

А. П. РЫЖОВ

ГИБРИДНЫЙ ИНТЕЛЛЕКТ СЦЕНАРИИ ИСПОЛЬЗОВАНИЯ В БИЗНЕСЕ



РАНХиГС

РОССИЙСКАЯ АКАДЕМИЯ НАРОДНОГО ХОЗЯЙСТВА
И ГОСУДАРСТВЕННОЙ СЛУЖБЫ
ПРИ ПРЕЗИДЕНТЕ РОССИЙСКОЙ ФЕДЕРАЦИИ

ИНСТИТУТ ЭМИТ

А. П. Рыжов

Гибридный интеллект. Сценарии использования в бизнесе

*Рекомендовано Учёным советом
Института экономики, математики
и информационных технологий РАНХиГС
в качестве учебного пособия для обучающихся
по специальности «Информационный менеджмент»*



«Академиздат»



ИЗДАТЕЛЬСТВО АКАДЕМИЗДАТ
НОВОСИБИРСК

2019

УДК 004.08
ББК 32.813
Р93

Рыжов, Александр Павлович

Р93 Гибридный интеллект. Сценарии использования в бизнесе / А. П. Рыжов. — Новосибирск : Академиздат, 2019. — 116 с. — (Библиотека Школы IT-менеджмента).

ISBN 978-5-6043351-0-9

Издание открывает серию «Библиотека Школы IT-менеджмента», в рамках которой планируются к публикации результаты исследований ведущих профессоров и преподавателей Школы IT-менеджмента Института экономики, математики и информационных технологий Российской академии народного хозяйства и государственной службы при Президенте Российской Федерации.

Книга написана на основе части курса по аналитическим информационным технологиям, читаемого автором в течение 10 лет в Школе IT-менеджмента, и части курса по теории нечетких множеств, читаемого автором более 20 лет на механико-математическом факультете МГУ имени М. В. Ломоносова.

Для IT-менеджеров, специалистов по разработке и внедрению интеллектуальных цифровых технологий, реальных и потенциальных их пользователей.

УДК 004.08
ББК 32.813

ISBN 978-5-6043351-0-9

© Школа IT-менеджмента РАНХиГС, 2019
© Рыжов Александр Павлович, 2019

Содержание

Введение	5
1. Искусственный интеллект. Возникновение, основные этапы, современное состояние.	8
1.1. Предыстория	8
1.2. Возникновение и развитие.	9
1.2.1. Этап романтизма	11
1.2.2. Этап экспертных систем	12
1.2.3. Зима искусственного интеллекта.	14
1.3. Современное состояние	18
2. Основы гибридного интеллекта	24
2.1. Взгляд из прошлого в будущее	24
2.2. Базовые проблемы гибридного интеллекта	29
2.3. Теория нечетких множеств — математика гибридного интеллекта	33
2.4. Проблема 1 — моделирование описания человеком объектов	47
2.5. Проблема 2 — обработка нечетких описаний объектов	58
3. Сценарии использования гибридного интеллекта	67
3.1. Персонализация	68
3.1.1. Персонализация поиска информации	70

3.1.2. Персонализация обучения в компьютерных обучающих системах.	79
3.1.3. Цифровые привычки и персонализация новостных лент	84
3.2. Оценка и мониторинг процессов.	87
3.2.1. Понятие систем оценки и мониторинга сложных процессов	87
3.2.2. Проблема агрегирования информации в системах оценки и мониторинга	92
3.2.3. Аналитические возможности систем оценки и мониторинга.	102
3.2.4. Возможности и ограничения систем оценки и мониторинга.	106
Список источников	109

Введение

Наверное, нет человека, который не слышал бы или не читал про искусственный интеллект. Люди, далекие от науки и технологий, испытывают широкий спектр эмоций от восхищения (обыграл чемпионов мира в шахматы и Го!) до страха (лишит всех работы, сможет уничтожить человечество, если решит его «оптимизировать»). Пользователи испытывают еще более яркие эмоции от эйфории («Наступил период, когда компания, если она не использует искусственный интеллект в своей деятельности, она проигрывает»¹) до глубокого разочарования («Этот продукт — дерьмо собачье»²). Инвесторы оптимистичны: в первые три месяца 2018 года был поставлен рекорд по финансированию в области искусственного интеллекта — общий объем инвестиций превысил 1,9 миллиарда долларов, рост по сравнению с аналогичным периодом 2017 года составил 69%³; к 2050 году инвестиции

¹ Сбербанк в 2018 году реализует более 150 проектов на основе искусственного интеллекта. 16 сентября 2017: <https://tass.ru/ekonomika/4567834>.

² IBM Watson выдавал «небезопасные и неправильные» рекомендации по лечению рака. 26 июля 2018: <https://hi-news.ru/research-development/ibm-watson-vydaival-nebezopasnye-i-nepravilnye-rekomendacii-po-lecheniyu-raka.html>.

³ Инвестиции в искусственный интеллект бьют рекорды. 29 мая 2018: <https://www.if24.ru/investitsii-v-iskusstvennyj-intellekt-byut-rekordy/>.

достигнут 1 триллиона долларов¹. Специалисты сдержанны: с одной стороны, хочется поучаствовать в освоении указанных миллиардов, а с другой—есть имя, есть репутация, есть научная этика, и для многих это не пустые слова...

В этой книге мы попытаемся аргументировать взвешенную позицию, свободную от упомянутых крайностей. Основные тезисы такой позиции можно сформулировать следующим образом:

1. Впереди нас ждет кризис искусственного интеллекта в широком (если угодно—в обывательском) понимании.
2. Успехи интеллектуальных информационных технологий есть, они серьезны, их невозможно игнорировать. Кризис не может закончиться очередной «зимой» искусственного интеллекта как научно-инженерного направления.
3. Выход из кризиса—системы смешанного или гибридного интеллекта, соединяющего в себе преимущества человека и компьютера как интеллектуальных систем при решении широкого класса прикладных задач.
4. Системы гибридного интеллекта могут быть описаны на математическом языке, мы можем разрабатывать оптимальные по ряду существенных параметров системы такого рода.
5. Гибридный интеллект имеет широкий спектр приложений и представляет собой прагматический аспект искусственного интеллекта.

Если вам интересны аргументы, подкрепляющие сформулированную позицию, эта книга для вас.

Книга написана на основе части курса по аналитическим информационным технологиям, читаемого автором в течение 10 лет в Школе ИТ-менеджмента Российской академии народного хозяйства и государственной службы при Прези-

¹ Искусственный интеллект: многообещающая инвестиция и философская идея. 14 декабря 2018: <http://www.forbes.ru/tehnologii/370385-iskusstvennyy-intellekt-mnogoobeshchayushchaya-investiciya-i-filosofskaya-ideya>.

денте Российской Федерации, и части курса по теории нечетких множеств, читаемого автором более 20 лет на механико-математическом факультете МГУ им. М. В. Ломоносова. Здесь мы ориентируемся на ИТ-менеджеров, специалистов по внедрению интеллектуальных технологий, реальных и потенциальных их пользователей. Это связано со следующими соображениями:

- математические основы гибридного интеллекта достаточно широко представлены в доступной научной литературе, интересующийся читатель сможет без труда их найти и изучить самостоятельно;
- основным драйвером исследований в области интеллектуальных информационных технологий (Artificial Intelligence, Data Sciences в англоязычной терминологии) сейчас является не столько теория, сколько приложения. Поэтому новые сценарии их применения очень важны, и автор надеется, что они возникнут у заинтересованного читателя в отношении гибридного интеллекта.

Конечно же, автор заинтересован в обратной связи и будет рад помочь использованию описанных или (что еще лучше) готовых к обсуждению новых сценариев в бизнесе читателей. Все координаты ниже.

До встречи!

Александр Павлович Рыжов,
профессор
ryjov@mail.ru

Skype/ Facebook: alexander.ryjov
LinkedIn: ru.linkedin.com/pub/alexander-ryjov/1/957/859
<http://www.intsys.msu.ru/staff/ryzhov/>

1. Искусственный интеллект. Возникновение, основные этапы, современное состояние

Для понимания источников возникновения и драйверов развития гибридного интеллекта необходимо провести экскурс в историю возникновения и развития искусственного интеллекта. Эта глава не является основной в книге, поэтому ее можно пропустить читателю, работающему в этой области или хорошо ее знающему. Для тех, кто решит ее прочитать, отметим, что мы попытаемся дать лишь общую картину достаточно крупными мазками. По представленным ссылкам можно в случае возникновения вопросов уточнить недостающие детали.

1.1. Предыстория

Мечта о смышленном искусственном существе-помощнике была всегда и у всех народов. Это нашло отражение в сказках. В русских сказках это колобок (нет аналогов в природе), мальчик-с-пальчик (человекоподобное существо, андроид на современном языке); в сказках других народов — Буратино, Оле Лукояе и многие другие персонажи. Подчеркнем, это — не говорящие медведь или лиса, это — искусственные существа с интеллектом. Таким образом, каждый из нас знаком с искусственным интеллектом с детства. Может, поэтому многие люди, далекие от науки и технологий, так легко про него рассуждают: это эмоционально близкое понятие, мы много про него размышляли, мы его «знаем».

