доллии (например, в Генераторе Программ, когда Доллия, – также как человеческое дитя, – будет осваивать репертуар возможных движений, по-разному сочетая некоторые базовые движения, и затем смотреть, «что получится»). Когда Доллия доберётся до литературного творчества, она также возьмёт набор каких-то уже существующих сюжетов и начнёт переставлять их различные элементы, выбирая лучшие варианты. (Тогда, чем больше в начальном сюжете сцен, чем разнообразнее множество задействованных элементов, и чем дольше и интенсивнее будут строиться и отбираться комбинации, тем лучше и оригинальнее может быть полученный результат).

.534. Впрочем есть одно место в операционной системе Доллии, где этот принцип случайного комбинирования, мы должны использовать задолго до того, как Доллия сможет обратиться к литературному творчеству. Каким образом Доллия вообще может обнаружить новые связи между вещами [объектами, ситуациями, явлениями и т.д.] во внешнем мире и своими чувствами? Каким образом вызвать в её поведении нечто новое и неожиданное, чтобы её поведение не было простым пассивным реагированием на события во внешнем мире, и чтобы она сама проявляла активность?

.535. С этой целью встроим в операционную систему Доллии специальный функциональный блок и назовем его, скажем, «Генератор Разнообразия». Пусть он непрерывно комбинирует определённые (для одних элементов) ситуации с широким спектром других известных Доллии элементов и передаёт эти комбинации в качестве исходного материала другим программам Доллии для последующего анализа. Тогда жизнь Доллии станет несравнимо активнее: в некотором месте реального внешнего мира (где-нибудь на природе, например, в лесу), может часами, или даже днями, не происходить ничего, что требовало бы от неё обязательного и неотложного реагирования,¹⁸³ но Генератор Разнообразия будет посылать её эти

за ними в этом следовать. Я вообще избегаю брать термины из английского языка (в знак «протеста» против его засилья) и предпочитаю первоисточники – греческий, латинский, в данном случае старофранцузский. Я сохраняю в Конспекте твои термины, но оставляю за собой право и впредь пользоваться словом «рандомгенератор». (За английским *variety* тоже стоит латинский субстантив «variatio» – «пестрота, разность, непостоянство», и далее верб «variare» «испестрить, преобразовать»; так что слово «Вариатор» от латинского; окончание -or – тоже латинское: *imperator* – повелитель, *victor* – победитель и т.д.).

¹⁸³ СМ. Видимо речь идёт об экологическом «убитом» лесу, где действительно ничего неожиданного произойти не может: даже муравей не пробежит, ни то что волки нападут ☺. А вот наблюдение за природой пусть даже в таком пустом лесу: там неяркий лесной цветочек пробился сквозь песок, а тут, по-видимому, ворона растерзала целлофановый мешок с остатками завтрака какого-то безалаберного туриста, – потребуют встроенного творческого процесса (процессора). **В.Э.:** Всё это так, все эти мелкие события не требуют обязательного действия субъекта. Если система ориентирована реагировать только на опасность, пищу и сексуального партнера (а таковы первоначальные установки всякого живого существа), то, несмотря на все эти муравьи, система будет стоять неподвижно часами. Это мы и можем наблюдать в природе у низших существ. Например, я, сколько ни посещал Рижский зоопарк, ни разу не видел, чтобы крокодил двигался. С таким же успехом в его клетку можно было поставить чучело (дешевле было бы: кормить не надо). А волки или тигры непрерывно бегают по клетке. Почему такая разница? ЧТО будет отличаться в их операционных системах? И вот, я отвечаю: у крокодила почти что нет рандомгенератора – совсем зачаточный (потому и снов он практически не видит). А у волков и тигров (соответственно, и их одомашнированных сородичей: собак и кошек) мощные генераторы, непрерывно создающие импульсы: побежать туда, обнюхать это и т.д. (и сны они видят, как, наверное, каждый сам наблюдал).

стимулы [вновь образованные комбинации для анализа] непрерывно. Она будет «вынуждена» непрестанно реагировать на все новые и новые «фантазии», тогда как за её пределами в мире все остается как и прежде; придется проверить все новые и новые, «приходящие на ум», связи между ситуациями и предметами.

.536. Функциональная роль Генератора Разнообразия и его необходимость нам очевидны: системы с Генератором Разнообразия будут гораздо более активны, чем системы без него, и их интеллектуальное развитие будет идти гораздо быстрее.¹⁸⁴ Аналогично, например, в генетике, появление эукариотов (то есть организмов с более чем одной хромосомой), значительно ускорило эволюционный процесс, так как позволило при оплодотворении по-разному комбинировать хромосомы матери и отца (и, тем самым, комплекты генов).¹⁸⁵

.537. Но теперь давайте посмотрим на отношения Генератора Разнообразия и Хроникера. Если деятельность Генератора Разнообразия и её результаты не попадут в Хронику и не будут зарегистрированы в памяти, то это будет «неосознанная» деятельность, и Доллия сама не будет знать ничего об этом. Когда система функционирует нормально, необходимости в отражении продукции Генератора Разнообразия в Хронике нет, поскольку есть более важные вещи, которые нужно фиксировать. Однако в те моменты, когда выключены все другие процессы, которые Х-селектор мог бы признать в качестве важных для записи в Хронику, в неё могут попасть также результаты работы Генератора Разнообразия. Тогда Доллия увидит сны: сцены, в которых в самом неожиданном порядке скомбинированы знакомые ей ситуации с объектами из совершенно других ситуаций.¹⁸⁶

.538. Собственно сны нам не важны для операционной системы Доллии, но нам абсолютно необходимы как её Генератор Разнообразия, так и Хроникер. Сны являются побочными эффектами работы этого аппарата в моменты, когда все «важные» процессы выключены.¹⁸⁷

.539. Системе важно, конечно, работа Генератора Разнообразия не во время сна, а во время бодрствования. Тот факт, что этот генератор работает в то время, когда человек бодрствует, каждый может проверить, немного потренировавшись. Я, например, теперь после небольшого количества попыток могу отключить все остальные свои «мысли», чтобы стала видимой продукция Генератора Разнообразия, – то есть не засыпая, могу «видеть сны».¹⁸⁸

.540. Функциональная роль Генератора Разнообразия, как мы видели, заключается в создании «разнообразного» и активного «внешнего мира» – такого, в котором есть большее количество различных комбинаций объектов и ситуаций, и в котором чаще «что-нибудь да и

¹⁸⁴ **СМ.** При первом прочтении данного параграфа это утверждение не показалось мне уж столь «очевидным». **В.Э.:** Речь здесь не об эволюции вида, а о развитии способностей отдельного индивида. Мощности рандомгенераторов отличаются и у отдельных людей. У одного непрерывно возникают всякие идеи (сделать то, попробовать это...), а другой тупо смотрит в даль (с появлением техники – в телевизор). Естественно, что первый вскоре окажется интеллектуально гораздо более развитым, чем второй (знающим, умеющим и т.д.).

¹⁸⁵ **В.Э.:** Об этом можно прочесть в {DVEA.641 = **МОИ** № 50}.

¹⁸⁶ **СМ.** Именно то, что объекты из знакомых ситуаций перемешаны с объектами из других ситуаций, и то, что происходит наложение ситуаций и внутри их происходит замена объектами из других ситуаций, часто приводит к полной неузнаваемости происходящих во сне событий, хотя их истоки лежат в известных (как на осознанном, так и неосознанном уровнях) данному человеку ситуациях.

¹⁸⁷ **СМ.** Значительная часть информации вокруг человека обрабатывается на неосознанном уровне (сознание человека слишком узко) и необязательно достигает осознанного уровня (например, оценка характера случайного собеседника или оценка в течении длительного времени изменяющейся обстановки), а во сне в результате случайной обработки она может попасть на осознанный уровень (в Хроникёр) и при отделении от других случайных наложений позволит выяснить причину, например, некоторого тревожного состояния. **В.Э.:** Об этом, пожалуй, стоит поговорить в отдельном параграфе.

¹⁸⁸ **В.Э.:** Было бы интересно, если бы и ты попробовал и описал свои наблюдения... Я, вот, прямо сейчас, после написания предыдущей фразы решил ради интереса попробовать: что мне на этот раз подкинет мой рандомгенератор? Лег на диван, прикрыл глаза рубашкой, чтоб свет не мешал, и стал расслабляться, но одновременно все-таки следя, что у меня перед глазами. Сначала ничего нет, только какие-то пятна и круги плывут. Но через некоторое время вдруг появляется картина: какая-то большая лестница, какие бывают, может быть, перед крымскими санаториями, а на ней какой-то толстый человек. Что за лестница, что за человек – не узнаю. Но чувствуется, что это может быть видно на какой-то фотографии или в каком-то фильме. Потом вдруг – отъезжающий от меня мотоциклист в каске... В общем – ерунда. Но видна работа генератора: вырывает из памяти какие-то случайные кадры и швыряет передо мной. И если к этим кадрам подключится Программатор и начнет генерировать программы моего поведения в этой (предполагаемой) ситуации, то уже пойдет полное сновидение с целой фабулой...

происходит», чтобы остальной части системы было бы из чего выбирать полезные комбинации, и были бы импульсы для активизации действий.¹⁸⁹ Таким образом, для остальной части системы Генератор Разнообразия выступает в качестве эмулятора внешнего мира. Поэтому и сны мы воспринимаем как данные нам извне, а не результатами нашей деятельности.

.541. Концепцию Фрейда о том, что сны есть «исполнение желаний», мы отвергаем (по крайней мере, в том виде, в котором Фрейд её изложил). «Старина Фрейд» слишком мало знал о проектировании операционных систем реального времени. Эти «исполнения желаний» не нужны операционной системе нашей Доллии, так же как не были нужны они и естественному отбору, когда он создавал людей. Но и нам, и естественному отбору абсолютно необходим Генератор Разнообразия, потому что без него невозможно создать хорошие, генерирующие новые идеи и активные системы.¹⁹⁰

.542. Отдаленная, приблизительная связь с концепцией сна по Фрейду, тем не менее, существует. Сны могут дать нам не только неожиданные комбинации известных нам ситуаций и объектов, но также сцены наших собственных действий в этих необычных ситуациях. Но такие сцены уже фигурировали в нашем проекте операционной системы: их создавал Анализатор Программ, прогнозируя получение ожидаемых результатов различных программ. Таким образом, в общем случае сны включают в себя не только результаты работы Генератора Разнообразия, но в их создании также участвуют и Генератор Программ и Проектор Программ. Если какую-то компьютерную программу, сгенерированную Доллией, обозначить словом «желание», а ожидаемый результат, полученный Проектором Программ, словами «исполнение желания», то у нас будет некоторое совпадение с Фрейдом, хотя в целом мы должны признать выводы Фрейда результатом преувеличения в одностороннем порядке роли одного из элементов, что и привело их к ошибочности.

.543. Для конкретного разговора о снах, мы введём понятие ещё одного функционального блока, назвав его «Сонатором»¹⁹¹. Генератор Разнообразия является центральным блоком Сонатора, впрочем Сонатор также использует Генератор Программ и Проектор Программ, в совокупности создавая те псевдореалистические сцены, в которых мы «участвуем» сами и, просыпаясь, называем их снами.

§29. О сновидениях

2011.08.04 15:56 четверг

В.Э.: Вернемся теперь к твоему подстрочному примечанию к Конспекту:

СМ. Значительная часть информации вокруг человека обрабатывается на неосознанном уровне (сознание человека слишком узко) и необязательно достигает осознанного уровня (например, оценка характера случайного собеседника или оценка в течении длительного времени изменяющейся обстановки), а во сне в результате случайной обработки она может попасть на осознанный уровень (в Хроникёр) и при отделении от других случайных наложений позволит выяснить причину, например, некоторого тревожного состояния.

Тут так сразу и не скажешь, правильно это или нет, потому что всё зависит от смысла, какой вкладывается в те или иные слова. В целом чувствуется влияние фрейдовских или постфрейдовских школ психологии, толкующих о «сознательном», «бессознательном», «осознании» и проводящих «психоанализ» снов для диагностики и, возможно, дальнейшего лечения.

Я вообще не люблю понятия «осознанный», «подсознание» и подобные как «неконструктивные» – не конструкторские. Человек, для которого они являются базовыми, не построит, не спроектирует систему искусственного интеллекта.

