

УДК 111 + 165 + 81`3

DOI: 10.18384/2949-5148-2024-1-33-433

ИНДИВИДУАЛЬНОЕ ПОНИМАНИЕ ВЕЩЕЙ: РАСПРЕДЕЛЁННЫЕ ЗНАЧЕНИЯ, ЭКЗЕМПЛЯРЫ И УНИЦЕПТЫ

Сериков А. Е.*Самарский национальный исследовательский университет имени академика С. П. Королёва
443086, г. Самара, ул. Московское шоссе, д. 34, Российская Федерация*

Аннотация

Цель. Исследование формирования и существования индивидуальных представлений людей о вещах и их связи с языковыми категориями.

Процедура и методы. Проанализированы сведения об исследованиях активности отдельных нейронов мозга, идеи теории когнитива и экзemplарных теорий категоризации, философская концепция уницептов.

Результаты. Научные данные свидетельствуют об уникальности представлений каждого человека о вещах. Отдельным элементам субъективного опыта каждого человека соответствуют распределённые гиперсети нейронов. Языковые категории опираются на субъективный опыт взаимодействия с конкретными экзemplарами. Коммуникация осуществляется не на основе всеобщих концептов, а на основе уникальных индивидуальных представлений – уницептов.

Теоретическая и/или практическая значимость. Впервые в обобщённом виде формулируется идея распределённого кодирования значений, объединяющая то, как понимается кодирование значений вещей в нейропсихологии, в теориях категоризации, в концепции уницептов.

Ключевые слова: концепт, прототип, распределённое кодирование, семантика, уницепт, экзemplар

INDIVIDUAL UNDERSTANDING OF THINGS: DISTRIBUTED MEANINGS, EXEMPLARS, AND UNICEPTS

A. Serikov*Samara National Research University,
ul. Moskovskoe shosse 34, Samara 443086, Russian Federation*

Abstract

Aim. To explore the formation and existence of humans' individual representations of things and their connection to linguistic categories.

Methodology. The research incorporates analysis of data on the activity of individual brain neurons, ideas from the cognitome theory and exemplar theories of categorization, as well as the philosophical concept of unicepts.

Results. Scientific evidence suggests the uniqueness of each individual's representations of objects. Distributed hypernetworks of neurons correspond to specific elements of subjective experience for each person. Linguistic categories are based on subjective experience of interacting with specific exemplars. Communication is based not on universal concepts, but rather on unique individual representations – unicepts.

Research implications. For the first time, this study presents, in a generalized form, the idea of distributed coding of meaning, which combines the understanding of meaning encoding in neuropsychology, categorization theories, and the notion of unicepts.

Keywords: concept, prototype, distributed coding, semantics, unicept, exemplar

Введение

Цель статьи состоит в исследовании вопросов о том, как формируются и существуют индивидуальные представления людей о вещах, как они связаны с языковыми категориями. В 1892 г. в классической статье «Смысл и денотат» («Смысл и значение») Готлоб Фреге предложил чётко различать имя или знак вещи (Name, Zeichen), с одной стороны, саму вещь как денотат или значение этого имени (Bedeutung), с другой, и смысл имени или понятие вещи (Sinn), с третьей. Уточняя, что такое смысл, Фреге проводит ещё одну важную границу: между всеобщим объективным смыслом вещи и её индивидуальным представлением, её субъективным внутренним образом. При этом ни сам Фреге, ни многочисленные его последователи в области семиотики особо не интересовались индивидуальными представлениями как таковыми, концентрируя своё внимание на объективных смыслах, в реальности которых никто из философов и учёных вплоть до недавнего времени не сомневался.

Однако в 2005 г. Рут Гаррет Милликэн публично заявила, что «Фреге ошибся, постулируя нечто общее за пределами *Bedeutung*, которое схватывается умом каждого компетентного говорящего, использующего одну и ту же недвусмысленную языковую форму» [7, р. 66]. Согласно Милликэн, всеобщих фрегеанских смыслов не существует, а существуют лишь сами вещи, существуют естественные языковые конвенции, не требующие всеобщего единого понимания языка, и существуют уникальные индивидуальные представления (понимания, реакции), названные ей относительно недавно *уницептами* [12].

Если Милликэн в чём-то права и смыслы вещей не столь всеобщы, как мы привыкли считать, вопросы об индивидуальных представлениях вещей приобретает особую актуальность. Как формируются и каким образом существуют эти представления с точки зрения современной науки? Одна и та же вещь в представлении разных людей – это одна и та же вещь или разные вещи?