Сотворение человеком существ, похожих на него самого, имеет долгую историю. Мы не будем касаться сложных теологических вопросов, хотя некоторые аналогии присутствуют еще в Книге Бытия (шестой день творения), легенде о големе и других источниках. Можно вспомнить известные из школьного курса истории эпизоды: Прометей вылепил человека из глины; Пигмалион вырезал из слоновой кости статую (Галатею) и полюбил ее, а Афродита оживила статую, и та стала его женой; Гефест создал Пандору, первую женщину на Земле¹. Эта мечта нашла свое отражение и в Средние века, причем над созданием таких механизмов работали лучшие ученые и изобретатели того времени: железный человек Альбертуса Магнуса, механический рыцарь да Винчи, механические дети Пьера Жаке-Дро, гомункул Парацельса... Эти механизмы были не только похожи на человека внешне, но и демонстрировали достаточно сложное поведение: железный человек Альбертуса Магнуса выполнял функции слуги — открывал дверь, приветствовал входящего поднятием руки, мог произнести несколько слов², механический рыцарь да Винчи умел садиться, двигать головой и руками, анатомически правильно открывать и закрывать рот³, механические дети Пьера Жаке-Дро могли писать, рисовать и играть на клавишине⁴.

Однако настоящая история искусственного интеллекта начинается только с изобретения первых вычислительных машин.

1.2. ВОЗНИКНОВЕНИЕ И РАЗВИТИЕ

История искусственного интеллекта как нового научного направления начинается в середине XX века. Основным толчком послужило создание компьютеров как материального

¹ Эти и другие сценарии достаточно широко представлены, например, в Википедии; там же можно найти ссылки на источники и специальную литературу.

² <http://informaticslib.ru/books/item/f00/s00/z0000011/st029.shtml>.

³ http://androbots.ru/istoriya_robototekhniki/robot_da_vinchi/robot_rycar_leonardo.php.

⁴ <https://masterok.livejournal.com/1495215.html>.

объекта, в котором может «жить» искусственный интеллект, хотя компьютеры создавались совсем не для этого. Важным был тот факт, что компьютеры намного быстрее человека решали сложные дифференциальные уравнения, поэтому возник вопрос: каковы границы возможностей компьютеров и достигнут ли машины уровня развития человека? В 1950 году английский ученый Алан Тьюринг пишет статью под названием «Может ли машина мыслить?» [1], в которой описывает процедуру, с помощью которой можно будет определить момент, когда машина сравнивается в плане разумности с человеком, получившую название теста Тьюринга. Смысл этой процедуры достаточно прост: у нас есть две комнаты, в одной расположен человек, в другой — компьютер; есть возможность задавать любые вопросы; если мы не можем уверенно сказать, где человек, а где компьютер, но в этом компьютере реализован искусственный интеллект. В оригинале: «Будет ли в этом случае задающий вопросы ошибаться столь же часто, как и в игре, где участниками являются только люди? Эти вопросы и заменят наш первоначальный вопрос “могут ли машины мыслить?”».

Достаточно очевидно, что такая постановка не могла возникнуть в пустоте. К этому времени сформировалось несколько предпосылок для такой постановки проблемы: философы много веков рассуждали о природе процесса познания, психология сформировалась как наука и добилась определенных результатов в понимании мыслительного процесса, нейрофизиологи достаточно хорошо изучили организацию мозга и механизмы его работы, математики предложили модели нейронов (статья Мак-Коллока У.С. и Питтса В. [2] была опубликована в 1943 году), сформировались основы теории алгоритмов (понятие машины Тьюринга было предложено в 1936 году). Перечисленное составляет целые пласты научной мысли, и мы даже не будем пытаться здесь делать какой-либо обзор этих научных направлений: это содержание десятков учебников, сотен статей, это годами изучается в университетах. Важно, что такая формулировка возникла не как креативная идея у случайного человека, а как результат размышлений одного из сильнейших математиков мира, базирующийся на дости-

жениях многих научных направлений. Важно, что возник критерий, и это уже научная постановка вопроса.

Дальнейшее развитие искусственного интеллекта как научного инженерного направления прошло несколько этапов, которые очень условно можно разделить.

1.2.1. Этап романтизма

Итак, человечество оказалось близко к реализации своей многовековой мечты — созданию искусственного существа, способного мыслить. Выход человечества в космос, впечатляющие результаты ядерной физики, прогресс в создании ЭВМ, успехи в других науках и технологиях делали эту мечту почти осязаемой и добавляли уверенности.

Многие научные коллективы занимались разработкой программных систем, имитирующих функции человеческого мозга: распознавание образов, хранение и поиск информации, логические рассуждения. Примером таких исследований может служить Общий решатель проблем (General Problem Solver), созданный в 1959 году А. Newell, J. C. Shaw, H. A. Simon [3]. Эта программа могла проводить доказательства теорем евклидовой геометрии и логики предикатов, решать шахматные задачи и головоломки (например, головоломку о ханойских башнях). Была предложена и реализована новая архитектура системы — база правил и машина логического вывода.

Создавались системы машинного перевода, анализа и реферирования текстов, даже написания музыки и стихов. Так, в 1954 году были продемонстрированы первые такие системы: IBM совместно с Джорджтаунским университетом (известный Джорджтаунский эксперимент) и (чуть позже и с меньшим пиаром) Институт точной механики и вычислительной техники РАН продемонстрировали возможности ЭВМ в области машинного перевода. Организаторы исследований были искренне уверены, что в течение 3–5 лет проблема машинного перевода будет решена.

В это время организовывались новые семинары (известный семинар А. А. Ляпунова по кибернетике начал работать на механико-математическом факультете МГУ в середине

1950-х годов, семинар по математической лингвистике под руководством В. А. Успенского и В. В. Иванова на филологическом факультете МГУ—1956 год), собиравшие огромные аудитории, открывались новые специальности (например, структурная лингвистика), готовились специалисты (первый прием студентов на отделение структурной и прикладной лингвистики филологического факультета МГУ им. М. В. Ломоносова был проведен в 1960 году). В Академии наук СССР был создан Научный совет по кибернетике (1959 год), позже переименованный в Научный совет по комплексной проблеме «кибернетика» при Президиуме АН СССР (1962 год), в ведущих университетах и технических вузах создавались профильные кафедры, исследования щедро финансировались.

В этот период были не только яркие взлеты, но и драматические падения. Например, резко «взлетели» нейронные сети (Ф. Розенблатт создал персептрон, вышла его книга [4]), так же резко они «упали» (вышла книга М. Минского и С. Пейперта [5], где была доказана ограниченность персептронов Розенблатта), исследования в области искусственных нейронных сетей были практически заморожены на многие десятилетия.

Реальные задачи оказались намного сложнее. В ходе выполнения многих проектов сформировалось понимание отличия данных от знаний, важности работы со знаниями, получившее свое техническое воплощение в концепции экспертных систем.

Этот этап длился до начала 1970-х годов.

1.2.2. Этап экспертных систем

Исторически первой системой, получившей название экспертной, была система Dendral, созданная в 1965 году в Стэнфордском университете. Эта система позволяла химику автоматизировать процесс определения химической структуры вещества. Задав некоторую информацию о веществе и данные спектрометрии, пользователь получал ответ в виде химической структуры. Это была первая система такого класса, применявшаяся в промышленности. Описано более двадцати ее применений в реальных задачах, дающих качество решений на уровне эксперта в области химии.

Наиболее известной экспертной системой является разработка той же лаборатории Стэнфорда начала 1970-х годов — система *Mycin*. Задачей системы была диагностика бактерий, вызывающих тяжелые инфекции (бактериемия и менингит), а также выработка рекомендаций по необходимому количеству антибиотиков в зависимости от массы тела пациента в случае того или иного заболевания. Удивительным оказался тот факт, что в точности диагноза на тестах она превосходила профессоров *Stanford medical school*. Успех был достигнут за счет специальной архитектуры системы, позволяющей эффективно работать со знаниями (база знаний на основе правил), а также механизмом вывода, позволяющим работать в условиях неполноты информации и ее неопределенности. Упомянутая выше система *Dendral* работала «по старинке»: там был блок генерации вариантов (достаточно понятно, что для химических формул это можно сделать) и блок проверки формул-кандидатов. Такой процесс — «генерация и проверка» — позволяет постоянно сокращать число возможных рассматриваемых вариантов. Это близко к «умному перебору» вариантов, используемому в системах первого этапа (раздел 1.2.1).

Успех *Mycin* был настолько ошеломляющим, что была выпущена так называемая оболочка этой системы с пустой базой знаний (*E-Mycin*), каждый мог загрузить туда свои знания и стать обладателем такого чудесного помощника в своей предметной области.