Что же касается фрейдовского «психоанализа» и его современных модификаций, то классический «психоанализ» Фрейда – просто бред (возможно, когда-нибудь дойдем до него); есть более разумные современные варианты, например, такие, как у Пека, но всё равно они преувеличивают роль анализа сновидений. На самом деле всё обстоит не так, как они пишут.

¹⁸⁹ **СМ.** В том числе и для построения номиналий.

¹⁹⁰ **СМ.** По сути дела, естественный отбор из имеющегося набора элементов создаёт новые устойчивые системы, которые полностью или частично использует для создания новых систем, которые снова проверяет на устойчивость.

¹⁹¹ **В.Э.:** В оригинале – сомниатор. «Somnium» на латинском языке – сновидение.

Во второй половине 1990-х годов, когда я особенно усердно занимался этим вопросом, разобрал «по косточкам» книги Фрейда¹⁹² и Пека¹⁹³, я одновременно проводил внимательные наблюдения над своими собственными сновидениями (теперь уже не делаю этого), сравнивая данную мне реальность сновидений с идеями этих авторов. И я могу сказать следующее.

Действительно, сновидение весьма часто сопровождается ярко выраженным, сильным чувством; особенно беспокоит и поэтому запоминается, когда это чувство тревоги, страха, угрозы и т.п., но чувство бывает и другим: ликующее, торжествующее и др. Пек считает, что это чувство вызвано сновидением, и что сновидение (из подсознания) что-то рассказывает, раскрывает об этом чувстве.

Ничего подобного! Из своего собственного опыта тех упомянутых наблюдений я могу сказать, что чувство первично. Оно НЕ порождено, не вызвано сновидением, как думает Пек, а имеется само по себе, и сновидение только сопровождает это (базовое) чувство, так сказать, аккомпанирует ему. В сновидениях всегда наблюдаются те чувства, которые были и без них, до них, наяву.

Сопровождая это базовое чувство всякими неожиданными комбинациями случайных картин, Сомниатор¹⁹⁴ не рассказывает нам в зашифрованном виде о причинах чувства (как это считает Пек), но он действительно генерирует исходное поле для отбора «новых идей» (в том числе и для объяснения причин чувства) – генерирует тем единственным способом, каким компьютеры могут создавать что-то новое.

Возвращаясь к тексту твоего Примечания, давай опустим эти неточные и расплывчатые термины «неосознанный», «осознанный» и будем говорить более по-конструкторски.

Итак, мозгом обрабатывается информация «вокруг человека», и часть этой информации и результатов обработки записывается в память, а часть – не записывается, отбрасывается. Это понятно.

Но то, что записано в память – это записано в Хронику, или у нас помимо Хроники есть еще какое-то файловохранилище? По тому, как модель нашей системы строилась на Рис.1 и вокруг него, Хроника была единственным хранилищем для такой информации, и всё, что запомнено, значит, прошло через Хроникер и находится в Хронике.

Правда, Хроника тоже неоднородна. Можно записать так, что запись легко доступна, а можно запихнуть её в какой-нибудь угол так, что потом эту запись не найдешь. Она там есть, но недоступна по стандартным алгоритмам поиска. Такая картина тогда была бы более точной и «конструкторской», чем использование терминов «осознанный», «неосознанный» и подобных.

Если Хроника у нас единственное хранилище, то, естественно, рандомгенератор может хватать кадры только из Хроники – но, правда, как из той её части, которая легко доступна для программ, ведущих логический поиск, так и из той части, которая по логическому поиску недоступна, – а рандомгенератору доступна именно потому, что он кадры не ищет, а хватает «наобум».

Ну вот, с такими уточнениями данное примечание можно принять.

§30. Письмо Сергея Марьясова от 16 августа 2011 года

no Sergey sevm@rambler.ru
kam Valdis Egle <valdis.egle@gmail.com>
datums 2011. gada 16. augusts 22:59
temats R-POTI-3
piegādātājs rambler.ru

¹⁹² В.Э.: Есть много изданий Фрейда, здесь главная работа «Толкование сновидений»: в моей библиотеке это издание Попурри, Минск 1997; кроме того сборники с множеством мелких работ, которые не перечислишь...

¹⁹³ В.Э.: В моей библиотеке это М.С. Пек. «Нехоженные тропы. Новая психология любви, традиционных ценностей и духовного роста», ЮНИТИ, Москва, 1996 (но вообще есть и другие русские издания и даже другие переводы). В оригинале: M. Scott Peck, M.D. *The Road Less Traveled. A New Psychology of Love, Traditional Values and Spiritual Growth*. Published by Simon & Schuster – New York, London, Toronto, Sydney, Tokyo, Singapore (1978).

¹⁹⁴ В.Э.: Мне не так легко перейти на твой термин «Сонатор», даже если бы я этого хотел, потому что у меня есть много старых (русских, но неопубликованных) текстов, где используются термины «сомниатор», «рандомгенератор» и т.п. Тогда мне пришлось бы все эти тексты пересматривать и менять.

Валдис, добрый день!

Просматривал книги по обсуждаемой нами тематике и нашёл набор интересных графиков <http://www.singularity.com/charts/page17.html>, похожий я видел в твоей книге R-DVESA (стр.32). Кроме того обнаружил, что первые главы многих книг доступны для просмотра на <http://www.amazon.com>.

Прилагаю два файла.

Всего наилучшего,

СМ

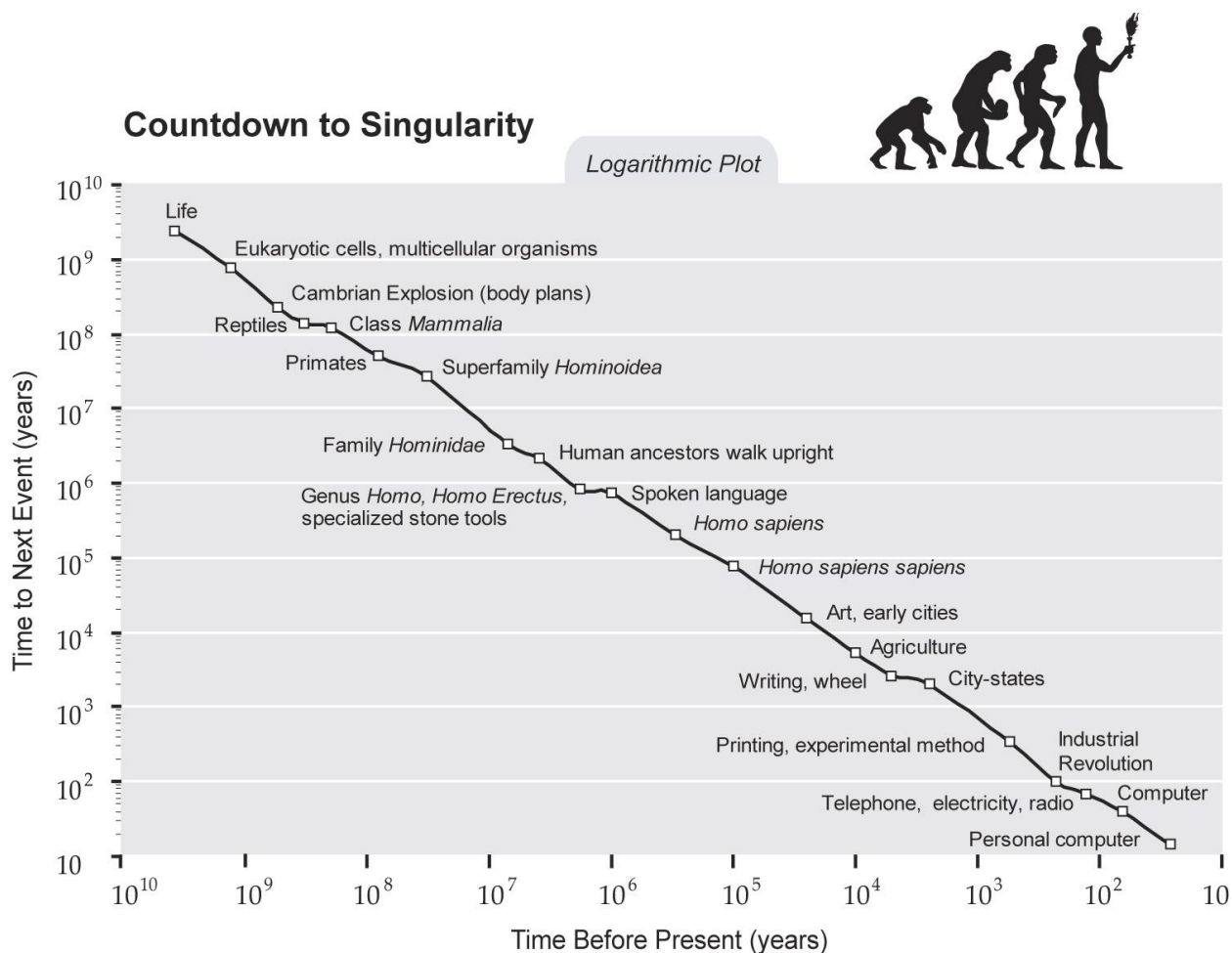


График с сайта <http://www.singularity.com/charts/page17.html>

Time Before Present	Time to Next Event	Event
3700000000	2400000000	Life
1300000000	750000000	Eucaryotic cells, multicellular organisms
550000000	220000000	Cambrian Explosion (body plans)
330000000	135000000	Reptiles
195000000	113500000	Class Mammalia
81500000	49000000	Primates
32500000	25500000	Superfamily Hominoidea
7000000	3100000	Family Hominidae
3900000	2100000	Human ancestors walk upright
1800000	800000	Genus Homo, Homo Erectus, specialized stone tools

1000000	700000	Spoken language
300000	200000	Homo sapiens
100000	75000	Homo sapiens sapiens
25000	15000	Art, early cities
10000	5000	Agriculture
5000	2490	Writing, wheel
2510	1960	City States
550	325	Printing, experimental method
225	95	Industrial Revolution
130	65	Telephone, electricity, radio
65	38	Computer
27	14	Personal Computer

Chart Countdown to Singularity, Events expressed as Time before Present (Years) on the X axis and Time to Next Event (Years) on the Y axis, Logarithmic Plot Page 17, Linear Plot page 18.

§31. Присоединенные файлы к письму от 16 августа (Селектор, Реактор и Навигатор)

С учётом статьи об инфантильности представлений (§27 данной книги) приходится проделать некоторую «работу над ошибками». Действительно на стр. 71 мне следовало ввести какие-то другие обозначения для описываемых мною стратегий, например, «счастье-восстановление», «счастье-творчество» и «печаль-забывание». Первые две стратегии мне интересны тем, что они весьма полезны для функционирования биологической системы под названием Человек, и в перспективе было бы интересно посмотреть, нужно ли реализовать что-то подобное в Доллии.

Что касается терминов Рандомгенератор и Сомниатор, то теперь, после того, как параграф 31 разобран, а о работе Сомниатора добавлен ещё один интересный параграф, применение твоих терминов для меня не является затруднительным. Кроме того, я не сторонник плодить множество разных названий для обозначения одних и тех же понятий, блоков и так далее. И так при чтении любой книги приходится подстраиваться под терминологию её автора, а тут ещё и многообразие в обозначении одного и того понятия, – это только лишнее напряжение для читателя. Но поскольку наша переписка является не только изложением некоторого материала, но и записью истории обсуждения каких-то вопросов («хроникой» ☺), то далее в переписке я буду придерживаться терминов, принятых в твоих книгах.

С прим. 175 согласен.

А вот прим. 173 могло бы стать интересным продолжением разговора о самопрограммировании и дереве программы. Ты пишешь: «...Когда следующий раз надо будет открывать дверцу, будет составляться новая программа – из тех же заготовок, возможно, учитывая опыт выполнения предыдущей программы...». Вопрос составления программ из заготовок мы в первом приближении уже рассмотрели, а вот как мог бы быть реализован «опыт выполнения предыдущей программы»? Я думаю, что накопление такого опыта тесно связано с формированием номиналий в мозгу человека и их развитием по мере накопления некоторого опыта. Движение запоминается не само по себе, а в связи с параметрами, характеристиками объекта или объектов, над которыми предстоит что-то сделать.¹⁹⁵ Одним из таких объектов является само человеческое тело: так, например, в занятиях йогой человек познаёт возможности своего тела: его мышц, лёгких, других внутренних органов.

¹⁹⁵ СМ. Валдис, мог бы ты описать, как строятся номиналии объектов и их пространственное восприятие?