Если каждую вещь мы понимаем по-своему, как возможно взаимное понимание людей, разговаривающих о вещах? Как кодируются индивидуальные представления о вещах? Не рождается ли в ходе обсуждения этих вопросов ещё одна относительно новая идея о том, что такое значение?

Когнитом и нейронное кодирование когнитивной информации

В научно-популярной статье «Нейроны бабушки» нейрофизиологи Родриго Киан Кирога, Ицхак Фрид и Кристоф Кох вспоминают, что в 1969 г. Джерри Леттвин выдумал историю о гениальном русском нейрохирурге Акакии Акакиевиче, который по просьбе пациента удалил у него нейроны, содержащие память о назойливой матери, а затем стал искать нейроны, в которых закодирована память о бабушке пациента, чтобы удалить и их. С тех пор в течение 40 лет существование «нейронов бабушки» часто обсуждалось в нейронауке как гипотетическая непроверяемая возможность, но в последнее время ситуация изменилась. Выяснилось, что действительно существуют отдельные нейроны (точнее, группы нейронов), которые активны в тот момент, когда человек воспринимает некий объект, его видео или фото, его звучащее или написанное наименование. В частности, у одного из пациентов был обнаружен «нейрон Дженнифер Энистон», который активизировался только при предъявлении фото этой актрисы, но не фото других актрис и актёров, знаменитостей, каких-либо мест или животных. У ещё одного пациента был обнаружен нейрон, активизировавшийся только при предъявлении фото или написанного имени актрисы Хэлли Берри. Другой нейрон реагировал исключительно на фото и имя (написанное или произнесённое) телеведущей Опры Уинфри, ещё один нейрон – только на изображения или имя Люка Скайуокера¹. Эти и некоторые другие от-

¹ Quiroga R. Q., Fried I., Koch C. Brain Cells for Grandmother // Scientific American. 2013. Vol. 308. № 2. P. 30–35.

крытия были сделаны в результате применения Фридом и его коллегами методики исследования активности отдельных нейронов мозга путём внедрения в них микро-электродов [17].

Опираясь на подобные эмпирические факты, Константин Владимирович Анохин развивает новый подход к построению теории сознания, одна из основных идей которого состоит в том, чтобы «рассматривать мозг не как коннектом – нейронную сеть, а как когнитом – нейронную гиперсеть, состоящую из нейронных групп со специфическими когнитивными свойствами» [1, с. 39]. Группы, о которых идёт речь, – это сети нейронов, которые разряжаются одновременно, когда активизируется тот или иной элемент опыта, знаний, памяти. Они называются *когами*. Коги образуют устойчивые связи, в основе которых лежат общие нейроны, одновременно входящие в разные коги. Коги и их связи совместно образуют глубокую (многослойную) нейронную гиперсеть, называемую *когнитомом* и понимаемую как нейронный коррелят разума.

В 2022 г. в докладе на междисциплинарном семинаре в МГУ Анохин излагал свои идеи в следующих терминах: «Разум обладает зернистой структурой, он состоит из когнитивных элементов – *когов*. Коги опосредуют каузальные (информационные) соотношения целостного когнитивного агента с его средой. Коги образуют между собой устойчивые когнитивные связи – *локи*. Локи отражают причинные связи элементов и процессов в среде и в соотношениях когнитивного агента с ней. Коги и локи образуют единую когнитивную сеть – *когнитом*. Когнитом является носителем всего субъективного опыта когнитивного агента»¹. Тот или иной ког может быть связан с образом определённого объекта или события, определённой ситуа-

ции, и его активация в некоторый момент времени будет передаваться другим когам, в которых закодирована возможная реакция на данный объект или событие в следующий момент времени. «Иначе говоря, если активирован ког “собака”, это означает, что во внешнем мире существует некий объект, существо, которое соответствует этому когнитивному элементу. Это элемент знания о том, “что” в мире для этого когнитивного агента. Но эта же активация этой когнитивной группы несёт знание также и то, “как”, т. е. не только причину, но и эффект»².

И терминология, и отдельные тезисы данной теории постоянно развиваются и уточняются, но одна из идей сохраняется: языковые категории и их значения в идиолекте каждого конкретного человека могут быть поняты как распределённо закодированные сетью нейронов его головного мозга. Самого Анохина прежде всего интересует проблема соотношения разума и мозга; о проблеме того, что такое значение, он специально не говорит. Но, решая проблему «разум / мозг», Анохин буквально задаёт вопросы о нейронном коде: «Как мозг кодирует ментальные феномены?» Поэтому представляется, что на периферии теории когнитивного агента возникает относительно новое (для философии, лингвистики, семантики в целом) понимание значения языковых категорий. Его можно назвать *распределённо закодированным значением*.