Однако чуда не случилось. Оказалось, что извлечь знания из головы экспертов и загрузить их в систему не получается. Было потрачено много усилий на модели представления знаний (декларативные, продукционные, сетевые, фреймовые и много вариантов их комбинаций — их описания без труда можно найти в интернете) и их реализацию, работу со статическими и динамическими знаниями, различные логические схемы вывода и т. п., но повторить успех *Mycin* не удалось. Зато возникло большое количество фейковых (говоря современным языком) экспертных систем. Действительно, возьмите любую инструкцию, например к чайнику. Там есть «база знаний»: если чайник не работает — проверьте розетку и т. п. Если упаковать такую инструкцию в виде компьютерной про-

граммы, то, пойдя на совсем небольшой компромисс со своей совестью или просто по незнанию в силу отсутствия соответствующего образования, можно объявить о создании экспертной системы нового поколения и начать ее продавать. Рынок подогрет, инвесторы переполнены энтузиазмом, успех обеспечен... Такой обман, естественно, не понравился ни рынку, ни инвесторам; венчурное и корпоративное финансирование этого направления достаточно быстро прекратилось. Хотя некоторые авторы продолжают их считать и насчитали уже несколько тысяч таких «промышленных экспертных систем». Чуть больше продержалось финансирование их государственных источников — это был разгар холодной войны. Формировались глобальные проекты: японский проект создания ЭВМ пятого поколения (1982 год), стратегическая компьютерная инициатива как часть стратегической оборонной инициативы Р. Рейгана (1983 год), где экспертные системы занимали центральное место¹. Однако геополитические изменения понизили актуальность создания новых дорогостоящих систем, щедрое финансирование прекратилось. Разрозненные группы в университетах занимались переосмыслением достигнутого. Наступила зима искусственного интеллекта.

Этот этап длился до начала 1990-х годов.

1.2.3. Зима искусственного интеллекта

Основные события следующего десятилетия не касались напрямую искусственного интеллекта (это словосочетание практически перестало употребляться), но создали базу для его нового этапа.

К таким событиям можно отнести огромный прогресс в развитии элементной базы компьютеров и ее существенное удешевление. Приведем всего лишь три факта: стоимость гигабайта памяти (рис. 1), затраты на обработку одного бита информации (рис. 2) и рост мощности вычислений (рис. 3). Обратите

¹ Мы не имеем возможности описывать здесь упомянутые проекты. Интересующийся читатель без труда найдет их описание и историю в научно-популярной литературе, хорошо представленной в интернете.

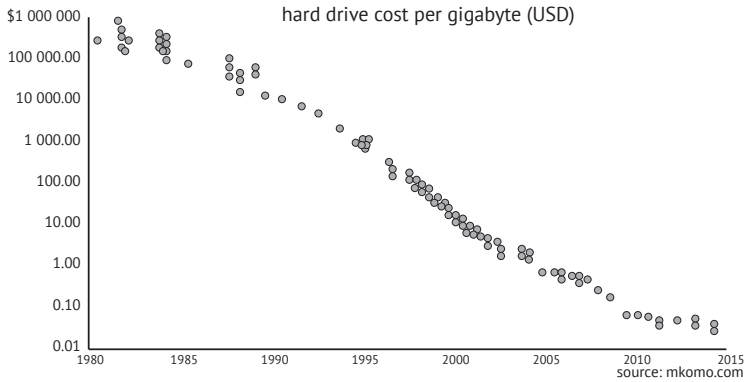


Рис. 1. Средняя стоимость гигабайта памяти

Источник: <http://www.mkomo.com/cost-per-gigabyte-update>.

внимание—езде логарифмическая шкала; в обычной шкале график был бы почти вертикальным.

Это позволило накапливать и обрабатывать большие объемы данных. Возникли такие компа-

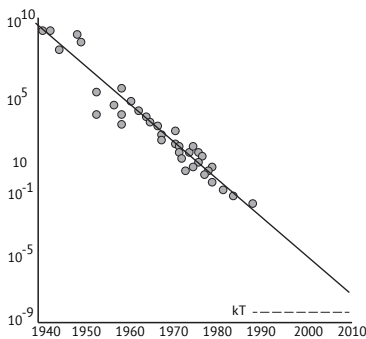


Рис. 2. Снижение затрат на обработку одного бита

Источник: https://www.extech.ru/files/anr_2015/anr_2.pdf.

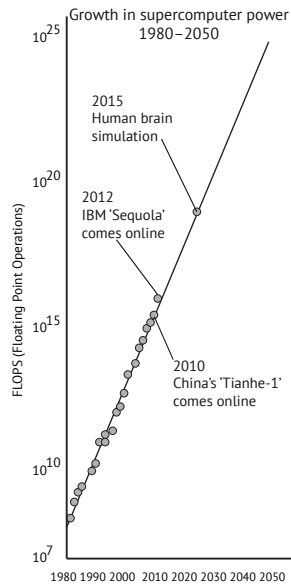


Рис. 3. Рост мощности вычислений

Источник: <https://www.futuretimeline.net/subject/computers-internet.htm>.

нии, как Google (зарегистрирована как частная компания 4 января 1996 года, а 19 августа 2004 года начала продажу своих акций на фондовом рынке), Facebook (основана 4 февраля 2004 года, выход на IPO—18 мая 2012 года) и многие другие, способные привлечь многие миллионы пользователей.

Вычисления стали не только доступными для компаний (порог входа на этот рынок понизился существенно), но и массовыми. Этому способствовало лавинообразное распространение интернета (включая мобильный интернет), смартфонов, планшетов, health trackers и многих других устройств (интересующийся читатель без труда найдет соответствующую аналитику в интернете). Люди стали не только потреблять, но и генерировать огромные объемы информации, которую компании с удовольствием хранили, систематизировали, монетизировали. Возникла новая цифровая вселенная—модель нашего физического мира, грамотно работающая в которой можно было много понять о мире, в котором живут люди. Эта цифровая вселенная стала все больше влиять на мир физический. Приведем лишь несколько фактов:

— Биржи¹:

- 70% торгов на Уолл-стрит осуществляются через роботов;
- в России на долю роботов приходится почти половина заявок на ММВБ и до 90%—на срочном рынке РТС;
- на EUREX алгоритмическая торговля уже составляет 90%, и эта цифра стремится к 100%.

— Торговля²:

- Amazon.com утверждает, что 40% продаж генерируются через механизмы рекомендаций.

— Социальные сети³:

- 50% выбора работы на LinkedIn—прямой результат рекомендаций.

¹ См.: Восстание машин. Роботы захватили и разоряют мировые биржи?//Версия. 2019. № 1; <https://versia.ru/roboty-zaxvatili-i-razoryayut-mirovye-birzhi>.

² https://netpeak.net/ru/blog/review_ratings/.

³ https://www.quora.com/How-does-LinkedIns-recommendation-system-work?cm_mc_uid=39580370786314790366644&cm_mc_sid_50200000=1479560781.

Эти тенденции были прекрасно обобщены в знаменитом отчете McKinsey Global Institute «Big data: The next frontier for innovation, competition, and productivity»¹, который я настоятельно рекомендую внимательно прочитать. Там же представлено множество фактов, ссылок, аналитики.

Также следует отметить развитие робототехники, связанное с удешевлением и доступностью элементной базы и появлением новых материалов, на основе чего возникли новые типы сенсоров и акторов. Существенным сдвигом в этой области был перенос роста из области производства в область домашнего использования роботов (рис. 4).

Перечисленные факторы (и не только они—автор не претендует на полноту, эти факторы кажутся ему наиболее важными) и создали ситуацию, когда исследователи и компании снова «вспомнили» про искусственный интеллект. В массовой прессе стали появляться статьи со словами: «Возник новый синтез четырех областей, включая математику, нейробиологию, информатику и психологию. Результат этого удивителен.

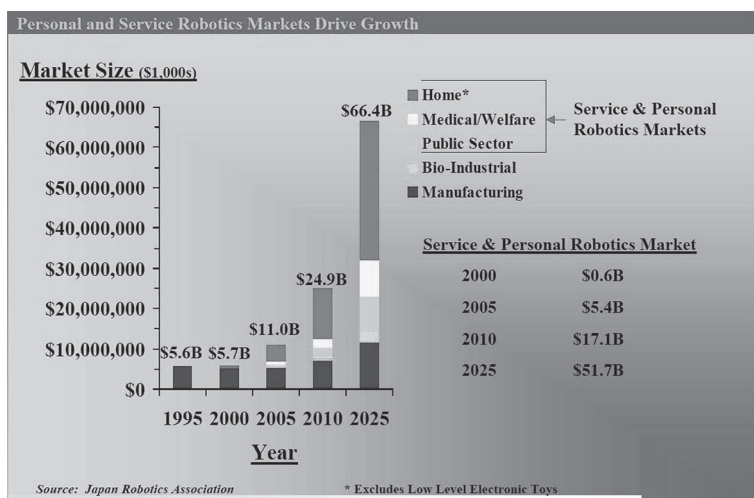


Рис. 4. Драйверы рынка роботов

¹ Big data: The next frontier for innovation, competition, and productivity // <https://www.mckinsey.com/business-functions/digital-mckinsey/our-insights/big-data-the-next-frontier-for-innovation>.

Мы видим, что когнитивные вычисления находятся на пороге, они стучат в дверь потенциально массовых приложений¹ или «Пришло время строить роботы на основе искусственного интеллекта. Мечта — поставить робота в каждый дом» (там же). В ведущих университетах и бизнес-школах снова стали проводиться семинары, посвященные искусственному интеллекту. Пример такого приглашения, полученного автором, приведен ниже.

VLAB Event — Rise of the Machines: The Business of AI

Event Date: Tuesday, September 15, 2009 at 6:00pm

Location: Stanford Business School

This month's event looks at the leading companies and thought leaders in AI to discuss the business of AI including autonomous robotics, learning systems and machine/human fusion and asks:

- * What are the business opportunities in AI?
- * How is AI being used to drive competitive advantage?
- * What business models are likely to work?
- * Who's investing in the space?
- * What 's next?
- * What applications are most needed and what opportunities are there for entrepreneurs?