Ты пишешь: *«Для каждого движения, каждого шага, каждого слова на низшем уровне непосредственно перед действием составляется программа этого действия, выполняется и тут же гасится (перекрывается новыми программами для этих же мышц)...»* Но музыканты, танцоры, спортсмены, рабочие у станка бесконечно большим количеством повторений одних и тех же движений добиваются идеального выполнения определённого набора движений. То есть, некоторая центральная ветвь дерева программ, отвечающая за это движение (и, или набора движений) сохраняется в памяти и выполняется, как программа, в которой большинство переменных заранее известно, и некоторые из них имеют фиксированное значение.

Дети не умеют открывать дверцы до определённого момента. А научившись, уже не забывают, как это делается, они используют «опыт» однажды успешно выполненной программы.

Я думаю, что ручку своего автомобиля я открываю практически без участия Хроникера. То есть, сколько я не пытался после каждой поездки в течение последних двух недель вспомнить, как я это проделывал, — мне не удалось. Я делал это «машинально»¹⁹⁶. После поездки я мог вспомнить, например, птичий помёт на стекле дверцы (к хорошей поездке ☺), или то, что думал о предстоящем маршруте, но вспомнить, думал ли я о ручке, не говоря о том, смотрел ли я на неё вообще, я не мог. Конечно, это нисколько не доказывает, что программа открытия дверцы моего автомобиля, которую я выполняю десятки раз в течение одной недели, не собирается из более мелких блоков за доли секунд до того, как я занимаю определённое положение около закрытой дверцы (это положение — результат выполнения другой моей регулярной программы: каждый автолюбитель знает насколько «удобно» открывается дверь его автомашины в обычных условиях, и насколько «трудно» это сделать, когда машина припаркована в луже или посреди блестящего гололёда). Но несомненно, что запоминаются какие-то её элементы или особенности сборки в отношении конкретных объектов (ручка моего автомобиля, его дверь, и сам стоящий автомобиль).¹⁹⁷



Защита от детского Рандомгенератора. Родители рассчитывают, что к моменту, когда дети научатся открывать такой замок, у них уже будет некоторая безопасная модель применения содержимого шкафчика.

¹⁹⁶ В.Э.: В терминах Рисунка 1 это означает, что Прожектор в это время был направлен на птичий помёт, на мысли о маршруте и т.п. Но ты можешь (если захочешь) его один раз направить и на ручку и свою руку. Сама программа же не может не существовать, потому что импульсы мышцам БЫЛИ посланы, и они не могут появиться «ниоткуда»; кто-то должен был (что-то должно было) сформировать и отправить эти импульсы.

¹⁹⁷ СМ. Хорошо помню, как папа учил меня открывать ручку двери купе железнодорожного купейного вагона (происходило это где-то в 70-х годах). Особенностью механизма было наличие в нём защиты от проскальзывания язычка дверного замка в случае горизонтальных ударов (сцепка вагонов, инерционное движение дверей при трогании или остановке поезда, а также дребезжание на стыках рельс). В механизме, когда он был новым (могу это только предположить ☺), нажатие на ручку двери одновременно прижимало механизм вниз и освобождало язычок для дальнейшего открывания, а из-за износа механизма дверь при нажатии на ручку соскальзывала или вправо или влево, язычок блокировался, и ручка на открывание не проворачивалась. Мне никак не удавалось подобрать правильное движение руки: дверь открывалась в одном случае из пяти. Оказалось, что в случае сильно разболтанного механизма, надо второй рукой зафиксировать (снизу) ось поворота дверной ручки в пространстве!

Я несколько раз пытался проделать опыт с просмотром продуктов Рандомгенератора (по аналогии с абз. 539, прим. 183), но ничего толком разглядеть не успевал – засыпал ☺. Вернее, я что-то начинал видеть, но если я пытался проанализировать, что я вижу, то изображение отступало, а если просто продолжал наблюдать в расслабленном состоянии, то успевал запомнить только сам факт появления картинки, а что было в ней изображено, запомнить не мог (или не успевал?). Видимо, это каким-то образом связано с моим нынешним общим психическим состоянием, потому что я хорошо помню, что в определённый период (годовой давности) я успевал заметить перед засыпанием, что мысли мои принимали странный (фантастический) оборот, появлялись картины, в которых я эти мысли уже обсуждал с какими-то людьми. Мне удавалось как-то с боку заметить, что такие странные мысли в нормальной жизни мне в голову не приходят, и какие-то действия, какие-то люди тоже не могут появиться, так как я собирался спать и уже должен лежать в кровати. Результатом такого бокового анализа был вывод, что я нахожусь в переходном состоянии ко сну, – после чего я полностью погружался в сон, и после этого в моей памяти сохранялись уже только какие-то обрывки утренних снов.

Конспект Сергея Марьясова (довольно подробный) параграфа 32 «Селектор, Реактор и Навигатор» из книги В. Эгле L-ARTINT

§ 32 Селектор, Реактор и Навигатор

.544. Теперь в операционной системе Доллии мы имеем два разных источника первичных импульсов: реальные события во внешнем мире и материалы¹⁹⁸ [импульсы], выданные Генератором Разнообразия.

.545. Далее, опытному разработчику очевидно, что для того, чтобы организовать реагирование Доллии на эти импульсы, нам потребуется две принципиально разные функциональные группы. Во-первых, понятно, что ни одна система не будет реагировать на абсолютно все импульсы, которые приходят непрерывным потоком (да в этом и нет необходимости). Поэтому нужно отобрать из всех этих импульсов те, на которые Доллия будет реагировать (а остальные она «пропустит мимо ушей»). Во-вторых, будут необходимы функции, которые непосредственно сгенерируют саму реакцию, как только некоторые импульсы пройдут через первый уровень «сита».

.546. Как и раньше, если выделены некоторые функциональные группы, то нужно ввести соответствующие функциональные блоки и присвоить им имена. Первый блок назовём, скажем, «Селектором», второй «Реактором». Взаимодействие Селектора и Реактора в первую очередь определит наблюдаемое извне поведение Доллии в конкретной ситуации (ее реакцию). Общее поведение будет определенным балансом между работой Селектора и Реактора: чем больше [импульсов] отберёт Селектор, тем больше будет загружен Реактор, и наоборот: чем меньше импульсов Селектор признает достойным внимания, тем более тщательно Реактор может сгенерировать реакцию и дольше её подготавливать.

.547. Однако, если Доллия будет просто реагировать на внешние или внутренние импульсы, то она будет «щепкой на волнах», которую несёт то в одну сторону, то в другую, в зависимости от полученных импульсов. Поэтому встроим в неё еще один блок, который должен «удерживать общий курс» среди всех этих волн (импульсов)¹⁹⁹, для реализации глобальной жизненной стратегии Доллии.²⁰⁰ По ассоциации с волнами назовём этот блок «Навигатором».

¹⁹⁸ **СМ.** В книге R-POTI-1 (стр.24) вводится несколько следующих определений: «...В общем случае всякая программа что-то берет (обозначим для будущих рассуждений это «что-то» словом «материал»), производит над этим материалом какие-то действия в определенной последовательности (обозначим состав этих действий и их последовательность словом «алгоритм») и в результате что-то получает (обозначим это результирующее «что-то» словом «продукт»)». До сих пор в своих конспектах параграфов книги L-ARTINT я этой терминологии, вообще говоря, не придерживался: в данном случае «материал» – исходные данные для дальнейшей обработки другими программами Доллии. **В.Э.:** Да, всё правильно, это продукты рандомгенератора и материалы для других программ. В том письме Фрейбергу, которое ты конспектируешь, я вообще-то слово «материал» использовал вне этих определений, в «общечеловеческом» смысле, как, например, в фразе «материалы уголовного дела показывают...».

¹⁹⁹ **СМ.** На стр. 56 данной книги мы говорили о том, что выбор «эмоционального состояния» подготавливает выбор некоторой стратегии поведения Доллии, которая может быть запущена на выполнение при наступлении некоторых пороговых условий. То есть, отчасти, Эмоциатор также выступает

.548. Совсем не реагировать на внешние и внутренние импульсы Доллия не может и не должна, но с другой стороны, необходимо удерживать «генеральный жизненный курс». Суммарное поведение Доллии будет определенным балансом между курсом, назначенным Навигатором, и реакциями на импульсы от Рандомгенератора (Генератора Разнообразия) или из внешнего мира.

.549. Ясно, что Реактор и Навигатор являются теми основными блоками, которые будут использовать ранее определенные нами функции Программатора, так как самопрограммирование нужно не само по себе, а для генерации действий, представляющих собой либо некоторую реакцию, либо осуществление общего курса.

.550. Давайте теперь подумаем, на каких принципах может работать Реактор. Как правило, в общем случае некоторая конкретная ситуация никогда не будет до конца ясной, и будет характеризоваться некоторым недостатком информации.²⁰¹ Тем не менее, реакция все же должна быть сгенерирована, так как чаще хоть какая-то реакция лучше, чем не делать ничего вообще. Поэтому встроим в Доллию возможность при необходимости быстро (хотя и без достаточных оснований) по некоторому внешнему²⁰² алгоритму принять некоторую модель текущей ситуации²⁰³, а затем, подготавливая свою реакцию, опираться на эту модель.

.551. Теперь Реактор Доллии может работать по одному из двух принципов: либо путем тщательного анализа ситуации (и на основе этого анализа создать максимально подробную и точную модель ситуации), либо «высосать модель из пальца» (без тщательного анализа ситуации). В соответствии с терминологией, принятой в психологии Юнга и его последователей, назовем первую стратегию «сензитивной»²⁰⁴, а вторую «интуитивной».

в роли некоторого «сита», защищающего Доллию от необходимости реагировать на все импульсы (внешние и внутренние), и разделение анализатора импульсов на блоки Эмоселектор и Селектор достаточно условное. **В.Э.:** Нет, это не так. Эмоселектор и Селектор – блоки функциональные. То есть, мы поступаем следующим образом: видим, что в системе требуется выполнение какой-то группы функций, и все те средства, которые эти функции осуществляют, мы относим к соответствующему функциональному блоку. Функция Эмоселектора – это обнаружить ситуации, в которых требуется изменение состояния системы, установить, какое именно состояние требуется, запустить эмогенератор и в конечном счете перевести систему в другое состояние (с другим набором параметров на входах и выходах процессоров). Функция же Селектора – это обнаружить ситуации, в которых требуется реакция системы, запустить Программатор для выработки программ этой реакции и в конечном счете отдать приказы мышцам на те или иные действия. Это разные функции и разные конечные результаты. Другое дело, что при физической реализации программ осуществление этих (отличающихся) функций может быть как-то совмещено в одном модуле, который работает «сразу по обоим направлениям». Но это уже не вопрос идеологии системы, а удобства реализации. С точки зрения идеологии: функции разные, и блоки разные. Кроме того, отработка обоих этих блоков по времени тоже не будет всегда совпадать. Иногда требуется реакция системы, а для этой реакции требуется перевод системы в другое состояние. Тогда они отработают (с положительным выходом) оба одновременно. Но в других случаях (и таких, в общем-то, большинство) реакция системы может быть сгенерирована БЕЗ перевода ее в другое (эмоциональное) состояние (в том же – как правило, спокойном, – состоянии). Тогда только Селектор выдаст сигналы, инициирующие что-то дальнейшее, а Эмоселектор «промолчит». И, с другой стороны, иногда никакая реакция системы не требуется (Селектор молчит), а Эмоселектор переводит систему в другое эмоциональное состояние (например, в «тоску без причины»; особенно это характерно для болезненных состояний, например, циклотимии или маниакально-депрессивного психоза).

²⁰⁰ **СМ.** О том, как формируется «глобальная жизненная стратегия», «генеральный жизненный курс» у человека достаточно хорошо описано в книге R-DVESA (Валдис Эгле, «Двеса», [§27](#), стр. 76–77). В перспективе было бы интересно посмотреть, как это могло бы быть реализовано у Доллии, и этот вопрос, наверное, перекликается с вопросом о стартовых блоках в прим. 156 данной книги.

²⁰¹ **СМ.** Строго говоря, любая реальная ситуация никогда не является до конца ясной, и всегда характеризуется некоторым недостатком информации: тем не менее операционная система человека легко с этим «уживается» и, может быть, в будущем имеет смысл об этом поговорить подробнее.