Сам этот термин нельзя считать совершенно новым. Один из математических подходов к обработке естественного языка (natural language processing) состоит в том, чтобы кодировать значения слов векторами (наборами координат) в многомерном пространстве. Модели, использующие этот принцип, называют *распределёнными семантическими моделями* (Distributional semantic models – DSMs). Согласно Рэндаллу Джемисону, Джонатану Эйвери, Брендану Джонсу и Микаэлю Джоунсу, эти модели используются не просто в целях

¹ Анохин К. В. Гиперсетевой мозг: новое решение mind-brain problem: доклад на междисциплинарном семинаре факультета психологии МГУ им. М. В. Ломоносова «Время. Субъект. Сознание. Деятельность» 11.05.2022 // YouTube: [сайт]. URL: <https://www.youtube.com/watch?v=k8IIRydqKu4> (дата обращения: 30.05.2023).

² Там же.

компьютерного перевода или поиска информации, но для того, чтобы «объяснить, как люди преобразуют статистический языковой опыт первого порядка в глубокое знание репрезентаций значений слов» [4, p. 119]. Перспективы развития DSMs данные авторы связывают с применением в них *экземплярного* подхода к пониманию значения в естественных языках.

Экземплярный подход к пониманию значений

В «Трактате о принципах человеческого знания» (1710) Джордж Беркли критически оценивал суждение Джона Локка о том, что имеющие «общий характер» слова являются знаками общих абстрактных идей. Согласно Беркли, невозможно вообразить «треугольник вообще» или «человека вообще», но только конкретные разновидности треугольника или типы человека, представления о которых и обобщаются соответствующими словами. Поэтому Беркли считает, что общих абстрактных идей не существует, а соответствующие слова являются знаками множеств, объединяющих схожие частные идеи.

В 1951 г. в статье «Две догмы эмпиризма» Уиллард Куайн показал, что кантианское различие между аналитическими и синтетическими суждениями, с одной стороны, и вера эмпиристов в возможность сведения эмпирических понятий исключительно к данным непосредственного опыта, с другой стороны, тесно взаимосвязаны. В этом контексте можно предположить, что Беркли заблуждался, когда полагал, что наши абстрактные идеи опираются исключительно на конкретные образы и образцы, но, с другой стороны, заблуждаются и те, кто полагает, что можно аналитически определить некое понятие, не опираясь ни на какой опыт и ни на какие образцы вообще.

Вопрос Беркли о представлении «треугольника вообще», «человека вообще» и сомнение Куайна в возможности чёткого различения между аналитическими и синтетическими суждениями можно пе-

реформулировать на языке современной нейрофизиологии следующим образом: существуют ли соответствующие общим абстрактным понятиям (чистым аналитическим понятиям) группы нейронов, при активизации которых не активируются нейроны конкретных понятий (нейроны эмпирических образцов)? Например, существуют ли *нейроны треугольника вообще*, *нейроны человека вообще* по аналогии с тем, как существуют *нейроны бабушки*? И могут ли нейроны треугольника вообще активизироваться без активизации нейронов равнобедренного треугольника, могут ли нейроны человека вообще активизироваться без активизации нейронов бабушки? В такой форме эти вопросы рано или поздно могут быть проверены эмпирически.

Идея о том, что наше понимание категорий естественного языка (а в более радикальной версии – и понимание научных понятий) с необходимостью опирается на знание конкретных образцов, во второй половине XX в. получила название *экземплярного подхода* к проблеме категоризации. В когнитивной психологии этот подход обычно противопоставлялся так называемому *прототипическому подходу*, в то время как оба этих подхода вместе взятые противопоставлялись *классическому логическому подходу* к определению понятий. Оба термина – «экземплярный» и «прототипический» – являются многозначными, поэтому на уровне ярлыков в этой области царит путаница, в которой следует разобраться.

В середине XX в. термин «прототип» употреблялся в когнитивной психологии в качестве синонима «схемы», понимаемой как статистическая средняя (центральная тенденция) конкретных образцов, а сами образцы – как отклонения от этой средней (вариации). В 1968 г. Майкл Познер и Стивен Кили опубликовали статью «О происхождении абстрактных идей», где сформулировали три альтернативные гипотезы о том, как может происходить запоминание и последующее распознавание образцов: 1) субъекты в ходе процесса

абстрагирования формируют и запоминают прототипическую схему, лежащую в основе образцов; 2) субъекты запоминают и прототипическую схему, и образцы; 3) субъекты запоминают только образцы. В своём исследовании Познер и Кили использовали в качестве прототипа / схемы бессмысленные паттерны из точек, а образцами были искажения этих паттернов на основании определённых статистических правил. Результаты исследования, в истолковании авторов, максимально согласовывались с гипотезой № 2, т. е. с одной из прототипических гипотез, отличных от экземплирной гипотезы № 3 [16].