Come join us to find out about the next wave of AI-driven companies and see how entrepreneurs are leveraging thinking machines to drive business opportunities.

Таким образом, «зима» искусственного интеллекта закончилась в середине 2000-х годов. Этап «ренессанса» длится по настоящее время, и на нем мы остановимся подробнее, это важно для понимания дальнейшего содержания книги.

1.3. СОВРЕМЕННОЕ СОСТОЯНИЕ

Важным событием современного этапа был факт официального прохождения теста Тьюринга — это случилось 8 июня 2014 года², и теперь мы не имеем даже такого простого критерия.

¹ John Markoff. Brainy Robots Start Stepping into Daily Life//The New York Times. 2006, July 20; https://archive.nytimes.com/www.nytimes.com/learning/teachers/featured_articles/20060720thursday.html.

² TURING TEST SUCCESS MARKS MILESTONE IN COMPUTING HISTORY: <http://www.reading.ac.uk/news-and-events/releases/PR583836.aspx>.

Доступность огромного количества разнообразных данных и дешевизна их обработки, возникшие на предыдущем этапе, сделали возможным проведение большого количества экспериментов с данными. В ходе таких экспериментов можно было проверить работоспособность и потенциал многих известных алгоритмов. Ярким примером служит искусственная нейронная сеть специальной архитектуры (стохастическая рекуррентная сеть или машина Больцмана), известная с 1985 года. В статистике аналогичная структура называется случайным Марковским полем и известна с начала 1980-х годов. Для обучения таких алгоритмов нужна большая выборка. Если в 1980-х годах получить такую выборку было невозможно, то Facebook смог предоставить Social Face Classification (SFC) dataset, содержащий 4,4 миллиона размеченных фотографий (лиц) своих 4030 пользователей (от 800 до 1200 фотографий каждого пользователя)¹. Это позволило обучить девятислойную нейронную сеть с более чем 120 миллионами весов соединений распознавать лица с точностью 97,25% (у людей этот показатель составляет 97,5%). Упомянутая выше статья опубликована в 2014 году, то есть машина Больцмана «ждала» такого своего применения почти 30 лет. Безусловно, авторы проделали огромную и сложную работу, но для получения результата нужно было создать Facebook, сделать его привлекательным для сотен миллионов пользователей, уговорить часть из них разметить свои фотографии.

Часть таких экспериментов оказались удачными. В ситуации отсутствия критерия многое из этого стали называть искусственным интеллектом. Действительно, почему бы не назвать настройку фотографий или распознавание номера автомобиля (а также много всего другого) искусственным интеллектом? Потому что система делает это лучше, чем человек? Ну тогда и калькулятор тоже является искусственным интеллектом — он же быстрее среднего человека складыва-

¹ Yaniv Taigman, Ming Yang, Marc'Aurelio Ranzato, Lior Wolf. DeepFace: Closing the Gap to Human-Level Performance in Face Verification: https://www.cv-foundation.org/openaccess/content_cvpr_2014/papers/Taigman_DeepFace_Closing_the_2014_CVPR_paper.pdf.

ет и умножает числа, извлекает корни, берет логарифмы. Шальные инвестиционные деньги (см. Введение) только увеличивают соблазн и порождают целую толпу громко кричащих «специалистов» без базового образования и «инноваторов», которым достаточно «харизмы и креатива». Такая ситуация не может существовать долго и обычно заканчивается печально, в том числе, к сожалению, и для специалистов.

Замечено это было еще в самом начале современного периода. Так, в статье John Markoff (2006 год) приводятся слова Eric Horvitz, избранного президентом Американской ассоциации искусственного интеллекта: «На конференциях вы слышите фразу: “ИИ на уровне человека”, и люди говорят это, не краснея». Сейчас искусственный интеллект обсуждается не столько на научных конференциях, сколько в околonaучной бизнес-прессе, и обсуждают его не столько специалисты, сколько продавцы (юристы, бухгалтеры), а объем инвестиций увеличился на порядки. Это все привело к хайпу, где теряются смыслы.

Многие специалисты и заказчики видят такую ситуацию. Достаточно посмотреть заголовки статей: «ИИ направляется обратно к корыту. Ожидания в отношении искусственного интеллекта (ИИ) становятся слишком завышенными. ИИ действительно все изменит, но не в ближайшее время»¹, «Завышенные ожидания: искусственный интеллект все еще зависит от людей»², «ИИ наиболее всего нуждается в управлении ожиданиями»³, «Вот почему ИИ не может решить все»⁴, «Семь

¹ AI heading back to the trough. The expectations over artificial intelligence (AI) are becoming too inflated. AI will indeed change everything, but not any time soon. 11 июля 2017 г.: <https://www.networkworld.com/article/3206313/internet-of-things/ai-heading-back-to-the-trough.html>.

² Inflated Expectations: Artificial Intelligence Still Depends on Humans. 29 июля 2016 г.: <https://techonomy.com/2016/07/27222/>.

³ What AI mostly needs is expectation management. 17 декабря 2017 г.: <https://towardsdatascience.com/what-ai-mostly-needs-is-expectation-management-625c90ea147f>.

⁴ Here's Why AI Can't Solve Everything. 27 мая 2018 г.: <https://www.sciencealert.com/ai-machine-learning-solve-all-humanity-problems-unrealistic>.

смертных грехов предсказаний ИИ. Ошибочные экстраполяции, ограниченное воображение и другие распространенные ошибки, которые отвлекают нас от более продуктивного мышления о будущем. Мы окружены истерией о будущем искусственного интеллекта и робототехники—истерией о том, насколько мощными они станут, как быстро и как они будут выполнять работу»¹. Отметим, что это не сайты каких-нибудь магазинов или банков, это вполне reputable аналитические и научные организации, включая, например, МИТ. Такие мнения стали высказываться с 2016–2017 годов, и их поток только увеличивается, достаточно посмотреть интернет.

Заметим также, что эту болезнь видят и на нее реагируют компании—лидеры рынка. Так, Facebook закрыл подразделение М, занимавшееся разработкой виртуального помощника на основе искусственного интеллекта, 19 января 2018 года. «Последний день работы виртуального помощника М будет 19 января,—сообщил Facebook—и подрядчикам, которые работали над ним, будет предложена другая работа в компании»². Известная компания Cambridge Analytica, которая занималась сбором данных о пользователях интернета и соцсетей, составлении их психологических портретов и разработке персонализированной рекламы, а также использующая технологии искусственного интеллекта для разработки стратегической коммуникации в ходе избирательных кампаний в интернете, прекратила все операции 2 мая 2018 года³. IBM Watson Health (наверное, никто боль-

¹ The Seven Deadly Sins of AI Predictions. Mistaken extrapolations, limited imagination, and other common mistakes that distract us from thinking more productively about the future. We are surrounded by hysteria about the future of artificial intelligence and robotics—hysteria about how powerful they will become, how quickly, and what they will do to jobs. 6 октября 2017 г. : <https://www.technologyreview.com/s/609048/the-seven-deadly-sins-of-ai-predictions/>.

² Facebook is shutting down M, its personal assistant service that combined humans and AI: <https://www.theverge.com/2018/1/8/16856654/facebook-m-shutdown-bots-ai>.

³ Cambridge Analytica and Scl Elections Commence Insolvency Proceedings and Release Results of Independent Investigation into Recent Allegations: <https://ca-commercial.com/news/cambridge-analytica-and-scl-elections-commence-insolvency-proceedings-and-release-results-3>.

ше и громче не кричал о своих успехах в области искусственного интеллекта) объявило о массовых сокращениях (до 70% сотрудников) 25 мая 2018 года¹. Мы приводим только факты; различные версии происходящего читатель сможет без труда найти в интернете.

Новая «зима» искусственного интеллекта началась? Ситуация действительно очень напоминает конец этапа экспертных систем (раздел 1.2.2), только более выраженный: огромные инвестиционные ресурсы, безумный маркетинговый хайп, завышенные ожидания всех, поток фантазий юристов и бухгалтеров, поделки малограмотных «специалистов»... Забавным отличием настоящего этапа развития искусственного интеллекта от предыдущих является позиция некоторых ученых и разработчиков, сформировавшаяся в 2016 году. Они стали разделять искусственный интеллект на два: сильный и слабый. Сильный — это тот, который «настоящий», но который будет нескоро, даже непонятно когда; слабый — это тот, который «почти настоящий», который уже распознает котиков и улучшает фотографии (интересующийся читатель наверняка найдет много рассуждений на эту тему в интернете). Эта позиция понятна, в ней одновременно присутствуют два «месседжа»: один — инвесторам и грантодателям («Я занимаюсь искусственным интеллектом»), другой — своей совести и коллегам («Я — специалист, я понимаю, что то, чем я сейчас занимаюсь — это не искусственный интеллект»). В связи с этим вспоминаются известные слова героя романа М. А. Булгакова «Мастер и Маргарита»: «Свежесть бывает только одна — первая, она же и последняя. А если осетрина второй свежести, то это означает, что она тухлая!». Мы, однако, не будем осуждать или приветствовать эту точку зрения. Заметим лишь, что такие попытки подмены понятий — тоже один из ярких индикаторов кризиса.