²⁰² **В.Э.:** В оригинале «по заменяющему алгоритму». То есть, имеется в виду, что если бы у Доллии была информация о ситуации такая, которую можно назвать более менее «полной», то Доллия отработала бы по алгоритму *A*. Но так как информации нет, ситуация не ясна, то Доллия не приходит к решению использовать алгоритм *A*, и запускает другой алгоритм *B*, который мы считаем (для фактической ситуации) «заменяющим».

²⁰³ **СМ.** Валдис, мог ли ты немного подробнее рассказать об этих алгоритмах: что они из себя представляют и откуда берутся.

²⁰⁴ **СМ.** Сензитивность (лат. *sensus* – чувство, ощущение) – характеристика органов чувств, выражающаяся в их способности тонко и точно воспринимать, различать и избирательно реагировать на

.552. Обе эти стратегии Реактора необходимы Доллии. Система, которая будет использовать только одну из них, не будет эффективной: если она будет руководствоваться только сензитивной стратегией, то она не сможет реагировать на неясные ситуации; а если она будет ориентироваться только на интуитивные стратегии, то сгенерированные реакции в совокупности будут гораздо хуже, чем у систем, которые имеют гибкость в применении двух стратегий в зависимости от ситуации. Итак, еще раз, требуется определенный баланс между сензитивной и интуитивной стратегиями действий Реактора.

.553. Итак, интуиция Доллии есть принятие некоторой модели ситуации без тщательного анализа, согласно некоторому замещающему алгоритму, руководствуясь только известными отдельными элементами ситуации. Никакой другой интуиции в компьютере не может быть (и, если принят постулат, что человек представляет собой биологический компьютер, то никакая другая интуиция не может существовать и в организме человека).

.554. Интуитивно принятая модель может оказаться правильной или неправильной. В первом случае хороший результат достигается с гораздо меньшими усилиями, чем в случае сензитивной стратегии. Поэтому (правда, не только потому) люди, которые больше привыкли пользоваться интуитивной стратегией, пытаются изобразить её более высокой формой психической деятельности, чем «сензитивное» мышление. Мы это [представление] отклоняем: интуитивное мышление есть деятельность в упрощенных, непроверенных моделях по неточным алгоритмам. Оно в сумме даст результаты хуже, чем «сензитивное» мышление. Также инсайт²⁰⁵ или озарение является подобным явлением.

.555. Поскольку мы различаем два типа генерации реакции, то, по своему обыкновению, будем считать, что они также реализуются двумя разными подблоками Реактора, и назовем один, скажем, «Сензитивреактором» (С-реактором), а второй для «Интуитивреактором» (И-реактором).

.556. И, наконец, в работе также и Селектора выделим две группы функций, и соответственно введём два подблока. Первый для рассмотрения и оценки собственно импульсов (внешних или внутренних), а второй для просмотра Хроники, чтобы оценить прошлый опыт, имеющийся у Доллии в результате этих [аналогичных, похожих?]²⁰⁶ импульсов. Оценка этих двух факторов будет участвовать в окончательном решении Селектора о том, надо ли инициировать генерацию реакции в данной ситуации. Первый подблок назовём «Импульс-анализатором» (И-анализатором), а второй «Хрониканализатором» (Х-анализатором).

*

В.Э.: Ну, Сергей, я вижу, ты делаешь успехи в освоении латышского языка ☺. Уже чувствуется некоторая легкость перевода...

§32. Некоторые особенности Рандомгенератора

2011.08.18 13:55 четверг

На этот раз ты задал мне целую кучу вопросов:

Как мог бы быть реализован «опыт выполнения предыдущей программы»?

Валдис, мог бы ты описать, как строятся номиналии объектов и их пространственное восприятие?

Валдис, мог ли ты немного подробнее рассказать об этих { .550 } алгоритмах: что они из себя представляют и откуда берутся.

По каждому из этих вопросов можно целые книги писать...

слабые, мало отличающиеся друг от друга стимулы. (Краткий словарь психологических терминов. <http://vocabulary.ru/dictionary/16>.) **В.Э.:** Да, таковы первоисток этого понятия, но есть еще дальнейшая история. Уже и сам Карл Юнг, но в особенности его последователи, такие как разработчики знаменитой американской психологической типологии MBTI, разделяли людей на таких, которые ориентированы (и полагаются в своих действиях) в основном на чувства, ощущения (на «сензитивность»), и таких, которые больше полагаются на «интуицию» (причем первые считались «приземленными», малоспособными, а вторые – способными, одаренными, в общем: «высшими»).

²⁰⁵ **СМ.** Инсайт – (от англ. *insight* – проницательность, проникновение в суть, понимание, озарение, внезапная догадка) – интеллектуальное явление, суть которого в неожиданном понимании стоящей проблемы и нахождении её. (Википедия. <http://ru.wikipedia.org/wiki/инсайт>).

²⁰⁶ **В.Э.:** Имеется в виду: «в результате импульсов такого рода».

Прежде, чем ответить на них так, как это возможно в рамках одного письма, создаваемого всего за несколько дней, я хотел бы еще немножко задержаться у Рандомгенератора и высказать о нем пару вещей, которые в прошлый раз я приберег для анализа тех картин, которые ты, возможно, опишешь по своим наблюдениям. Но так как ты конкретные картины не описал, то приходится говорить на основе моих собственных наблюдений.

Первая вещь, которую я хотел бы отметить, это то, что в кадрах, предоставляемых рандомгенератором, как правило, есть движение. Но – одно элементарное движение. Я никогда не видел, чтобы двигались многие объекты; движется всегда что-то одно: либо поступательно в каком-нибудь направлении, либо туда-сюда вибрируя.

Например, первый кадр, который я четко зафиксировал и проанализировал в 1990-х годах, когда этими вопросами стал вплотную заниматься, был такой. Тоже тогда писал, устал, прилег, закрыл глаза, и вдруг передо мной картина: узкая щель между высокими домами, и в этой щели проехала машина: легковая, светлая, слева направо. (Эти картины существуют очень короткое время – доли секунды, поэтому, естественно, некогда их рассматривать и анализировать; их надо сперва просто запомнить, то есть, перезаписать в память в другое место, откуда их потом можно по необходимости вызывать – вспоминать).

Ну, и так как я в то время как раз изучал Фрейда и формировал позицию Веданской теории по сновидениям, то я эту проехавшую машину стал всесторонне изучать: что это такое и откуда она взялась? В конце концов я узнал эту щель: это вид из одного переулка в городе Цесис (где живет моя мама и брат). Переулок выходит на улицу Rīgas; машина проехала по этой улице, а из переулка видна была в узкой щели между домами. По этому переулку я в 1970-х годах ходил из маминной квартиры на вокзал. Но потом я перестал использовать этот путь, ходил другими маршрутами и никогда больше не ходил по той улочке. Таким образом, в 1990-х годах, когда рандомгенератор мне подсунил машину в щели, этому кадру было минимум 20 лет. (Да и машина была явно советская – более всего похожая на «Победу»).

Вот какую чепуху (абсолютно неважную и не значимую) способна хранить наша Хроника и выдавать через десятилетия! Не удивительно, что мы потом не в состоянии опознать, что это такое, и многие склонны приписывать такие видения «внешним силам».

Но рассказал я об этом из-за движения: в кадре было движение – двигалась машина. Кадры памяти похожи на .gif файлы (а не, скажем, на .jpg). Там нет многочисленных движений, нет целого «видеоклипа», но есть какое-то одно простое движение, которое ловится, выделяется и запоминается нашей системой восприятия. Эта система ориентирована на отслеживание и обнаружение единичных движений. (Что вполне понятно с точки зрения Естественного отбора).

Вторая вещь, которую я хотел бы отметить, это то, что рандомгенератор не всегда выдает просто кадры из памяти. Уже и раньше я наблюдал комбинированные кадры, но в следующую же ночь после отправки тебе предыдущей версии этой книги, утром проснувшись и пытаясь еще раз заснуть, я увидел такую картину. Там были детские качели, возле них слева стояли дети, числом около пяти, нормальные дети. А собственно в качелях сидела какая-то образина и, как будто глумясь, открывала и закрывала рот. Она не качалась; движение должно быть одно, и в данном случае это было движение рта образины.

Качели узнать было нетрудно: я каждый вечер прогуливаюсь по паркам Риги и прохожу мимо двух детских площадок: в Верманском саду и на Эспланаде, а также подолгу сижу на скамейках возле этих площадок, наблюдая за детьми. Но немыслимой образины в качелях там, естественно, никогда не было. Тогда я стал думать: что это за образина, откуда она взялась? В конце концов узнал-таки ее. Это кадр из мультфильма – рисованного, американского, типа диснеевских. Движения рта образины – это она говорит в фильме. Я вообще-то не смотрю мультфильмы, но жена смотрит, и иногда приходится кинуть взгляд на телевизор.

И вот, рандомгенератор не просто взял эти два кадра из памяти. Он их скомбинировал в один кадр – работа с программистской точки зрения вообще-то довольно сложная! То есть, он действительно генератор, а не просто передатчик кадров.

Причем чувствуется, что кадры между собой связаны: их объединяет один мотив – детский.

Как можно догадаться, алгоритм отработки Рандомгенератора в данном случае должен был быть таким:

- 1) взять из Хроники случайный кадр № 1 (с детской площадкой, детьми и качелями);
- 2) определить его как обладающий признаком «детский»;

- 3) произвести в Хронике поиск другого кадра № 2, тоже обладающего признаком «детский»;
- 4) выделить и вырезать из кадра № 2 некоторый (скорее всего, главный) объект;
- 5) выделить в кадре № 1 некоторое место (собственно качели);
- 6) поместить вырезанный из кадра № 2 объект в выделенное место кадра № 1;
- 7) суммарный кадр подsunуть Селектору как заменитель реальности.

§33. Некоторые особенности низшего самопрограммирования

Вернемся теперь к дверце машины и к программе, которая управляет рукой, дверцу открывающей.

Эта программа – звено самой низшей инстанции в программном дереве; ниже уже программ нет; дальше идут непосредственно импульсы мышцам.

Попробуем себе представить, как эта программа может реально выглядеть, что из себя представляет. Импульсы идут по нервам к мышцам. Наш мир детерминирован; эти импульсы не могут появиться «ниоткуда». Прежде чем импульсы пошли, их «причина» была закодирована в каких-то клетках мозга (человека или Доллии). Вот эти клетки и содержат «последнюю программу дерева». Пожалуй, можно с уверенностью сказать, что эта «последняя программа» почти полностью соответствует потокам импульсов – а, точнее, наоборот: потоки импульсов почти полностью соответствуют программе. Поэтому по потокам импульсов мы можем судить о самой программе.

Что же мы знаем о потоках импульсов при открывании дверцы машины?

Во-первых, мы знаем, что даже к одной мышце импульсы идут по нескольким нервным волокнам и, значит, даже программа, управляющая одной мышцей, имеет некоторую «протяженность». Это целый участок мозга, хотя и небольшой.

Во-вторых, мы знаем, что в открывании дверцы задействованы многие мышцы: надо и тело поворачивать, и ногу удобнее поставить, и руку протянуть, и пальцами двигать. Значит, «протяженность» программы еще больше увеличивается.

На уровне «последней программы» вообще нет четких границ, где начинается «программа открывания дверцы» и где кончается «последний шаг к машине». На самом деле здесь одна сплошная программа непрерывного управления всем телом на протяжении длительного времени. На более высоком уровне такая граница была гораздо четче, потому, что имелось намерение именно открыть дверцу, а не делать что-то другое, но далее оно расплзлось в сплошное управление телом.

В-третьих, мы знаем, что в ходе этого «сплошного управления телом» потоки импульсов к тем или другим мышцам то возрастают, то спадают – меняются во времени. Соответственно, значит, меняется и «последняя программа»; она постоянно корректируется в ходе самопрограммирования на низшем уровне.

Предположим, что Доллия стояла около машины и разговаривала с кем-то. Потом вдруг ей «пришло в голову» открыть дверцу машины. Это «пришло в голову» означает, что аппарат ее самопрограммирования построил в дереве общей программы ее деятельности узел «открыть дверцу», т.е. некоторую программу среднего уровня или несколько ниже среднего. Как теперь этот узел дерева (который еще только узел, а не программа для мышц) превратить в реальный поток импульсов к мышцам?

Очевидно, во-первых, что у Доллии должна быть библиотека заготовок программ для разных действий, и она в этой библиотеке должна найти заготовку для открывания дверцы, вызвать ее оттуда, «активировать». Если Доллия когда-нибудь раньше открывала дверцу (а, тем более, если она это делала часто), то такая заготовка в библиотеке будет. (А если раньше не делала, то ей придется заготовку создавать: «обучаться», как дверца открывается, что, разумеется, гораздо труднее, чем просто вызвать уже существующую заготовку).