Экземплирный подход зародился в рамках исследований по дифференцировочному обучению животных (*discrimination learning*, научение реакции различения). В 1975 г. в работе «Теория контекста в дифференцировочном обучении» Дуглас Медин сравнивает *целостные* и *компонентные* модели стимулов, воздействующих на обучаемое животное. Например, с точки зрения целостной модели, крыса в лабиринте воспринимает белый треугольник слева как единый стимул, а с точки зрения компонентной модели, она реагирует независимо на белый цвет, форму треугольника, левое расположение. Компонентные модели обладали рядом преимуществ, но обычно опирались на предположение, что отдельные компоненты стимула воздействуют на животное независимо. Идея Медина состоит в том, чтобы рассматривать отдельные компоненты как зависимые друг от друга; тогда один из компонентов может быть понят как *контекст* для другого [10].

В статье 1978 г. «Контекстная теория обучения классификации», написанной Медином совместно с Маргерит Шаффер, эта идея используется при объяснении того, как люди выполняют задачи категоризации (классификации). Гипотеза заключается в том, что люди распознают предметы по признакам, не являющимися независимыми. Значение одного из параметров, на основании которого предъявляемый предмет относится к некоторой катего-

рии, взаимодействует со значением других (контекстных) параметров. Контекстная теория противопоставляется авторами классическому логическому подходу к «хорошо определённым понятиям», т. е. таким, которые определены путём перечня достаточных и необходимых существенных признаков. Ссылаясь на исследования Элеанор Рош с соавторами, обнаружившими в начале 1970-х гг. эффекты базового уровня абстракции и прототипические эффекты в структуре языковых категорий, Медин и Шаффер пишут, что «большинство естественных категорий не имеют хорошо определённых правил или фиксированных границ, разделяющих альтернативные категории. Скорее, члены различаются степенью, в которой они признаны хорошими (типичными) примерами данной категории, и многие естественные концепты не могут быть определены путём указания на множество критических признаков» [11, р. 208]. Отталкиваясь от этого эмпирического факта, Медин и Шаффер формулируют необходимость объяснения того, как происходит изучение и использование «плохо определённых» категорий. При этом они подвергают сомнению обоснованность выводов Познера и Кили. «Познер и Кили (1968) предлагали, что на основе опыта знакомства с экземплярами, у людей формируется впечатление о центральной тенденции некой категории, и что категориальные суждения должны быть основаны на этой центральной тенденции, или прототипе. Поскольку нет универсального согласия, что формирование прототипа лежит в основании концептуального обучения в этой области, появление всё больших доказательств того, что естественные категории и концепты не являются хорошо определёнными, привело к возрастанию интереса в развитии теорий концептуального поведения по правилам, допускающим исключения» [11, р. 208]. Основная идея предлагаемой Медином и Шаффер контекстной теории соответствует гипотезе № 3 Познера и Кили и заключается в том, что «классификационные суждения основаны на извлечении сохра-

нённой в памяти информации об экземплярах» [11, р. 208–209]. Позже такой подход стали называть *экземплярной точкой зрения*.

Здесь наметилось видимое противоречие в использовании термина «прототипический подход». С одной стороны, этот термин обозначал подходы на основе идеи прототипической схемы, понимаемой как центральная тенденция. С другой стороны, он начал обозначать подходы, основанные на открытых Рош прототипических эффектах, которые можно было объяснить, используя как идею центральной тенденции, так и идею экземпляров. Дело в том, что и сама Рош в качестве автора теории прототипов критиковала идею центральной тенденции. Борис Митрофанович Величковский приводит (в его переводе на русский язык) следующую цитату из статьи Рош 1973 г. «О внутренней структуре перцептивных и семантических категорий»: «Ни модель формирования понятий в терминах заучивания “правильной” комбинации дискретных атрибутов, ни модель процесса абстракции в терминах абстрагирования центральной тенденции ... некоторого произвольного сочетания признаков не являются адекватными объяснениями природы и развития естественных категорий ... Предлагается ... следующая альтернатива: существуют ... формы, которые перцептивно более заметны, чем все другие стимулы в данной области ... эти наиболее заметные формы являются “хорошими формами” гештальтпсихологии»¹.

В приведённой выше цитате из статьи Медина и Шаффер видно, что эти авторы опираются на открытые Рош прототипические эффекты, но понимают прототипы как экземпляры и противопоставляют свою теорию экземпляров теориям, основанным на понимании прототипов как схемы и центральной тенденции. В этом контексте думается, что полезным было бы терминологическое различие

прототипов-схем и прототипов-экземпляров.