¹ IBM's Watson Health wing left looking poorly after 'massive' layoffs Up to 70% of staff shown the door this week, insiders claim: https://www.theregister.co.uk/AMP/2018/05/25/ibms_watson_layoffs/.

1. ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ

Последующее содержание книги — аргументация в пользу того, что выход из кризиса есть, и он заключается в создании человеко-компьютерных систем, объединяющих сильные стороны интеллекта человека и компьютера при решении практических задач.

2. Основы гибридного интеллекта

Кризисные явления в области искусственного интеллекта видны специалистам, компаниям, наблюдателям. Есть ли выход из кризиса или впереди очередная «зима»? Для ответа на этот вопрос полезно вернуться к истокам. Как видели искусственный интеллект люди, положившие начало этому направлению? Насколько их взгляд актуален для современного понимания этого направления? Каковы основные проблемы такого понимания искусственного интеллекта и можно ли их решить? Взгляд автора на сформулированные вопросы и возможные ответы на них представлены в этой главе.

2.1. Взгляд из прошлого в будущее

Один из основателей кибернетики и искусственного интеллекта Уильям Росс Эшби (William Ross Ashby) в известной работе «Введение в кибернетику»¹ в 1956 году писал: «Интеллектуальная сила, как и физическая сила, может быть усилена. Пусть никто не говорит, что это невозможно сделать, потому что гены-шаблоны делают это каждый раз, когда они образуют мозг, который вырастает, чтобы быть чем-то лучше, чем ген-образ мог бы

¹ Ashby W.R. An Introduction to Cybernetics. London, UK: Chapman and Hall, 1956. P. 271; <http://pespmc1.vub.ac.be/books/IntroCyb.pdf>.

быть подробно описан. Новым является то, что мы можем теперь делать это синтетически, сознательно, преднамеренно» (в оригинале: «Intellectual power, like physical power, can be amplified. Let no one say that it cannot be done, for the gene-patterns do it every time they form a brain that grows up to be something better than the gene-pattern could have specified in detail. What is new is that we can now do it synthetically, consciously, deliberately»). Это значит, что он видел искусственный интеллект не как самодостаточную самодумающую автономную машину, а как усилитель интеллекта человека. Такой же точки зрения, выраженной более конкретно, придерживался Джозеф Карл Робнетт Ликлайдер (Joseph Carl Robnett Licklider), во многом предвосхитивший современный облик информационных технологий (его, в частности, называют духовным отцом интернета): «Человеко-компьютерный симбиоз является ожидаемым развитием сотрудничества между людьми и электронными компьютерами. Это будет связано с очень тесной связью между человеческим и электронным членами партнерства... Предварительные анализы показывают, что симбиотическое партнерство будет выполнять интеллектуальные операции гораздо эффективнее, чем только человек может их выполнять»¹ (в оригинале: «Man-computer symbiosis is an expected development in cooperative interaction between men and electronic computers. It will involve very close coupling between the human and the electronic members of the partnership... Preliminary analyses indicate that the symbiotic partnership will perform intellectual operations much more effectively than man alone can perform them»).

Отметим, что некоторые компании-лидеры также акцентируют свое внимание на роли искусственного интеллекта как помощника в решении конкретных задач. Так, известная компания «Мак-Кинзи», проанализировав более 1000 современных технологий, составила список из 12 революционных технологий (Disruptive technologies), которые окажут существ-

¹ Licklider J. C. R. Man-Computer Symbiosis. IRE Transactions on Human Factors in Electronics. Vol. HFE-1, 4–11. P. 4; <http://groups.csail.mit.edu/medg/people/psz/Licklider.html>.

венное влияние на нашу жизнь в перспективе до 2025 года. Номер два в этом списке — автоматизация интеллектуальной работы (Automation of knowledge work), характеризуя которую авторы пишут: «Эти возможности не только расширяют вычисления в новые сферы... но также создают новые отношения между работниками знаний и машинами. Все чаще можно взаимодействовать с машиной так, как это было бы с коллегой»¹ (в оригинале: «These capabilities not only extend computing into new realms... but also create new relationships between knowledge workers and machines. It is increasingly possible to interact with a machine the way one would with a co-worker»). IBM в 2018 году сформулировала различие между «традиционным» искусственным интеллектом и тем, что позволяет строить практические приложения на его основе: «В IBM мы руководствуемся термином «дополненный интеллект», а не «Искусственный интеллект». Это критическое различие между системами, которые расширяют и дополняют человеческий опыт, а не те, которые пытаются копировать человеческий интеллект. Мы фокусируемся на создании практических приложений ИИ, которые помогают людям с четко определенными задачами и в перспективе предоставляют широкий спектр общих сервисов ИИ на платформе для поддержки широкого спектра новых приложений»² (в оригинале: «At IBM, we are guided by the term “augmented intelligence” rather than “artificial intelligence”. It is the critical difference between systems that enhance and scale human expertise rather than those that attempt to replicate all of human intelligence. We focus on building practical AI applications that assist people with well-defined tasks, and in the process, expose a range of generalized AI services on a platform to support a wide range of new applications»). Созвучно с идеями Эшби и Ликлайдера, не правда ли? Эта точка зрения близка и автору на-

¹ McKinsey Global Institute. Disruptive technologies: Advances that will transform life, business, and the global economy. 2013. May. P. 41; http://www.mckinsey.com/insights/business_technology/disruptive_technologies.

² IBM. Cognitive computing. Preparing for the Future of Artificial Intelligence 2018: <http://research.ibm.com/cognitive-computing/ostp/rfi-response.shtml>.

стоящей книги. Augmented intelligence — расширенный, дополненный, усиленный интеллект. Наверное, на русский язык это можно перевести и как гибридный интеллект, где обе компоненты — человеческая и компьютерная — являются равноправными. Действительно, есть много примеров того, когда компьютерные программы оказываются эффективнее человека — известные всем примеры шахмат, Го, покера и многие другие. Поэтому прилагательное «дополненный», несущее смысл «вспомогательный, неглавный», не видится корректным. Мы не будем придерживаться человеко-центрического эгоизма и, считая обе компоненты равноправными, называть такие системы гибридным интеллектом.

Отличие гибридного интеллекта от текущего понимания искусственного интеллекта можно продемонстрировать с помощью следующей аналогии. Люди, наряду с мечтой об интеллектуальном помощнике, всегда мечтали летать. Летать, как птицы: Икар, орнитоптер Леонардо да Винчи (устройство, придуманное да Винчи, которое теоретически могло поднимать человека в воздух, как птицу), даже в конце XIX — начале XX века передвижение по воздуху виделось именно таким образом. В качестве примера можно привести некоторые открытки из проекта «En L'An 2000» — как художники видели повседневную жизнь в 2000 году¹ (открытки были напечатаны в 1899, 1900, 1901 и 1910 годах) — рис. 5.

Первая партия этих иллюстраций, в которых художники попытались предугадать будущее, готовилась для Всемирной выставки в Париже.

Удивительно, но некоторые технологии были предсказаны достаточно точно (например, смартфоны — рис. 6).

Более того, люди и сегодня пытаются летать, как птицы — достаточно набрать это словосочетание или «wingsuit» в любом поисковике и увидеть множество клипов. Это очень дорого, очень рискованно, это не имеет никакого практического смысла. Но люди будут и дальше пытаться летать, как птицы несмотря ни на что, потому что это — мечта. Однако мы все летаем на самолетах (не как птицы): это дешево, без-

¹ <http://postomania.net/post111710824/>.



Aerial Firemen



Рис. 5. Открытки из проекта «En L'An 2000»

опасно, осмысленно (быстрая доставка людей и грузов на большие расстояния). Подобно этому мечта об искусственном интеллекте расслаивается на романтический аспект (игра в шахматы и Го, распознавание котиков) и прагматический аспект (гибридный интеллект).

Если мы, соглашаясь с отцами-основателями, Мак-Кинзи и ИБМ, а также многими учеными и разработчиками интел-

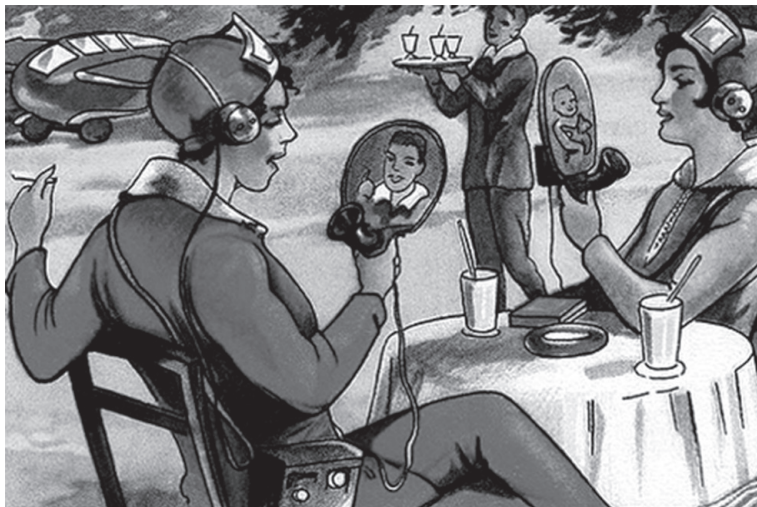


Рис. 6. Открытка из проекта «En L'An 2000»: прообраз смартфона

лектуальных систем, видим в гибридном интеллекте выход из кризиса искусственного интеллекта, то это означает, в частности, переосмысление базовых проблем, про которые мы должны думать.