Но это еще только заготовка программы; она не может быть готовой, окончательной программой, потому что окончательная программа должна учитывать ту конкретную позу, в которой Доллия сейчас, в данный момент, стоит по отношению к машине, и эти позы (а также, возможно, другие обстоятельства) в разных случаях будут разными. На базе этой заготовки, корректируя и пополняя ее, Доллия должна составить окончательную программу своих действий по открытию дверцы и «загрузить» ее в поле общего, «сплошного» управления телом, а потом

еще и корректировать программу в ходе выполнения, особенно если возникают какие-то непредвиденные обстоятельства (например, ручку заклинило).

Ты спрашиваешь: «Как мог бы быть реализован «опыт выполнения предыдущей программы»?».

В первую очередь – в виде корректировки заготовки. Заготовка ведь не мраморная глыба, раз и навсегда данная. Она тоже создавалась (в процессе обучения), и этот процесс ее формирования в общем-то продолжается и дальше, хотя и не так интенсивно, как в начале. Во вторую очередь – просто в запоминании информации, которая будет использована в следующий раз в тот момент, когда из заготовки будет делаться окончательная программа. (Впрочем, эти «очереди» названы в порядке их важности; хронологически, скорее, будет наоборот: сначала Доллия просто запомнит информацию, а если ситуация будет повторяться, то скорректирует заготовку).



Рис.6. Принципиальная схема создания программы открытия дверцы автомобиля

Ты пишешь:

Но музыканты, танцоры, спортсмены, рабочие у станка бесконечно большим количеством повторений одних и тех же движений добиваются идеального выполнения определённого набора движений. То есть, некоторая центральная ветвь дерева программ, отвечающая за это движение (и, или набора движений) сохраняется в памяти и выполняется, как программа, в которой большинство переменных заранее известно, и некоторые из них имеют фиксированное значение.

Правильно, такая «автоматизация» действий означает, что всё больше и больше элементов переносятся в заготовку, а всё меньше и меньше остается допрограммировать, когда из заготовки делается окончательная программа. Но это одновременно означает, что эта программа сможет

быть выполнена только во всё более и более узких условиях: никаких тебе других поз, как у Доллии, когда она открывает дверцу, никаких других меняющихся обстоятельств.

Принципиальная схема составления программы низшего уровня изображена на Рис.6.

§34. Как строятся номиналии объектов и их пространственное восприятие?

2011.08.20 15:00 суббота

Ты задал мне вопрос, вынесенный в заглавие этого параграфа. Но там фактически два вопроса: «Как строятся номиналии?» и «Как организуется пространственное восприятие?».

Номиналии в Веданской теории – это внутрикомпьютерные объекты, соответствующие внешним объектам (именуемым реалиями) или такие, о которых предполагается, что они соответствуют внешним объектам (в таком случае их реалии «воображаемы», «идеальны» в философском смысле, а не материальны); внутрикомпьютерная обработка номиналий в обоих случаях почти не отличается.

Номиналии строятся настолько разнообразными путями, что охватить все варианты не только в одном параграфе с описанием механизмов пространственного восприятия, но и вообще в одном коротком рассказе практически невозможно. Вот, я выглянул в окно, увидел на улице машину – и у меня в голове ее номиналия. Как она строилась? Отображение картины внешнего мира на сетчатке глаза, передача этого изображения в мозг, отработка программ выделения отдельных объектов из общей картины... Тут тогда надо описывать все эти процессы.

Но вот я вспоминаю машину, которая в 1970-х годах проехала в щели Цесисского переулка и была мне подкинута Рандомгенератором в 1990-х годах. Теперь я вспоминаю (вызываю из памяти) не ту, первую запись 1970-х годов (ее я из памяти вызвать не могу; ее не найти), а вторичную ее перезапись, сделанную в 1990-х, которую я найти и вызвать могу. Таким образом, у этой несчастной (и совершенно никому не нужной и не важной) «Победы» в моем мозге целая куча номиналий: и та затерявшаяся запись 1970-х годов, и тот дубликат 1990-х годов – это в долгосрочной памяти, а когда я непосредственно «вспоминаю» ее, то еще одна номиналия возникает в оперативной памяти... И где тут граница: считать ли все эти образы одной номиналией или разными?

А образина в качелях? Она тоже номиналия; я описал выше предполагаемый алгоритм, по которому она должна была создаваться. А подобных алгоритмов куча; по необозримому множеству алгоритмов мозг может создавать всевозможные номиналии воображаемых объектов. (Почитай «Мифологический словарь»!). Чего стоит одна только математика с ее тысячами понятий, то есть – номиналий! А другие науки?

Так что вопрос о способах построения номиналий в его общем виде пока придется отложить, в дальнейшем возвращаясь к нему на конкретных примерах построения конкретных номиналий. Номиналия – это очень фундаментальное понятие Веданской теории (и поэтому номиналии очень разнообразны), и главная задача этого понятия – отделить то, что существует во внешнем мире, от того, что существует во «внутреннем мире» и подчеркнуть, что всякая «фантазия» и «воображение» всё равно имеет материальную структуру, материальный объект во «внутреннем мире».

А здесь нам лучше будет больше внимания уделить второй части твоего вопроса.

§35. Пространственное восприятие

Построение у Доллии пространственного восприятия – это один из самых сложных вопросов во всем конструировании Доллоса, как я это многократно уже отмечал (см., напр., {PENRO1}²⁰⁷; вообще в комментариях к Пенроузу о пространственном восприятии говорится многократно {PENRO3}²⁰⁸, {PENRO5}²⁰⁹ и др.).

Конструируя пространственное восприятие у Доллии, мы в первую очередь должны определиться: копируем ли мы человеческую систему, или отклоняемся от нее. Для начала примем первый вариант и попытаемся скопировать человеческую.

²⁰⁷ МОИ № 14, стр.32, сноска 67.

²⁰⁸ МОИ № 15, стр.10, сноска 19.

²⁰⁹ МОИ № 16, стр.60, сноска 43.

У человека нет единой системы восприятия пространства – этих систем минимум четыре, и в нормальных условиях производится постоянное согласование их результатов, то есть, отыскивается интерпретация, наложение результатов одной системы на результаты другой, дающее максимум совпадений. Эти системы такие:

- 1) зрительная;
- 2) слуховая;
- 3) двигательная;
- 4) «внутренняя» (или «абсолютная» или «абстрактная»).

Начать лучше с последней. Абстрактная система пространства базируется на «гироскопе» внутреннего уха (возможно, и на других датчиках тоже), которые всегда (в земных условиях, не считая космические корабли с невесомостью) могут точно установить, где «верх», а где «низ», т.е. направление гравитационного поля Земли. Это дает Доллии (и человеку) одну «абсолютную» ось: вертикаль. Абстрактная система присоединяет к вертикали еще две оси, перпендикулярные ей и перпендикулярные между собой. Таким образом получается «трехмерное евклидовое пространство».

Следует обратить внимание на то, что это пространство Доллия создала САМА (!). Это никакое не внешнее пространство, а «внутреннее», порожденное одним только вводом трех осей. (Для простоты будем говорить о трех осях, где положение объекта характеризуется типа «выше–ниже», «вперед–назад», «правее–левее» по оси, хотя наряду с «декартовыми координатами» у человека наверняка используются и угловые; оба эти способа характеристики локализации объектов эквивалентны с математической и с логической точек зрения).

Наличие этого «внутреннего пространства» (т.е. наличие трех осей) означает, что Доллия способна присваивать пространственные характеристики различным объектам (пока мы говорим только о воображаемых объектах). Вот, она лежит (с закрытыми глазами) и думает (воображает): «Справа от меня стол, а за столом шкаф». Это означает, что в ее мозге строится номиналия стола и номиналия шкафа, и в этих номиналиях (как структурах данных) фиксируются пространственные характеристики реалий (материальных или воображаемых), соответствующих этим номиналиям.

Таким образом, «внутреннее пространство» Доллии, это не какой-то объект, полностью готовый и полностью данный; это ее потенциальная способность характеризовать объекты их расположением по осям.

Так как способ характеристики объектов по их взаимному расположению относительно осей представляет собой в общем-то алгоритм (используемый в различных строящихся Доллией программах наряду с другими алгоритмами), то мы можем определить «внутреннее пространство» Доллии как потенциальный продукт этого алгоритма. Внутреннее (абсолютное, абстрактное) пространство: «это все те места, которые можно охарактеризовать данным способом», и оно потенциально бесконечно.

В Доллии мы можем встроить кодировку пространственного положения объектов относительно осей и прямо в числовом виде в таких единицах измерения как метры, но в Природе, конечно, этого нет. Там кодировка относительна: «тот объект дальше, чем этот и правее», «а тот еще дальше и еще правее», «а тот совсем далеко и высоко высоко» и т.п.

Наличие этого алгоритма кодирования, присвоения объектам характеристик по трем осям, определяет людскую уверенность, что «наше пространство – это трехмерное евклидово пространство». Молотое-перемолотое в литературе «ньютоновское абсолютное пространство» – это на самом деле не что иное, как вот это внутреннее пространство Доллии (и людей), возникающее исключительно потому, что алгоритм кодирования взаимного расположения объектов у нее (и у нас) базируется на трех перпендикулярных осях.

Меня уже довольно давно (лет 15) занимает вопрос: «А было ли обязательным то обстоятельство, что Естественный отбор встроил в нас аппарат кодирования пространства, базирующийся именно на трех осях? А можно ли было встраивать аппарат, базирующийся, например, на четырех осях?» (Тогда «наше» пространство было бы четырехмерным евклидовым). У меня нет уверенного ответа на этот вопрос, но есть сильное подозрение, что можно было. Естественный отбор просто выбрал минимальное количество осей, дающее эффективный

результат.²¹⁰ (Было бы интересно это проверить экспериментально в системах искусственного интеллекта).

Но не будем отвлекаться. У Доллии, как и у людей, три оси, то есть, пространственное расположение характеризуется тремя независимыми величинами (хранящимися в номиналии, но относящимися к реальности).

Имея такое «внутреннее пространство», Доллия теперь открывает глаза и видит красочный и богатый внешний мир. Теперь всё, что она видит, локализуется в ее внутреннем пространстве (т.е. всему присваиваются характеристики по указанному алгоритму Трех осей). Зрительная система восприятия пространства налагается на внутреннюю, и они состыковываются. (Ср. эксперименты с очками, упомянутые в {PENRO5}).

Зрительная система восприятия пространства базируется на взаимном расположении изображений объектов на сетчатке глаза и на повороте глаз при сфокусированном взгляде на предмет при бинокулярном зрении. На основе этой информации и присваиваются пространственные характеристики каждой видимой точки (или области) по алгоритму Трех осей, строится общая пространственная картина. (Бывают и ошибки, т.н. иллюзии).

На всё это налагается еще слуховая система восприятия пространства, которая базируется на временной разнице, с какой одинаковые звуки приходят в разные уши. Тут тоже делается попытка согласовать результаты этих вычислений (аналоговых) с результатами предыдущих систем.

И, наконец, двигательная система. Она «измеряет» пространство теми усилиями мышц, какие необходимо приложить, чтобы нужную точку достигнуть. Именно эти усилия в первую очередь и кодируются в заготовках программ низшего уровня (таких, как заготовка программы открытия дверцы машины). Именно эти усилия и корректируются тогда, когда из заготовки перед выполнением делается окончательная программа.

Эта же система локализует в пространстве части собственного тела. (Ведь, чтобы рассчитать правильное движение руки к какому-нибудь предмету, необходимо знать не только то, где находится предмет, но и то, где находится рука).

Эта система тоже накладывается на предыдущие, и все они соединяются в некоторый сложный ансамбль (сложенный в нормальных условиях, а при дефектах наблюдаются разные интересные эффекты, описанные в психиатрии,²¹¹ когда системы выходят из согласия; например, после принятия наркотиков «Алисе» кажется, что она «10 футов высотой» {R-ALICE}²¹², её двигательная система пространства вышла из согласия со зрительной).

§36. Откуда берутся алгоритмы

2011.08.21 12:30 воскресенье

Остался еще только один из заданных тобой вопросов. Он задавался в таком контексте:

.550. ...встроим в Доллию возможность при необходимости быстро (хотя и без достаточных оснований) по некоторому внешнему алгоритму принять некоторую модель текущей ситуации²¹³, а затем, подготавливая свою реакцию, опираться на эту модель...