В 1981 г. Медина и Эдвард Смит выпустили книгу «Категории и понятия», в которой различили три теоретические точки зрения на структуру категорий: *классическую, вероятностную*, примером которой является теория прототипов, и их собственную – *экземплярную*. Согласно вероятностной точке зрения (the probabilistic view) принадлежность объекта к классу определяется схожестью с прототипом. Непрототипические объекты относятся к категории лишь с некоторой вероятностью. Один из вариантов этой точки зрения называется подходом, основанном на признаках (featural): категории представляются в терминах признаков, значимых в разной степени и с разной вероятностью; объект принадлежит к категории, если обладает неким критическим количеством признаков с учётом их вероятности; прототип – это не реальный объект, а абстрактное обобщение наиболее вероятных признаков разных объектов. Другой вариант вероятностной точки зрения – это подход, основанный на *размерностях* (dimensional): значимые признаки понимаются как переменные, формирующие многомерное пространство, в котором прототипы занимают определённое место, а объект принадлежит к категории в зависимости от близости к этому месту. Третий вариант данной точки зрения называется *целостным* (holistic) подходом: вероятность принадлежности объектов к классу зависит от схожести с прототипами, которые понимаются как целостные образы конкретных прототипических представителей. Суть *экземплярной точки зрения* (the exemplar view) заключается в том, что категории в обыденном восприятии связаны не столько с какими-либо абстрактными обобщениями или образами, сколько с имеющимися знаниями о конкретных представителях или подвидах данной категории. «Представление о понятии состоит из отдельных описаний некоторых его экземпляров (или конкретных объектов, или подмножеств)» [18, р. 144].

¹ Величковский Б. М. Когнитивная наука: основы психологии познания: учебное пособие: в 2 т. Т. 2. М.: Смысл, 2006. С. 35.

Обобщающая информация включается в знания о категориях, но играет существенно меньшую роль, чем это предполагается вероятностной точкой зрения. Точка зрения экземпляров позволяет достаточно просто объяснить, почему большинству испытуемых с трудом удаётся определять категории путём перечисления необходимых и достаточных признаков. Она объясняет существование дизъюнктивных категорий, членство в которых для одних объектов связано с одними признаками, а для других – с другими. Например, имя существительное определяется как часть речи, отвечающая на вопросы «Кто?» или «Что?» и обозначающая или вещи, или вещества, или живые существа, или явления и события, или признаки.

Есть нечто общее между экземплярным подходом к категоризации, с одной стороны, и теорией когнитива, с другой. Согласно теории экземпляров, каждая категория (концепт) связана с множеством экземпляров. С точки зрения теории когнитива, любой конкретный эпизод индивидуального опыта, любой концепт и каждое из известных индивиду его значений, каждый из известных индивиду случаев употребления слова – всё это может быть рассмотрено как знание об экземплярах, связанных с тем или иным концептом и закодированное в соответствующих когах и локах. Если тезисы этих теорий свести воедино, то получается следующее: значение абстрактных категорий в идиолекте закодировано множеством взаимосвязанных нейронов в мозге данного индивида, и это множество должно быть связано с несколькими другими множествами нейронов, которые кодируют память об экземплярах данной категории, воспринимаемых индивидом в качестве её прототипов-экземпляров или в качестве основы для прототипов-схем.

Сегодня споры по поводу экземплярного подхода продолжаются, т. е. они по-прежнему актуальны. Один из последних примеров – тематический выпуск в октябре-декабре 2020 г. сдвоенного 40-го тома журнала “First Language”, полностью по-

свящённого обсуждению экзemplарной модели усвоения языка. В начале тома опубликована проблемная статья Бена Эмбриджа «Против сохраняемых абстракций: Радикальная экзemplарная модель усвоения языка» [6], затем 18 статей различных авторов, критически анализирующих и комментирующих тезисы Эмбриджа, а в конце тома – его ответная статья «Абстракции, сделанные из экзemplаров или “Вы правы, а я изменил свою точку зрения”: Ответ комментаторам» [5]. Итоги дискуссии видны из названия заключительной статьи: начав с радикального тезиса о том, что при усвоении языка мы запоминаем лишь конкретные высказывания, их значения и контекстуальные детали, и не сохраняем в своей памяти абстрактные лингвистические формы, Эмбридж под давлением критических аргументов пришёл к выводу, что хотя мы и сохраняем в памяти конкретные экзemplары, в процессе их использования мы также заново репрезентируем (re-represent) их на различных уровнях абстракции.