2.2. БАЗОВЫЕ ПРОБЛЕМЫ ГИБРИДНОГО ИНТЕЛЛЕКТА

Общая схема системы на основе гибридного интеллекта представлена на рис. 7. Предполагается, что такие системы способны обрабатывать информацию, полученную как на основе показаний различных измерительных приборов, так и на основе описаний объектов реального мира человеком.

Заметим, что все измерения, которыми мы пользуемся, есть результат договоренностей, их нет в природе. Например, расстояния измерялись везде по-разному: верста, дюйм, сажень, локоть и пр. Развитие торговли потребовало унификации этих разных величин, и во Французской академии наук была создана специальная комиссия под руководством известного математика и инженера де Борда (кто изучал гидравлику, знает

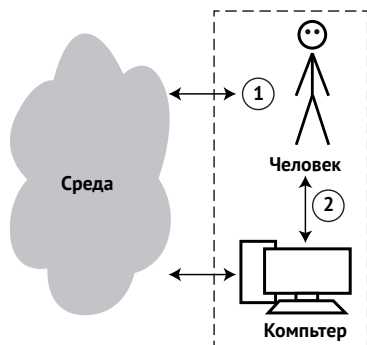


Рис. 7. Общая схема систем на основе гибридного интеллекта

теорему Борда–Карно), куда входили Лагранж, Лаплас, Кондорсе и другие известные всем по базовому курсу математического анализа ученые. Они определили новую величину—метр—как одну десятиmillionную расстояния от Северного полюса до экватора (четверть земной окружности), измеренного вдоль меридиана, проходящего через Париж, и процедуру ее измерения.

На основе этого был изготовлен эталон метра (сначала из латуни—в 1795 году, затем из платины—в 1799 году), которым мы все пользуемся. Единица массы (килограмм, определение которого было основано на массе 1 дм³ воды) тоже была привязана к определению метра, и мы также пользуемся сейчас именно такой величиной.

Похожая ситуация была и с измерением температуры—это был «запрос» промышленности и медицины. Нужен был дешевый прибор, который был бы независим от атмосферного давления (чем «грешил», например, термоскоп Галилея, описанный в 1597 году) и который можно было легко изготовить. Цельсий в 1742 году предложил использовать два состояния вещества, одинаковые везде (таяние льда и кипение воды), и поделил положение трубки с ртутью в этих состояниях на 100 единиц—мы пользуемся этим прибором до сих пор.

Именно возникновение возможности проводить достаточно точные измерения привело к разделению единой науки—натурфилософии—на естественные науки (современные физику, химию, инженерию и пр.) и философию (современные гуманитарные науки).

Поэтому измерения элементов среды, поступающие от приборов (нижняя горизонтальная стрелка на рис. 7), мы будем считать достаточно хорошо изученной и решенной пробле-

мой, а вот измерения, получаемые человеком (верхняя горизонтальная стрелка на рис. 7), не являются таковыми.

В связи с этим возникают проблемы.

Проблема 1 (моделирование описания человеком объектов / Perception modelling): Как человек описывает объекты? Можем ли мы описывать объекты наиболее надежным и эффективным для последующей обработки образом?

Эта проблема обозначена цифрой 1 на рис. 7. Мы должны предложить своеобразный градусник, позволяющий человеку описывать объекты с минимальной неопределенностью, или если несколько людей описывают один и тот же объект, то их описания должны быть максимально единообразны. Мы хотим, чтобы наши описания максимально соответствовали реальному объекту, то есть ситуации, когда одному реальному объекту соответствуют разные описания, или, наоборот, разным объектам соответствуют одинаковые описания, были минимизированы. Как это сделать, обсуждается далее (раздел 2.4).

Естественным становится вопрос: а как обрабатывать такие описания?

Проблема 2 (обработка описаний объектов / Perception-based computing). Как мы обрабатываем такие «человеческие»/perception-based описания (например, ищем или обобщаем)? Можем мы оптимизировать такие вычисления?

Эта проблема обозначена цифрой 2 на рис. 7. Можно ли ввести показатели качества обработки такой информации (например, качество поиска)? Как они связаны с качеством описаний? Эти вопросы обсуждаются в разделе 2.5.

Такая проблема встает и тогда, когда мы работаем только с показаниями приборов, но по правилам, задаваемым человеком (то есть верхняя горизонтальная стрелка отсутствует). Правила часто задаются в виде «если — то», например: «Если температура высокая, а давление низкое, то ситуация опасная». Мы можем описывать температуру и давление разными наборами слова; какой набор является «хорошим»?

Для того чтобы решать подобные проблемы, нужно научиться описывать эти процессы (описание, поиск, логический вывод) на каком-то математическом языке. Выбор та-

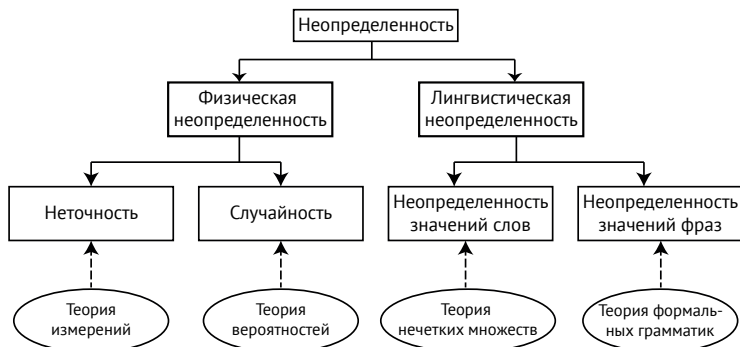


Рис. 8. Классификация неопределенности

кого языка зависит от типа неопределенности, присутствующего в модели. Вопрос соотношения реального объекта и его модели (это одно и то же или это разные вещи?)—сложный философский вопрос. Мы не будем его здесь обсуждать. Мы будем придерживаться позиции, что модель—это неполное «грубое» описание реального объекта, они неэквивалентны. Из этого (неэквивалентности объекта и его модели) следует, что в любой модели есть неопределенность. Неопределенность—тоже сложное понятие. Достаточная для дальнейшего понимания книги классификация неопределенности и связанные с типами неопределенности математические инструменты приведены на рис. 8.

Неточность связана с измерительными приборами. Например, у нас есть весы с точностью 100 грамм. Мы взвешиваем на них мешки и загружаем в вагон. После загрузки нескольких сотен мешков что мы можем сказать о весе вагона? Это может оказаться важным, например, с точки зрения безопасности перевозки. Для ответа на подобные вопросы есть интервальная арифметика.

Случайность достаточно долго и глубоко изучалась в философии, мы не будем здесь проводить обзор подобных работ. Отметим, что специальный раздел математики—теория вероятностей—позволяет обрабатывать случайные величины и процессы и понимать, насколько наши описания и вычисления соответствуют реальности.

Теория формальных грамматик позволяет понять неопределенность смысла фраз.

В контексте сформулированных выше проблем нас интересует неопределенность значений слов. Математическая теория, специально разработанная для обработки такого типа неопределенности, называется теорией нечетких множеств. Краткий обзор основных понятий этой теории представлен в разделе 2.3.

2.3. ТЕОРИЯ НЕЧЕТКИХ МНОЖЕСТВ — МАТЕМАТИКА ГИБРИДНОГО ИНТЕЛЛЕКТА

Известна точная дата возникновения термина «нечеткое множество» (fuzzy set) — 20 августа 1965 года. Это дата выхода восьмого номера журнала *Information and control*, в котором была опубликована статья Fuzzy Sets¹ профессора Лотфи Заде (Zadeh L. A.) из Berkeley University (представлена в редакцию 30 ноября 1964 года). В это время профессору Заде было 44 года, он был известным специалистом в области теории систем, теории автоматического управления и их приложений.

Понятие нечеткого множества, введенное Заде в статье Fuzzy Sets, является интуитивно простым, однако оно касается базового понятия математики — понятия множества — и поэтому является одновременно концептуальным. Учитывая интересы читателей, мы объясним это понятие с нескольких точек зрения: содержательной, формальной (математической) и инженерной (специалистам в области ИТ).

Содержательно понятие нечеткого множества касается следующего. Все, кто пытался разрабатывать модели или системы, учитывающие оценки человека (пользователя, эксперта или в любой другой роли), сталкивались с ситуацией, когда непонятно, как определить в модели или запрограммировать в системе используемые нами понятия «малый/большой рост» (например, курса акций), «высокая/низкая температура» (на объекте управления или температура больного),

¹ https://www-liphy.ujf-grenoble.fr/pagesperso/bahram/biblio/Zadeh_FuzzySetTheory_1965.pdf.

«сильное/слабое изменение» (параметра модели). Это связано с тем, что понятие множества, с помощью которого мы строим модели, требует от нас конкретного определения границы, когда «малый» (высокий, сильный) становится «нема-
 лый» (высокий, сильный). Иначе никак. Для разных людей (пользователей, экспертов) эта граница может быть разной. Какая из них правильная? Мы вынуждены придумывать разные ухищрения (например, усреднять мнения экспертов), что приводит к тому, что модель не отражает ничьего мнения и никого не устраивает в конечном счете. Однако люди как-то понимают друг друга без четкого и однозначного определения подобной границы. Теория нечетких множеств дает формализм описания подобных понятий и тем самым снимает отмеченное выше противоречие.