Во-первых, вспомним, что такое вообще алгоритм. Алгоритм – это «идея программы». Не существует программы без алгоритма; раз есть программа, значит, у нее есть какой-то алгоритм: плохой ли, хороший ли, изящный или неуклюжий – но есть.

Во-вторых, раз система производит какие-то действия, значит, у нее есть программа, реализующая эти действия (а, стало быть, и алгоритм этой программы).

²¹⁰ В.Э.: Подобным образом у нас именно два глаза потому, что это минимальное количество глаз, позволяющее определять расстояние до предмета (путем вычисления разности поворота глаз). У нас именно два уха потому, что это минимальное количество ушей, позволяющее определить направление, откуда идет звук (путем вычисления разности во времени прихода звуковых волн).

²¹¹ В.Э.: О метаморфозах и расстройствах телесной схемы см., напр., {L-EGLIT} <http://vekordija.narod.ru/L-EGLIT.PDF>, <https://yadi.sk/i/CaOFDIO33KELdi>.

²¹² <http://vekordija.narod.ru/R-ALICE.PDF>, https://yadi.sk/i/32_hK8q53KENDT.

²¹³ СМ. Валдис, мог ли ты немного подробнее рассказать об этих алгоритмах: что они из себя представляют и откуда берутся.

Поэтому всё это трио «действие–программа–алгоритм» неразрывно связаны, и на какое из этих слов в том или ином тексте ставится акцент, – это вопрос лишь формы выражения, а не существа дела.

Итак, система получила какой-то внешний стимул (что-то произошло, возникла какая-то ситуация), и ей надо как-то отреагировать на этот стимул (произвести какие-то действия, значит составить программу этих действий; значит, найти алгоритм для этой программы).

Как она это делает? Конечно, самый простой и массовый способ – это просто взять заготовку программы из «базы программ» (вместе с её алгоритмом) и выполнить эту программу. Тогда всё оригинальное творчество системы при генерации данной реакции сводится к тому, чтобы определить, какую именно заготовку надо взять, и к тому, чтобы эту заготовку чуточку подкорректировать соответственно конкретным условиям.

Естественно, заготовка должна быть предварительно в эту библиотеку помещена. Это делается на протяжении всей жизни в ходе обучения и тренировок (предыдущих выполнений этой же программы – или, точнее, программ, сгенерированных на основе этой заготовки).

Стандартный способ улучшения алгоритмов – это «метод проб и ошибок». Система меняет какое-нибудь звено в алгоритме и пробует: что получится? Плохо получилось – больше не пробует. Удачно получилось – пробует еще раз, и еще, и закрепляет в заготовке. Так и наращивает постепенно свои алгоритмы. Многие животные подсматривают, как это делают другие их сородичи и подражают им, т.е. учатся у них. (И есть даже один вид животных, которые организовали специальную институцию для обучения своих детёнышей, называемую ими «школой»).

Вот, так и создаются постепенно всё более и более мощные библиотеки алгоритмов или, что то же самое, базы программ.

Принципиальная схема тех вещей, о которых говорилось в п.550, изображена на Рис.7.

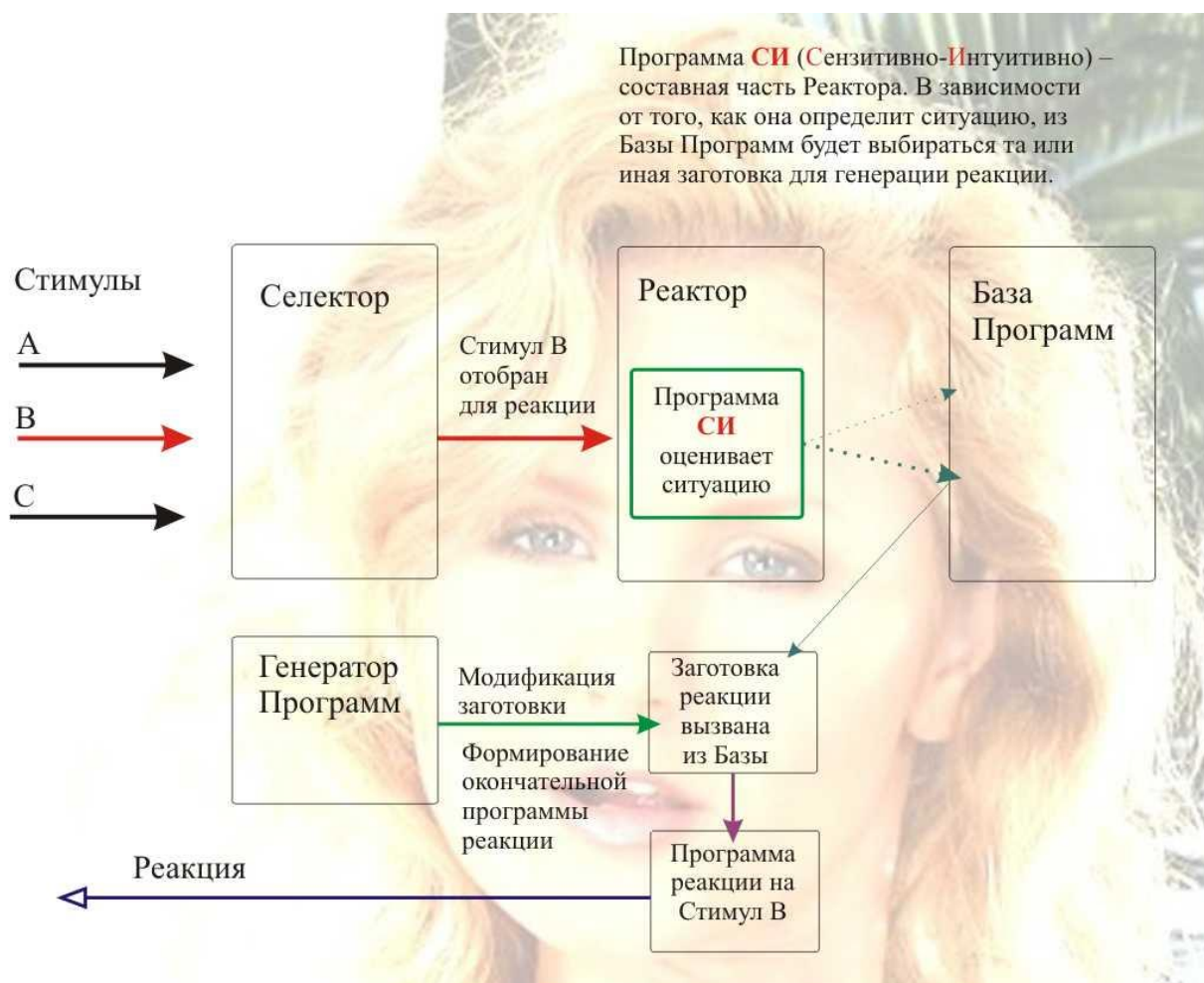


Рис.7. Принципиальная схема генерации реакции на внешний стимул

Итак, Сергей, я ответил на все твои вопросы, насколько это было возможно в пределах разумного расхода времени. Теперь ты законспектировал все «технические» параграфы выбранного тобой документа. Вводную и заключительную части этого документа я перевел на русский язык сам и помещаю здесь ниже:

§37. Вводная и заключительная части Конспекта

Перевод Валдиса Эгле (довольно точный)
параграфа 27 «Незнакомая операционная система»
из книги L-ARTINT

.472. Посвящается Иманту Фрейбергу,
профессору информатики
Квебекского университета в Монреале

§27. Незнакомая операционная система

1999.06.21 19:03 понедельник

.473. Основной постулат Веданской теории, как известно, таков: умственная деятельность человека представляет собой функционирование биологического самопрограммирующегося компьютера (или системы обработки информации, называемой мозгом). Теперь рассмотрим с точки зрения компьютерной науки или информатики некоторые главные выводы из такого постулата, вытекающие, если такой постулат был бы принят. Как это методологически правильно при разборе выводов, сам постулат мы здесь не оцениваем (и, следовательно, не оспариваем), оставляя вопрос о его принятии или отвержении для решения по другим критериям.

.474. Если принят такой постулат, то мы сразу можем смотреть на человеческое существо как на такую материальную систему, которая действует под управлением определенной операционной системы. Тогда, как специалисты по информатике, мы можем задать себе вопрос: «Как должна быть построена эта операционная система, чтобы наблюдались все те эффекты, которые мы видим в деятельности человека?».

.475. Здесь ситуация похожа на ту, в какой мы находились бы, если бы получили какой-нибудь неизвестный нам компьютер с работающей в нем неизвестной операционной системой (допустим, его во вражеской стране «сперли» шпионы нашего государства и теперь прислали нашим специалистам для изучения). Какова система команд (инструкций) этого компьютера, мы не знаем, исходных текстов (*source texts*) мы не имеем, деассемблировать программы мы не можем... Можем только пускать их на выполнение, смотреть, как программы ведут себя в разных ситуациях, и думать, как они в таком случае могут быть устроены.

.476. В нашем бюро имеются некоторые слишком молодые (или, наоборот, слишком старые) сотрудники, которые рьяно берутся за дело и начинают изучать: Вот, этот блок компьютера изготовлен из такого-то вещества; вот, когда компьютер выполняет такие-то действия, то вот по этому кабелю проходят импульсы такой-то формы, а когда он выполняет те действия, то вот по этому проводу проходят волны вот такой-то формы... И если вынуть из процессора вот эту микросхему, то наблюдается вот какой эффект, а когда выполняют вот это действие, то активной является вот эта область оперативной памяти...

.477. Не отрицая, что такое рвение наших сотрудников может иметь некоторую ценность, мы всё же, видимо, смотрели бы на их деятельность с ухмылкой. Изучая такие вещи и накапливая сведения такого рода, мы, кажется, сможем узнать о внутреннем устройстве незнакомой операционной системы весьма мало. (Однако при изучении человека психология, физиология и психофизиология до сих пор шли именно по такому пути, и он считался «подлинно научным»; поэтому не следует, видимо, удивляться, что в этих исследованиях в конце концов достигнуто так мало, и психическая деятельность человека всё еще остается «тайной за семью печатями»).

.478. Сами мы, напротив, проектировали и создавали операционные системы²¹⁴; мы хорошо знаем, что на самом деле всё то, что изучается нашими старательными сотрудниками, имеет совсем случайный характер. С таким же успехом при действиях указанного вида активной могла оказаться совсем другая область оперативной памяти (не существенно, где именно размещена данная подпрограмма), с таким же успехом для постройки процессора могли использовать совсем другие материалы, пускать сигналы по другим проводам и в виде импульсов другой формы.

.479. Когда мы проектируем операционные системы (и другие большие компьютерные программы), мы действуем прямо противоположно. «Нисходящее конструирование программ»²¹⁵ Йодана (*Yourdon*) для нас является элементарным и обыденным инструментом в нашей работе, без владения которым вообще нельзя считаться специалистом по информатике, и когда мы начинаем проектировать свою систему, то нас не интересуют ни материалы, из которых будут изготовлены компьютеры, ни системы их инструкций, ни кабели связи, ни даже языки программирования. Мы сперва все вопросы решаем концептуально.

.480. Исходя из функций, предусмотренных для системы, мы (на совершенно абстрактном уровне) решаем, какие вообще нужны будут функциональные блоки, каким образом они между собой будут взаимодействовать, какую информацию друг другу будут посылать (конечно, не на уровне сигналов, а принципиально), какие нужны будут структуры данных и т.д. Потом мы начинаем крупные функциональные блоки всё больше и больше конкретизировать, всё больше и больше детализировать, всё больше и больше структурировать, и только в самом конце приходим к тем вопросам о материалах, кабелях связи, импульсах и других деталях (к тому же даже и тогда еще мы эти вещи считаем маловажными, где одно решение очень легко можно заменить другим).

.481. Если человеку, который привык мыслить в таком духе, поручат исследовать какую-то чужую и незнакомую операционную систему, то совершенно естественным ему покажется браться сперва не за изучение проводов и импульсов, а предпринять нечто такое, что можно было бы назвать перепроектированием системы. Он начнет думать, как ОН поступил бы, если ему нужно было бы проектировать систему с данными функциями, какие нужны были бы функциональные блоки, принципиальные структуры данных и т.д.; он начнет думать, какие технические решения неизбежны (по-иному и нельзя делать, чтобы система работала), и где могут быть несколько равносильных решений (и какие именно). И кажется, что в конце концов у него понимание неизвестной операционной системы окажется намного глубже, чем у «исследователей проводов», хотя он, может быть, не делал ничего другого, как только сидел у своего письменного стола и смотрел в потолок.