В ходе подобных споров выясняется, в частности, что следует различать *радикальный* и *умеренный* варианты экзemplарного подхода. Радикальный вариант отрицает роль абстракций при обучении и выполнении задач категоризации полностью, умеренный вариант допускает определённую роль абстракций, но также настаивает на необходимости существования конкретных образцов. При этом можно по-разному понимать, что такое экзemplар. Например, в умеренных подходах абстракции нижнего уровня понимаются как экзemplары по отношению к абстракциям верхнего уровня. В контексте вопроса о распределённом кодировании значений концептов можно высказать следующее дополнительное соображение: нельзя понимать экзemplары только как образцы каких-то конкретных материальных вещей, ведь существуют также образцы теоретических объяснений, примеры определений и связанных с ними вычленяемых существенных признаков предметов, образцы классификаций из

учебников, конкретные примеры употребления слов и языковых выражений.

Естественные конвенции и уницепты

В конце 1960-х гг. Дэвид Льюис разработал теорию языковых конвенций, ставшую основой для всех последующих дискуссий в этой области [7; 8]. Согласно этой теории, конвенции – это альтернативные варианты использования языка, существующие на основе почти всеобщей правдивости говорящих, почти всеобщем доверии слушающих, почти всеобщем следовании конвенции, а также рациональном понимании и предпочтении абсолютного большинства членов популяции того, что другие члены популяции предпочитают следовать данной конвенции, что позволяет эффективно решать проблемы координации. В 1998 г. Милликэн отбрасывает большинство этих предположений и предлагает понятие *естественных* языковых конвенций. «Естественная конвенциональность складывается из двух довольно простых взаимосвязанных характеристик. Во-первых, естественные конвенции состоят из паттернов, которые “воспроизводятся” ... Во-вторых, тот факт, что эти паттерны распространяются, основан частично на значении прецедентов, а не на свойственной им, к примеру, превосходящей способности выполнять определённую функцию» [13, p. 162].

Идеи Милликэн в своей основе очень просты, лучше согласованы с эмпирическими наблюдениями и позволяют предлагать теоретические объяснения, невозможные на основе традиционных представлений [2; 3; 9]. Как и конвенции в теории Льюиса, естественные конвенции Милликэн имеют место там, где возможны альтернативные варианты. Но выбор этих вариантов совсем не обязательно должен строиться на рациональных предположениях большинства членов популяции. Достаточно простого подражания прецедентам и того, чтобы основанное на нём общение было результативно хотя бы иногда, т. е. время от времени приносило бы

пользу. Например, чтобы один человек общался другому правдивую информацию, помогающую обоим решить их практические проблемы. Или чтобы после просьбы о помощи она была правильно понята, и нужная помощь была действительно оказана. Говоря техническим языком, для того, чтобы языковые конвенции сохранялись, они должны быть функциональны, а для этого при их использовании время от времени должны выполняться условия истинности или обоснованности. Время от времени – значит не всегда и даже не в большинстве случаев, а просто в некоторой доле случаев, достаточной для того, чтобы эти конструкции не забывались и воспроизводились.

Но что особенно важно, такое понимание конвенций хорошо согласуется с теми описанными выше фактами, что представления людей о вещах являются индивидуальными и уникальными. В своей новой книге «За пределами понятий: уницепты, язык и естественная информация» Милликэн предложила для обозначения таких уникальных представлений удачный неологизм – *уницепты* [12].

Если существует одна и та же вещь или множество похожих вещей, у встречавшего такие вещи человека есть для них имя или несколько синонимичных имён, и связанное с этими именами распределённое значение – уницепт. Уницепт – часть идиолекта, уницепты одних и тех же вещей у всех людей разные. Некоторые имена и определения вещей конвенциональны, конвенциональные имена и определения являются элементами уницептов тех людей, которые данным конвенциям следуют. Но даже конвенциональные имена и определения связаны у разных людей с разными экземплярами, и поэтому их уницепты различаются. При этом нет таких конвенций, которым следуют все или даже абсолютное большинство. Кроме того, большинство людей не задумываются об определениях и не помнят их, поэтому многие уницепты не включают в себя определения.

В этом контексте следует уточнить понятие идиолекта. Традиционное понима-

ние заключается в том, что идиолект – это индивидуальная подсистема естественного языка, включающая лексические и грамматические формы, которыми владеет данный индивид. Допустим, индивид знает далеко не все слова и идиомы родного языка, и знаком не со всеми значениями тех слов, которые знает, но те, с которыми он знаком, – это всеобщие значения, т. е. фрегеанские смыслы, зафиксированные в словарях. Однако, если Милликэн права, уницепты разных людей, значения индивидуальных лексиконов не совпадают, а в словарях представлены лишь некоторые конвенциональные значения, эксплицированные и сформулированные лексикографами.