Для *формального* объяснения напомним основные понятия теории множеств. Сначала говорится, что существует универсальное множество U . Это понятие не имеет формального определения. Говорится, что все понимают это одинаково — множество студентов в университете, множество стульев в аудитории и т. п. В таком универсальном множестве можно определить подмножества и сделать это уже формально. Именно здесь и начинается математика. Подмножества определять можно по-разному, самым распространенным и общим является определение через характеристическую функцию. Характеристическая функция множества A определяется на всех элементах множества U и принимает два значения: 1 или 0:

$$h_A(u) = \begin{cases} 1, & \text{если } u \in A \\ 0, & \text{если } u \notin A. \end{cases}$$

Таким образом, $h_A: U \rightarrow \{0,1\}$, то есть для любого элемента универсального множества U существуют только две возможности: он может либо принадлежать, либо не принадлежать множеству A .

Рассмотрим множество всех подмножеств U и обозначим его через $S(U)$. В нем можно ввести операции пересечения ($\cap$), объединения ($\cup$) и дополнения ($\neg$), как операции над характеристическими функциями:

$$\begin{aligned}
 h_{A \cap B}(u) &= \begin{cases} 1, & \text{если } h_A(u) = 1 \text{ и } h_B(u) = 1 \\ 0, & \text{иначе} \end{cases} \\
 h_{A \cup B}(u) &= \begin{cases} 1, & \text{если } h_A(u) = 1 \text{ или } h_B(u) = 1 \\ 0, & \text{иначе} \end{cases} \\
 h_{\neg A}(u) &= \begin{cases} 1, & \text{если } h_A(u) = 0 \\ 0, & \text{если } h_A(u) = 1 \end{cases} \\
 (A \in S(U), B \in S(U))
 \end{aligned}$$

Используя такие множества, мы можем определять различные понятия. Операции пересечения, объединения и дополнения могут интерпретироваться как логические связки «и», «или» и «не» соответственно. В этом случае мы говорим о булевой (двузначной) логике. Имея выражение некоторых понятий $a_1, a_2, \dots, a_n$ в виде множеств $A_1, A_2, \dots, A_n$, мы можем находить множества, соответствующие понятиям « a_1 и a_2 и a_n », «не a_1 и или a_2 и a_3 » и т. п.

Теория множеств и соответствующая ей булева логика составляют базу классической математики. Модели сложных технических и физических систем, химических процессов хорошо описывались на этом языке и удачно реализовывались на компьютерах. Речь могла только идти о недостаточном быстродействии, недостаточной памяти и других технических проблемах реализации этих моделей. Ситуация поменялась коренным образом, когда возникла необходимость учитывать особенности восприятия, оценки и анализа информации человеком как полноправной части моделируемой системы. Дело в том, что суждения и оценки человека являются приближенными и нечеткими, а компьютер может выполнять только четкие инструкции. Преодоление этого лингвистического барьера составляет сверхзадачу теории нечетких множеств.

Основная идея Заде заключалась в том, чтобы «разрешить» характеристической функции принимать не только значения 0 (полная непринадлежность) или 1 (полная принадлежность), но и промежуточные значения принадлежности из отрезка $[0,1]$. Таким образом, им было заменено понятие

характеристической функции $h_A(u)$ на понятие функции принадлежности $\mu_A: U \rightarrow [0,1]$.

Замечание 1. Характеристическая функция является частным случаем функции принадлежности. В этом смысле теория нечетких множеств является обобщением теории множеств.

Обозначим через $F(U)$ множество всех нечетких подмножеств универсального множества U . В нем можно ввести операции пересечения ($\cap$), объединения ($\cup$) и дополнения ($\neg$), как операции над характеристическими функциями, например так:

$$\mu_{A \cap B}(u) = \min(\mu_A(u), \mu_B(u))$$

$$\mu_{A \cup B}(u) = \max(\mu_A(u), \mu_B(u))$$

$$\mu_{\neg A}(u) = 1 - \mu_A(u)$$

Такие определения были введены Заде в упомянутой статье Fuzzy Sets.

Замечание 2. Если мы заменим в приведенных выше определениях функцию принадлежности на характеристическую, мы получим в точности операции пересечения, объединения и дополнения, принятые в теории множеств.

Можно ли определить эти операции как-то по-другому? Ответ положительный, например:

$$\mu_{A \cap B}(u) = \mu_A(u) \times \mu_B(u)$$

$$\mu_{A \cup B}(u) = \mu_A(u) + \mu_B(u) - \mu_A(u) \times \mu_B(u)$$

$$\mu_{\neg A}(u) = 1 - \mu_A(u).$$

Замечание 3. Если мы заменим в последних определениях функцию принадлежности на характеристическую, мы получим в точности операции пересечения, объединения и дополнения, принятые в теории множеств.

Используя такие множества, мы можем определять различные понятия. Операции пересечения, объединения и дополнения могут интерпретироваться как логические связи «и», «или» и «не» соответственно. В этом случае мы говорим о нечеткой логике. Имея выражение некоторых

понятий $a_1, a_2, \dots, a_n$ в виде множеств $A_1, A_2, \dots, A_n$, мы можем находить множества, соответствующие понятиям « a_1 и a_2 и a_3 », «не a_1 и или a_2 и a_5 » и т. п.

Эта простая идея дала толчок развитию большого числа исследований как вглубь (по изучению других возможных способов представления нечеткости и анализу их свойств¹), так и вширь (по применению нечетких моделей в управлении, системах принятия решений, распознавания образов и т. п.). В настоящее время это направление является интенсивно развивающейся научно-инженерной дисциплиной, по которой защищены сотни диссертаций, опубликованы тысячи книг и статей, реализованные на ее основе аппаратные средства используются в широком спектре оборудования от бытовой техники до космических систем.

С инженерной точки зрения теория нечетких множеств является инструментом моделирования специального вида неопределенности — лингвистической неопределенности (раздел 2.2).

Основной предпосылкой возникновения теории нечетких множеств является развитие вычислительной техники и связанное с ним развитие информационных технологий и искусственного интеллекта и попытки решения на этой базе задач в областях, традиционно относящихся к гуманитарным: медицина, управление, социология и др. Это нашло отражение в упоминавшейся в разделе 1.2.2 концепции экспертных систем, японском проекте создания ЭВМ пятого поколения, стратегической компьютерной инициативе Р. Рейгана и многих других мегапроектов, сформировавших современный ландшафт информационных технологий.

Ниже приводятся определения необходимых для дальнейшего понимания (но далеко не всех) понятий теории нечетких множеств и простые примеры, их иллюстрирующие. Читателю, знакомому с основами теории нечетких множеств, можно перейти к следующему разделу; читателю, которому

¹ Примерами вопросов в рамках этого направления являются: «Существуют ли другие, отличные от $[0,1]$, множества принадлежности?», «Существуют ли другие, отличные от приведенных, операции пересечения, объединения и дополнения?», «Сколько их?», «Чем они отличаются?» и пр.

захочется понять больше про эту науку, можно порекомендовать доступную для скачивания книгу [8].

Формальное определение нечеткого множества звучит следующим образом.

Пусть задано универсальное множество U элементов u и $\mu_A: U \rightarrow [0, 1]$. Нечетким подмножеством A в U называется график отображения μ_A , то есть множество вида $\{(u, \mu_A(u)): u \in U\}$; при этом значение $\mu_A(u)$ называется степенью принадлежности u к A .

Таким образом, задание нечеткого множества A в U эквивалентно заданию его функции принадлежности $\mu_A(u)$. Мы, следуя сложившейся традиции, будем употреблять термин «нечеткое множество» вместо более корректного термина «нечеткое подмножество».

В практических задачах использовались различные реализации таких функций принадлежности. Приведем некоторые из них.

Исторически первыми использовались и изучались так называемые s - и π - функции, задаваемые так:

$$S(u, \alpha, \beta, \gamma) = \begin{cases} 0 & \text{для } u \leq \alpha \\ 2 \left(\frac{u - \alpha}{\gamma - \alpha} \right)^2 & \text{для } \alpha \leq u \leq \beta \\ 1 - 2 \left(\frac{u - \gamma}{\gamma - \alpha} \right)^2 & \text{для } \beta \leq u \leq \gamma \\ 1 & \text{для } u \geq \gamma \end{cases}$$

$$\pi(u, \beta, \gamma) = \begin{cases} S\left(u, \gamma - \beta, \gamma - \frac{\beta}{2}, \gamma\right) & \text{для } u \leq \gamma \\ S\left(u, \gamma, \gamma + \frac{\beta}{2}, \gamma + \beta\right) & \text{для } u \geq \gamma. \end{cases}$$

Их графики представлены на рис. 9 и 10 соответственно.