.482. Именно таким методом к изучению умственной деятельности человека подходит Веданская теория. Нашу позицию можно выразить так: «Человек? Ладно, оставим пока человека в покое; его изучали долго и усердно. Решим сначала такую задачу. Нам дана механическая кукла, – но очень очень тонкая кукла –, заполненная огромным количеством всевозможных механизмов, совершенно эквивалентных всем органам человека. У этой куклы в голове имеется чрезвычайно мощный компьютер, объем памяти и быстроедействие которого полностью эквивалентны соответствующим параметрам человеческого мозга. Нет только одного – операционной системы для этого компьютера. И вот, нам как специалистам по информатике дано задание: спроектировать эту операционную систему так, чтобы в деятельности куклы наблюдались все те же самые эффекты, которые наблюдаются в деятельности человеческого существа».

.483. Наша деятельность, конечно, будет «чисто спекулятивной». Мы не будем сидеть в лабораториях у микроскопов и осциллографов, не будем изучать энцефалограммы и статистические таблицы психологических тестов. Мы просто сядем за стол, поднимем глаза к потолку и начнем думать, какая операционная система должна быть в голове у куклы (ради удобства присвоим ей имя – ну, скажем, Доллия) – значит, начнем думать, как нам строить операционную систему Доллии, чтобы получить все те эффекты, что начальство от нас требует.

.484. Говорят, что знаменитый английский физик новозеландского происхождения, лауреат Нобелевской премии Резерфорд (*Rutherford*) однажды поздно вечером зашел в свою лабораторию и увидел там одного своего сотрудника, который усердно что-то делал. Резерфорд у него

²¹⁴ **В.Э.:** Тут 12 лет назад, в 1999 году, я немножко польстил профессору Фрейбергу, приравнив его опыт своему. А теперь, когда он поступил так по-свински, у меня уже нет желания ему льстить, и я скажу, как есть: я делал операционные системы, а Имант Фрейберг не делал.

²¹⁵ Yourdon Edward. «Techniques of program structure and design». Prentice-Hall, Inc., Englewood Cliffs, New Jersey, 1975.

спросил: «Вы всегда так долго работаете?». Юноша, ожидая похвалы, гордо ответил: «Да!». Но Резерфорд помрачнел и только сказал: «Вполне достаточно работать четыре часа в день. В остальное время надо думать».

.485. Наука о человеке работала очень долго. Теперь не помешает и немножко подумать.

(Здесь находились законспектированные С. Марьясовым §§ 28–32)

Перевод Валдиса Эгле (довольно точный)
параграфа 33 «Результаты проектирования» из
книги L-ARTINT

§33. Результаты проектирования

.557. Итак, мы в основных чертах спроектировали операционную систему для механической, управляемой компьютером куклы Доллии: выделили главные функциональные блоки или группы функций, определили их принципиальное взаимодействие, рассмотрели разные технические решения.

.558. Это система, которая, начиная с данных ей «стартовых кирпичиков» в виде элементарных программных блоков, сама себя будет программировать дальше, постепенно накапливая на разных уровнях всё более сложные программы; будет способна анализировать свои будущие программы с точки зрения возможных последствий их выполнения, для этой цели генерируя их потенциальные результаты в виде специальных структур данных («картин»); она будет способна своевременно настраиваться на ту или иную работу, заранее активизируя необходимые процессы («эмоции»); она будет накапливать в специальной хронике информацию о своей предыдущей деятельности и таким образом сможет эту деятельность и ее результаты анализировать и учитывать в дальнейшем программировании своей деятельности («сознание»); она будет иметь свой внутренний генератор импульсов к деятельности, который будет поставлять свою продукцию наравне с импульсами внешнего мира; она будет способна генерировать новые комбинации прежде полученных элементов и ситуаций, постепенно отбирая и накапливая те, что дают всё более хорошие модели; в ней будет существовать динамическое равновесие между, с одной стороны, отбором тех импульсов, на которые требуется реакция, и, с другой стороны, генерацией реакции на отобранные импульсы; равновесие между импульсами внутреннего и внешнего генератора с одной стороны, и общим стратегическим жизненным курсом с другой; равновесие между сензитивной и интуитивной стратегиями при генерации реакции.

.559. На этом первом, концептуальном уровне мы выделили ряд необходимых для системы групп функций и ввели в нее соответственно 19 функциональных блоков и подблоков. Некоторые, может быть, скажут, что декларирование всех этих блоков, введение разных «генераторов», «анализаторов» и подобных дерзких названий является слишком спекулятивным и ничего не может дать. Им я отвечу, что у них просто нет опыта проектирования больших компьютерных систем. Я поступаю лишь так, как поступает любой хороший специалист, когда он должен спроектировать систему с заданными параметрами.

.560. Конечно, для дальнейшего проектирования и реализации операционной системы нужно будет все эти блоки конкретизировать и детализировать, как мы это обычно делаем, проектируя компьютерные системы, но и уже сейчас, спроектированная на этом, принципиальном уровне, операционная система Доллии совершает целый переворот в двух «старых» науках, связанных с умственной деятельностью человека, — в математике и психологии — совершает, как только мы принимаем дополнительный постулат о том, что умственная деятельность человека и есть функционирование такой «операционной системы», построенной по этим же или по похожим принципам.

.561. В математике упомянутый постулат вместе со знаниями о структуре и работе операционной системы впервые за более чем 4000-летнюю историю этой науки дает ей реальный предмет: математике не приходится уже начинать с абстрактных понятий, таких как «число», «арифметическое действие», «функция» и др., которые хоть и кажутся людям необыкновенно ясными, но о которых никто не может сказать, что же они из себя представляют, откуда и каким путем возникают.

.562. Зная, что вся математика возникла и была развита в этом «компьютере Доллии», мы можем все математические объекты свести к внутренним структурам этого компьютера. «Абстрактные понятия» математики теперь превращаются в потенциальные продукты разных программ «компьютера Доллии», которые получены, анализируя эти программы «сбоку», без их

выполнения – значит, поступая именно так, как это было необходимо для реализации самопрограммирования Доллии. Тогда соотношения между математическими объектами, которые прежде выглядели почти что мистическими, превращаются в реальные соотношения между потенциальными продуктами этих программ.

.563. Становится также ясно, каким образом изучение этих соотношений может быть полезно человеку, – ибо на самом деле полезны сами эти программы, – и обнаруживается ответ на тот вопрос, который в прошлом задавали сотни ученых («..каким образом изучение каких-то там абстрактных утверждений может давать такую огромную пользу в практической жизни!?!..»), и на который до сих пор не смог ответить никто.

.564. Предъявление и точное определение фактического предмета математики по большей части не изменяет ничего в результатах науки математики. Подавляющее большинство математических фактов и выводов сохраняются, становясь лишь реальнее, осязаемее и яснее (сохраняется всё, что было действительно полезно, ибо отображало объективную реальность в соотношениях между продуктами мозговых программ). Однако в некоторых местах математики более точное определение предмета этой науки приводит к иным результатам, чем это признано в теперешней математике. Так это имеет место с канторовской теорией о бесконечностях, которая полностью разрушается. (Но исчезает только то, что нигде и никогда не использовалось в реальной жизни, ибо это и невозможно было использовать: это не соответствовало никакой реальности).

.565. В психологии «операционная система Доллии» (вместе с нашим основным постулатом – значит, суммарно: Веданская теория) дает совершенно новое, не встречавшееся ранее представление о всех психологических делах; она впервые превращает психологию из чисто эмпирической науки (которая только накапливает и систематизирует факты наблюдений) в теоретическую науку, которая руководствуется какими-то предварительными предпосылками и может из них получить определенные выводы. И это мы сейчас рассмотрим подробнее.²¹⁶

§38. Общая схема Доллоса

2011.08.21 17:41 воскресенье

Завершая Конспект данного документа, я приведу еще ту схему {ARTINT.605}, с которой Сергей Марьясов 26 апреля 2011 г. (§3) начал разговор на эту тему (Рис.8).

Я даю эту схему такой, какой она выглядит в книге ARTINT и какую я её создал теми средствами, какие были в моем распоряжении в 1990-х годах. Лишь латышские надписи я заменил на русские в соответствии с Конспектом.

В этой схеме все функциональные блоки, вводившиеся в Конспекте, сведены воедино, чтобы их можно было обозреть все сразу.

Там есть один элемент, не упоминавшийся в Конспекте; это «База Выводов».

Согласно концепции, принятой в этой схеме (и во всем документе), Память делится на три части:

- Хроника, в которой запоминается (если отобрана Х-селектором) вся «внешняя» информация: что было, что увидено, что услышано, что я делал и т.п.;
- База Программ, в которой хранятся все накопленные заготовки программ, на основе которых формируются окончательные программы для выполнения;
- База Выводов хранит «внутреннюю» информацию, вычисленную уже самой системой, а не данную ей извне, как в Хронике; там хранятся все построения, сделанные самой системой (ее убеждения, взгляды и т.п.).

Конечно, во многом Хроника и База Выводов переплетаются, например, теорему Пифагора подавляющее большинство систем получают сначала как внешние данные и зафиксировывают в Хронике, но, проверив её доказательство и убедившись в его справедливости, система в своей Базе Выводов отметит этот факт. А если она сама придет к какой-нибудь новой теореме, то эта теорема у нее будет уже целиком в Базе Выводов.

Таков смысл этого понятия.

²¹⁶ В.Э.: Далее следовал второй документ той серии – о психологии, адресованный профессору психологии Монреальского университета в Канаде – Вайре Вике-Фрейберге.

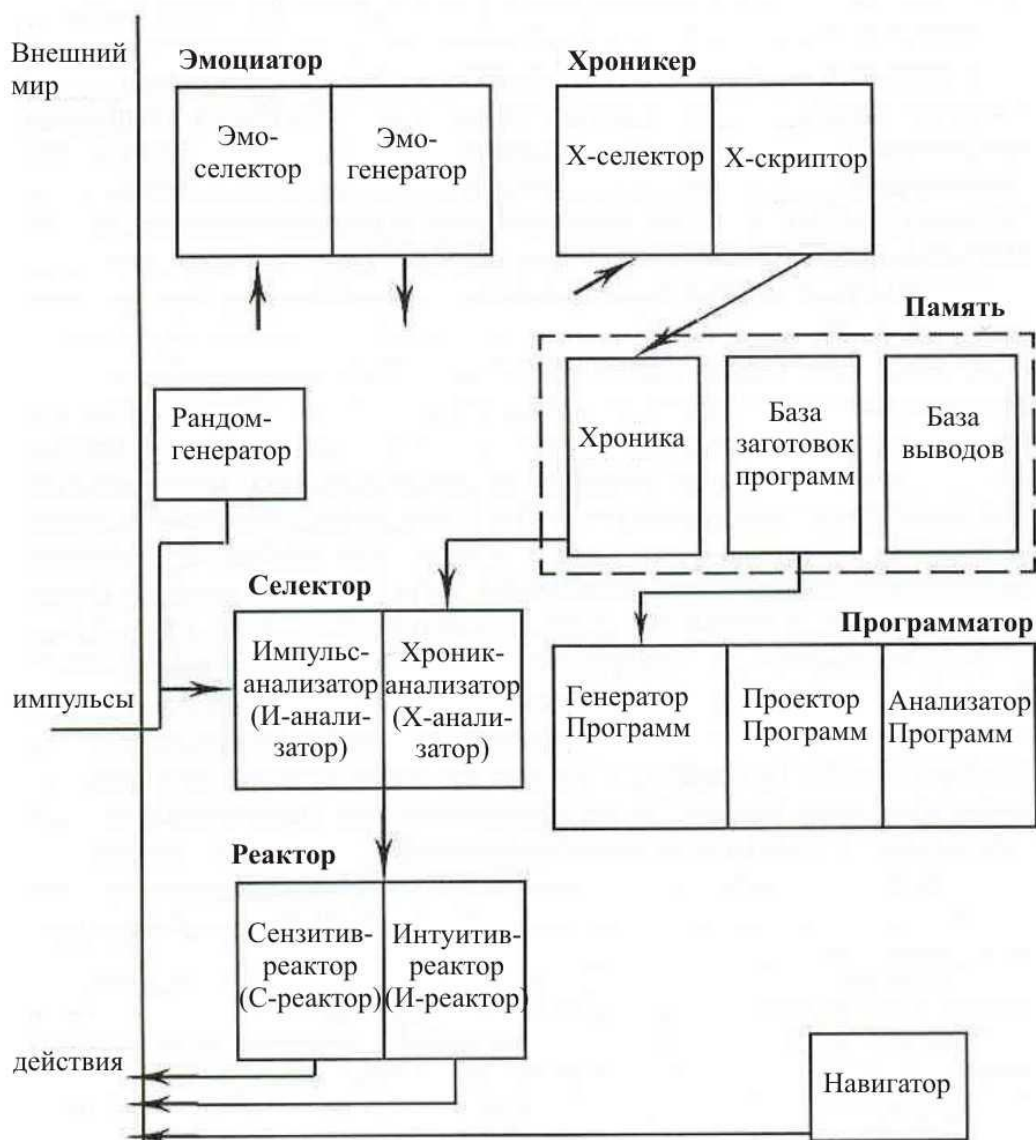


Рис.8. Суммарная схема первого этапа проектирования Доллоса (ОС Доллии)

Приведенные в этой схеме (и введенные в самом документе) понятия дают фундамент для «конструкторского» разговора и рассуждения о различных вещах, связанных с интеллектом человека. Они вытесняют или оттесняют множество «традиционных» (но неконструкторских) понятий, таких как «сознание», «подсознание», «эго», «личность» и др.