Заключение

Итак, существует нечто, что можно назвать индивидуальным распределённым кодированием значения языковых категорий. Близкие по значению языковые категории и кодирующие их коги должны быть связаны общими нейронами, участвующими в обоих когах. Поэтому связанные по значению категории можно представить как закодированные совместно единой гиперсетью. Активация одной из категорий должна активировать ассоциативно связанные с ней другие категории, значе-

ния которых в данном идиолекте не могут существовать отдельно друг от друга, но взаимно определяют друг друга коннотативно. Такой взгляд хорошо согласуется и с экземплярной теорией языковых категорий, и с представлением об уницептах. В сознании вещи и их виды представлены, с одной стороны, как память об их конкретных экземплярах и образах, с другой стороны – как память об экземплярах их абстрактных определений или описаний. В принципе, распределённые значения существуют просто потому, что существуют индивиды с их особым опытом и образцами значений.

Что же такое общие значения слов и связанных с ними вещей? Если они и существуют, то лишь как конвенции, причём такие, которые не являются общеобязательными. Несмотря на реальное существование вещей, каждый из нас воспринимает и понимает эти вещи по-своему. Поэтому, когда мы говорим о вещах или делаем с ними что-то совместно с другими людьми, мы понимаем друг друга лишь частично. Это частичное понимание и есть то, что мы привыкли считать пониманием, другого просто не бывает.

Статья поступила в редакцию 05.12.2023.

ЛИТЕРАТУРА

1. Анохин К. В. Когнитом: в поисках фундаментальной нейронаучной теории сознания // Журнал высшей нервной деятельности. 2021. Т. 71. № 1. С. 39–71.
2. Сериков А. Е. Понятие конвенции Дэвида Льюиса и его современная рецепция // Аспирантский вестник Поволжья. 2023. Т. 23. № 3. С. 43–48.
3. Сериков А. Е. Понимание правил языка в биосемантике Рут Гаррет Милликэн // Семиотические исследования. 2023. Т. 3. № 3. С. 13–21.
4. An Instance Theory of Semantic Memory / R. K. Jamieson, J. E. Avery, B. T. Johns, M. N. Jones // Computational Brain & Behavior. 2018. Vol. 1. No 2. P. 119–136.
5. Ambridge B. Abstractions Made of Exemplars or ‘You’re All Right, and I’ve Changed My Mind’: Response to Commentators // First Language. 2020. Vol. 40. No. 5–6. P. 640–659.
6. Ambridge B. Against Stored Abstractions: A Radical Exemplar Model of Language Acquisition // First Language. 2020. Vol. 40. No 5–6. P. 509–559.
7. Lewis D. Convention: A Philosophical Study. Harvard: Harvard University Press, 1969. 213 p.
8. Lewis D. Languages and Language // Minnesota Studies in the Philosophy of Science / ed. by K. Gunderson. Minneapolis: University of Minnesota Press, 1975. P. 3–35.
9. Matczak M. Ruth G. Millikan’s Conventionalism and Law // Legal Theory. 2022. Vol. 28. No. 2. P. 146–178.
10. Medin D. L. A Theory of Context in Discrimination Learning // The Psychology of Learning and Motivation. Vol. 9 / ed. by G. H. Bower. New York: Academic Press, 1975. P. 263–314.

11. Medin D. L., Schaffer M. M. Context Theory of Classification Learning // *Psychological Review*. 1978. Vol. 85. P. 207–238.
12. Millikan R. G. *Beyond Concepts: Unicepts, Language, and Natural Information*. Oxford: Oxford University Press, 2017. 240 p.
13. Millikan R. G. Language Conventions Made Simple // *The Journal of Philosophy*. 1998. Vol. 95. № 4. P. 161–180.
14. Millikan R. G. On Meaning, Meaning, and Meaning // Millikan R. G. *Language: A Biological Model*. Oxford: Oxford University Press, 2005. P. 53–76.
15. Murphy G. L. Is There an Exemplar Theory of Concepts? // *Psychonomic Bulletin & Review*. 2016. Vol. 23. P. 1035–1042.
16. Posner M. I., Keele S. W. On the Genesis of Abstract Ideas // *Journal of Experimental Psychology*. 1968. Vol. 77. № 3. Pt. 1. P. 353–363.
17. *Single Neuron Studies of the Human Brain: Probing Cognition* / ed. I. Fried, U. Rutishauser, M. Cerf, G. Kreiman. Cambridge, MA: MIT Press, 2014. 382 p.
18. Smith E. E., Medin D. L. *Categories and Concepts*. Cambridge, MA: Harvard University Press, 1981. 203 p.