Однако здесь только их определения, а конкретная их «работа» должна проявляться в конкретных разговорах и рассуждениях на конкретные темы по конкретным вопросам – если таковые будут заданы читателями или соавторами ПОТИ.

§39. «Получай, фашист, гранату...»

Итак, Сергей, ты перевел с латышского языка и законспектировал некоторый документ. Что это был за документ? Ты его обозначал как «Параграф такой-то книги В. Эгле L-ARTINT». Но книга ARTINT – это сборник разнородных материалов из разных книг, который я скомпоновал в марте 2009 года в ответ на просьбу докторантки (по-советски: аспирантки) Рижского технического университета Даце Апшвалки предоставить ей еще какие-нибудь материалы по Веданской теории помимо тех, что в тот момент имелись в Интернете.

Тот документ, который ты законспектировал, на самом деле было написанное в июне 1999 года мое первое послание Иманту Фрейбергу, профессору информатики Университета Квебека в Монреале, который за несколько месяцев до того покинул Канаду и вернулся на родину в Латвию, из которой он уехал 10-летним ребенком в 1944 году. (Потом были еще и другие

послания как ему, так и его жене – профессору психологии Университета Монреаля, тоже в это же время вернувшейся в Латвию).

История с ними описана в §41 книги {РОТИ-1 = МОИ № 41}, и нет нужды её здесь повторять. Здесь я отмечу только два ключевых момента этой истории:

1) Я послал Фрейбергу документ, который сейчас перед тобой и перед читателями, и просил Фрейберга в меру его сил поспособствовать тому, чтобы Веданская теория была рассмотрена Латвийской наукой.

2) Фрейберг обратился в Полицию безопасности, чтобы меня арестовали и посадили в тюрьму.

ВСЁ.

Только эти два ключевых момента имеют значение. Остальное всё – фон.

Теперь я спрашиваю тебя – и всех читателей: ЧТО было такое в моем послании и в моей просьбе, чтобы со мной поступить ТАК?

Почему нельзя было поступить по-человечески: ответить, что-то сказать, не доводить дело до эксцессов?

Праздные вопросы: спрашивать надо Фрейберга, но он никогда не давал никаких пояснений о своем поведении.

Это поведение Фрейберга я квалифицирую как паранойяльное.

Но Фрейберг вел себя так НЕ потому, что он – психически больной параноик.

Он вел себя так потому, что он – расист.

Я для него являлся представителем низшей расы, с которым нечего разговаривать, нечего опускаться до того, чтобы ему отвечать, нечего с ним соблюдать общечеловеческие нормы поведения, потому что он же – не человек!

Расизм в том и состоит и заключается, что по какому-то признаку выделяется группа «нечеловеков», и в отношении представителя этой группы полностью игнорируются его личные качества. Не важно – умён он или дурак, честен или вор; не важно, что он там написал: логично это, нелогично – какое это имеет значение?! Он же – не человек! Вот если бы он был НАШИМ, то есть человеком – ну, тогда, конечно, совсем другое дело! Тогда Фрейберг расплылся бы в улыбках...

А я ненавижу расизм.

Тем более, когда он обращен против меня.

В конце того §41 {РОТИ-1 = МОИ № 41} я сказал, что я сделаю «с некоторыми профессорами» и «академиками» Латвии».

И вот, сейчас я это сделаю с Фрейбергом.

«Получай, фашист, гранату!»...

У меня была уже подготовлена «Доска Позора» в отдельном файле, Фрейберг уже висел на ней вниз головой, но я всё еще сомневался: переносить это в книгу РОТИ-3 – или не переносить?

Не потому сомневался, что мне его стало жалко.

Мне не жалко этого канадско-латышского подонка.

Но я думал: «Может быть не стоит загрязнять Переписку ПОТИ его историей?»

Я сомневался, и в конце концов решил бросить жребий в виде односантимной монетки: как Судьба решит, так и будет! А я снимаю с себя ответственность.

И боги Олимпа решили:

ВЫСТАВИТЬ ЕГО НА ПОЗОР!

Что ж, я подчиняюсь решению богов.

Justitia in suo cuique tribuendo cernitur – говорил Марк Туллий Цицерон: Справедливость состоит в том, чтобы воздать каждому по заслугам.

Так установим же справедливость и воздадим по заслугам этому псевдоученому. Пусть он висит вниз головой в раскаленной траурной рамке перед всей Россией и перед всем миром – и висит вечно, как папа римский у Данте в аду.

ДОСКА ПОЗОРА



Имант Фрейберг

– профессор информатики
Университета Квебека в Монреале,
почетный доктор
Латвийской Академии Наук.

**Этот человек опозорил
Канадскую и Латвийскую науку**

§40. Заключение

2011.08.22 00:16 ночь на понедельник

И еще раз – Итак, Сергей!

Наша книга достигла 100-й страницы, то есть того предела, у которого её полагается закрывать. Одновременно она сейчас имеет и некоторую логическую завершенность: сюда полностью перенесен и дополнен комментариями, вопросами и ответами документ, написанный 12 лет назад, хоть и недостойному адресату, но сам текст-то достойный.

Добавление в этот том каких-либо новых материалов только испортило бы его. Поэтому я закрываю книгу РОТИ-3. Она готова.

Если ты мне теперь напишешь, то это будет означать, что ты открываешь новую книгу РОТИ-4 в соавторстве со мной. Твоя воля – продолжить или нет.

С закрытием книги РОТИ-3 закрывается и первая трилогия Переписки ПОТИ. «Сакральное» число – три! Три человека побывали у меня в соавторах:

- Николай Шуйкин,
- Дмитрий Манин и
- Сергей Марьясов.

Я могу сказать, что только один из них показал себя ученым. Это был Сергей Марьясов.

Николая Шуйкина и Дмитрия Манина я, к сожалению, признать учеными не могу.

Главная, первейшая, высочайшая заповедь ученого (настоящего ученого, а не карьериста от науки!) – это **ВЫЯСНИТЬ ИСТИНУ**.

Только Сергей Марьясов пытался ВЫЯСНИТЬ ИСТИНУ.

У Николая Шуйкина и Дмитрия Манина такого желания не было.

Речь не идет о том, признать или не признать справедливой Веданскую теорию. Речь идет о желании выяснить, верна она или нет.

Ну, на Николая Шуйкина и сердиться-то невозможно – что с него возьмешь?

А вот Дмитрий Манин в принципе обладает интеллектуальным потенциалом, который мог бы его сделать ученым. Но выступление его в этой Переписке было чрезвычайно слабым. Подвели высокомерие, капризность и неряшливость.

Эта Переписка по-прежнему адресуется Комиссии РАН по борьбе с лженаукой, как это указано на титульном листе каждого тома. Я не посылаю Комиссии ни эти тома, ни интернетовские ссылки на них, а просто выставляю книги в Интернет. У меня нет достоверной информации, нашли ли члены Комиссии их там, или нет. Но думаю, что уже нашли.

Пока я ничего не требую от Комиссии РАН. Но придет время, когда я потребую от них ответа на некоторые вопросы.

Валдис Эгле

22 августа 2011 года

P.S. 2011.08.29: Реакция Сергея Марьясова на §39 настоящего тома отображается в §1 книги {РОТИ-4 = **МОИ** [№ 44](#)}.

Данная книга является протоколом реального общения авторов, т.е. «хроникой текущих событий», хранилищем писем, записок, файлов и т.д., и поэтому мнение авторов по обсуждаемым в книге вопросам может не совпадать, и каждый из авторов единолично несёт ответственность за свои высказывания, утверждения, идеи и размещаемые в своих текстах (письмах) материалы.

Научно-популярное издание
«Мысли об Истине»
Выпуск № 43
Сформирован 10 сентября 2017 года

Все читатели приглашаются принять участие в создании альманаха МОИ и присылать свои статьи и заметки для этого издания по адресу: Marina.Olegovna@gmail.com. Если присланные материалы будут соответствовать направлению Альманаха и минимальным требованиям информативности и корректности, то они будут опубликованы в нашем издании.

Основной вид существования Альманаха МОИ – в виде PDF-файлов в Вашем компьютере. Держите все выпуски МОИ в одной папке. Скачать PDF-ы можно с разных мест в Интернете, и не важно, откуда номер скачан. В Интернете нет одной фиксированной резиденции МОИ.

Содержание

Эгле В. и Марьясов С. Переписка о Теории интеллекта.....	2
Третья переписка в ПОТИ.....	3
Глава 1. Как завязалась эта переписка.....	3
§1. Первые письма.....	3
Глава 2. Проект искусственного интеллекта.....	13
§2. Первая постановка задачи.....	13
§3. Письмо Сергея Марьясова от 26 апреля 2011 г.....	14
§4. Ответ.....	15
Глава 3. Американские философы.....	16
§5. Передовица на сайте.....	16
§6. Джон Сирл. «Разум мозга – компьютерная программа?».....	17
§7. Черчленды. «Искусственный интеллект: Может ли машина мыслить?».....	29
§8. Продолжение ответа.....	40
§9. Письма Сергея Марьясова от 2 и 17 мая 2011 г.....	40
§10. Постулат 1 – Естественный отбор.....	43
§11. Постулат 2 – обработка информации.....	44
§12. Самопрограммирование (опять незаконченное).....	45
Глава 4. Конспект письма профессору Фрейбергу.....	45
§13. Письмо Сергея Марьясова от 24 мая 2011 г. (Хроникер).....	45
§14. Проектор Хроникера.....	50
§15. О памяти.....	51
§16. Письма Сергея Марьясова от 31 мая и 7 июня 2011 года (Gsvano).....	52
§17. Письмо Сергея Марьясова от 10 июня 2011 года (Эмоциатор).....	55
§18. Письмо Сергея Марьясова от 20 июня 2011 года (обобщение).....	59
§19. Письмо Сергея Марьясова от 24 июня 2011 года (Программатор).....	60
§20. Интенсивность процессов.....	63
§21. Краткие письма.....	65
§22. Эмоциатор.....	66
§23. Программатор.....	67
§24. Дерево программы.....	69
§25. Письмо Сергея Марьясова от 18 июля 2011 года (еще раз Эмоциатор).....	71
§26. Потоки сигналов.....	73
§27. От вещи к слову!.....	74
§28. Письмо Сергея Марьясова от 2 августа 2011 года (Вариатор).....	75
§29. О сновидениях.....	80
§30. Письмо Сергея Марьясова от 16 августа 2011 года.....	81
§31. Присоединенные файлы к письму от 16 августа (Селектор, Реактор и Навигатор).....	83
§32. Некоторые особенности Рандомгенератора.....	87
§33. Некоторые особенности низшего самопрограммирования.....	89
§34. Как строятся номиналии объектов и их пространственное восприятие?.....	91
§35. Пространственное восприятие.....	91
§36. Откуда берутся алгоритмы.....	93
§37. Вводная и заключительная части Конспекта.....	95
§38. Общая схема Доллоса.....	98
§39. «Получай, фашист, гранату...».....	99
§40. Заключение.....	101
Содержание.....	103