REFERENCES

1. Anohin K. V. [Cognitome: In Search of Fundamental Neuroscience Theory of Consciousness]. In: *Zhurnal vysshej nervnoj deyatel'nosti* [I.P. Pavlov Journal of Higher Nervous Activity], 2021, vol. 71, no. 1, pp. 39–71.
2. Serikov A. E. [David Lewis's Convention and Its Contemporary Perception]. In: *Aspirantskij vestnik Povolzh'ya* [Postgraduate Bulletin of the Volga Region], 2023, vol. 23, no. 3, pp. 43–48.
3. Serikov A. E. [Understanding Language Rules in Ruth Garrett Millikan's Biosemantics]. In: *Semioticheskie issledovaniya* [Semiotic Studies], 2023, vol. 3, no. 3, pp. 13–21.
4. Jamieson R. K., Avery J. E., Johns B. T., Jones M. N. An Instance Theory of Semantic Memory. In: *Computational Brain & Behavior*, 2018, vol. 1, no 2, pp. 119–136.
5. Ambridge B. Abstractions Made of Exemplars or 'You're All Right, and I've Changed My Mind': Response to Commentators. In: *First Language*, 2020, vol. 40, no. 5–6, pp. 640–659.
6. Ambridge B. Against Stored Abstractions: A Radical Exemplar Model of Language Acquisition. In: *First Language*, 2020, vol. 40, no. 5–6, pp. 509–559.
7. Lewis D. *Convention: A Philosophical Study*. Harvard, Harvard University Press, 1969. 213 p.
8. Lewis D. Languages and Language. In: Gunderson K., ed. *Minnesota Studies in the Philosophy of Science*. Minneapolis, University of Minnesota Press, 1975, pp. 3–35.
9. Matczak M. Ruth G. Millikan's Conventionalism and Law. In: *Legal Theory*, 2022, vol. 28, no. 2, pp. 146–178.
10. Medin D. L. A Theory of Context in Discrimination Learning. In: Bower G. H., ed. *The Psychology of Learning and Motivation*. Vol. 9. New York, Academic Press, 1975, pp. 263–314.
11. Medin D. L., Schaffer M. M. Context Theory of Classification Learning. In: *Psychological Review*, 1978, vol. 85, pp. 207–238.
12. Millikan R. G. *Beyond Concepts: Unicepts, Language, and Natural Information*. Oxford, Oxford University Press, 2017. 240 p.
13. Millikan R. G. Language Conventions Made Simple. In: *The Journal of Philosophy*, 1998, vol. 95, no. 4, pp. 161–180.
14. Millikan R. G. On Meaning, Meaning, and Meaning. In: Millikan R. G. *Language: A Biological Model*. Oxford, Oxford University Press, 2005, pp. 53–76.
15. Murphy G. L. Is There an Exemplar Theory of Concepts? In: *Psychonomic Bulletin & Review*, 2016, vol. 23, pp. 1035–1042.
16. Posner M. I., Keele S. W. On the Genesis of Abstract Ideas. In: *Journal of Experimental Psychology*, 1968, vol. 77, no. 3, pt. 1, pp. 353–363.
17. Fried I., Rutishauser U., Cerf M., Kreiman G., eds. *Single Neuron Studies of the Human Brain: Probing Cognition*. Cambridge, MA, MIT Press, 2014. 382 p.
18. Smith E. E., Medin D. L. *Categories and Concepts*. Cambridge, MA, Harvard University Press, 1981. 203 p.

ИНФОРМАЦИЯ ОБ АВТОРЕ

Сериков Андрей Евгеньевич – кандидат философских наук, доцент кафедры философии Самарского национального исследовательского университета имени академика С. П. Королёва;
e-mail: serikov.ae@ssau.ru

INFORMATION ABOUT THE AUTHOR

Andrew Ye. Serikov – Cand. Sci. (Philosophy), Assoc. Prof., Department of Philosophy, Samara National Research University;
e-mail: serikov.ae@ssau.ru

ПРАВИЛЬНАЯ ССЫЛКА НА СТАТЬЮ

Сериков А. Е. Психофизиологическая проблема: единство материального и идеального в человеке // Современные философские исследования. 2024. № 1. С. 33–43.
DOI: 10.18384/2949-5148-2024-1-33-43

FOR CITATION

Serikov A. Ye. Individual Understanding of Things: Distributed Meanings, Exemplars, and Unicepts. In: *Contemporary Philosophical Research*, 2024, no. 1, pp. 33–43.
DOI: 10.18384/2949-5148-2024-1-33-43