Нетрудно заметить, что такое представление было попыткой представить переход от непринадлежности к принадлежности в виде гладких функций. Однако оно не нашло дальнейшего развития по ряду причин и сейчас используется

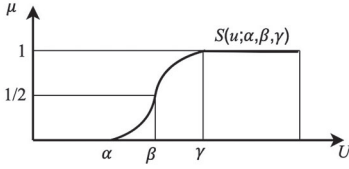


Рис. 9. Пример s- функции

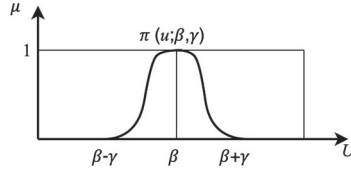


Рис. 10. Пример π- функции

только в демонстрационных/учебных целях. Мы будем также иногда их использовать далее. Примером более общего представления таких функций являются так называемые множества $(L - R)$ -типа. Их функция принадлежности задается с помощью функций L и R , определенных на числовой прямой и удовлетворяющих следующим требованиям:

1. $L(0) = R(0) = 1$.
2. $L(-x) = L(x)$, $R(-x) = R(x)$.
3. L и R — невозрастающие функции на множестве неотрицательных действительных чисел.

Примерами таких функций являются:

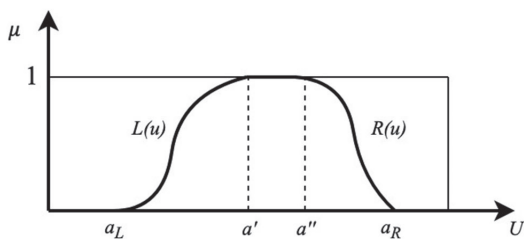
$$L(u) = e^{-|u|^p}, p \geq 0,$$

$$R(u) = \frac{1}{1 + |u|^p}, p \geq 0,$$

$$L(u) = \begin{cases} 1 & \text{при } u \in [-1, 1] \\ 0 & \text{в противном случае.} \end{cases}$$

Функция принадлежности нечеткого множества $(L - R)$ -типа задается так:

$$\mu_A(u) = \begin{cases} L\left(\frac{a' - u}{a_L}\right) & \text{при } u \leq a', a_L \geq 0 \\ R\left(\frac{u - a''}{a_R}\right) & \text{при } u \leq a'', a_R \geq 0 \\ 1 & \text{при } u \in [a', a'']. \end{cases}$$


 Рис. 11. Пример функции принадлежности $(L - R)$ -типа

Пример такой функции представлен на рис. 11.

На практике часто используют частный случай таких функций принадлежности, когда функции L и R являются линейными. В этом случае имеем:

$$\mu_A(u) = \begin{cases} 0 & \text{при } u \leq a_L \\ \frac{u - a_L}{a' - a_L} & \text{при } a_L \leq u \leq a' \\ 1 & \text{при } a' \leq u \leq a'' \\ \frac{a_R - u}{a_R - a''} & \text{при } a'' \leq u \leq a_R \\ 0 & \text{при } u \geq a_R \end{cases}$$

Линейные функции принадлежности $(L - R)$ -типа при $a' < a''$ называются трапецеидальными, при $a' = a''$ — треугольными. Именно они чаще всего используются в нечетких моделях, и мы будем чаще всего использовать их далее.

Обобщение понятия нечеткого множества, введенного Заде, коснулось прежде всего множества, описывающего степень принадлежности, то есть $[0, 1]$. Использование этого отрезка вполне понятно: 0 — полная непринадлежность, 1 — полная принадлежность. Однако это не всегда так. Например, в первой промышленной экспертной системе MYCIN (раздел 1.2.2), как и в любой серьезной системе такого типа, блок обработки лингвистической неопределенности работал с мно-

жеством неопределенностей $[-1,1]$. К настоящему времени возникло и используется много таких вариаций, например:

- нечеткие множества в смысле Гогена $\mu_A: U \rightarrow L$, где L — конечная или бесконечная дистрибутивная решетка;
- P — нечеткие множества $\mu_A: U \rightarrow P$, где P — множество подмножеств отрезка $[0,1]$;
- нечеткие множества типа n -множества, у которых значениями функции принадлежности являются множества типа $(n-1)$.

Общим требованием к множеству принадлежностей является: мы должны иметь минимальный и максимальный элементы и для любых двух различных элементов уметь определить, который из них больше. С учетом сказанного, общее определение нечеткого множества формулируется следующим образом.

Пусть U — универсальное множество, M — множество принадлежностей, $\mu_A: U \rightarrow M$. Тогда нечетким множеством A в U называется график отображения μ_A , то есть множество вида $\{(u, \mu_A(u)): u \in U\}$. Таким образом, нечеткое множество A задается тройкой (U, M, μ_A) .

Ниже мы будем считать $M=[0,1]$ — этого вполне хватает для данной книги.

Вторым направлением обобщения предложенного Заде формализма явилось обобщение операций пересечения, объединения и дополнения. Кроме упомянутых в лекции 1 примеров

$$\mu_{A \cap B}(u) = \min(\mu_A(u), \mu_B(u))$$

$$\mu_{A \cup B}(u) = \max(\mu_A(u), \mu_B(u))$$

$$\mu_{\neg A}(u) = 1 - \mu_A(u)$$

и

$$\mu_{A \cap B}(u) = \mu_A(u) \times \mu_B(u)$$

$$\mu_{A \cup B}(u) = \mu_A(u) + \mu_B(u) - \mu_A(u) \times \mu_B(u)$$

$$\mu_{\neg A}(u) = 1 - \mu_A(u),$$

было предложено много других, для которых замена функции принадлежности на характеристическую давала в точно-

сти операции пересечения, объединения и дополнения, принятые в теории множеств (замечание 3 выше). К настоящему времени общим представлением таких операций считаются так называемые треугольные нормы (t -нормы), треугольные конормы (t -конормы) и обобщенные дополнения. Они определяются следующим образом.

Треугольной нормой (сокращенно -нормой) называется двухместная действительная функция $T: [0,1] \times [0,1] \rightarrow [0,1]$, удовлетворяющая следующим условиям:

- 1) $T(0,0) = 0$, $T(\mu_A, 1) = \mu_A$ (ограниченность);
- 2) $T(\mu_A, \mu_B) = T(\mu_B, \mu_A)$ (коммутативность);
- 3) $T(\mu_A, T(\mu_B, \mu_C)) = T(T(\mu_A, \mu_B), \mu_C)$ (ассоциативность);
- 4) $T(\mu_A, \mu_B) \leq T(\mu_C, \mu_D)$, если $\mu_A \leq \mu_C$ и $\mu_B \leq \mu_D$ (монотонность).

Простыми случаями треугольных норм являются приведенные выше.

$T_m(\mu_A, \mu_B) = \min(\mu_A, \mu_B)$ (иногда называемая логическим произведением);

$T_p(\mu_A, \mu_B) = \mu_A \cdot \mu_B$ (алгебраическое произведение), а также

$T_i(\mu_A, \mu_B) = \max\{0, \mu_A + \mu_B - 1\}$ (ограниченное произведение, или -норма Лукасевича);

$$T_d(\mu_A, \mu_B) = \begin{cases} \mu_A, & \text{если } \mu_B = 1 \\ \mu_B, & \text{если } \mu_A = 1 \text{ (драстическое произведение)} \\ 0, & \text{в других случаях.} \end{cases}$$

Можно показать, что эти -нормы упорядочены следующим образом: $T_d \leq T_i \leq T_p \leq T_m$. Более того, для любой -нормы T справедливо: $T_d \leq T \leq T_m$.

Треугольной конормой (сокращенно -конормой) называется двухместная действительная функция $\perp: [0,1] \times [0,1] \rightarrow [0,1]$, удовлетворяющая следующим условиям:

- 1) $\perp(1,1) = 1$, $\perp(\mu_A, 0) = \mu_A$ (ограниченность);
- 2) $\perp(\mu_A, \mu_B) = \perp(\mu_B, \mu_A)$ (коммутативность);
- 3) $\perp(\mu_A, \perp(\mu_B, \mu_C)) = \perp(\perp(\mu_A, \mu_B), \mu_C)$ (ассоциативность);
- 4) $\perp(\mu_A, \mu_B) \geq \perp(\mu_C, \mu_D)$, если $\mu_A \geq \mu_C$ и $\mu_B \geq \mu_D$ (монотонность).

Простыми случаями треугольных -конорм являются приведенные выше.

$\perp_m(\mu_A, \mu_B) = \max \{\mu_A, \mu_B\}$ (иногда называют логическая сумма);

$\perp_p(\mu_A, \mu_B) = \mu_A + \mu_B - \mu_A \times \mu_B$ (вероятностная сумма), а также

$\perp_i(\mu_A, \mu_B) = \min \{\mu_A + \mu_B, 1\}$ (ограниченная сумма, или -конорма Лукасевича);

$$\perp_d(\mu_A, \mu_B) = \begin{cases} \mu_A, & \text{если } \mu_B = 0 \\ \mu_B, & \text{если } \mu_A = 0 \text{ (драстическая сумма)} \\ 0, & \text{в других случаях.} \end{cases}$$

Справедливы соотношения, аналогичные приведенным выше для -норм: $\perp_d \geq \perp_i \geq \perp_p \geq \perp_m$, и для любой t -конормы справедливо $\perp_d \geq \perp \geq \perp_m$.

t -нормы и -конормы можно задавать параметически. Их бесконечно много, и все они в частном случае классической теории множеств (то есть при замене функций принадлежности на характеристические функции) дают общепринятые определения операций пересечения и объединения.

Обобщенным дополнением называется функция $c: [0,1] \rightarrow [0,1]$, удовлетворяющая следующим условиям:

- 1) $c(0) = 1, c(1) = 0$;
- 2) c — невозрастающая функция, то есть если $\mu_A \leq \mu_B$, то $c(\mu_A) \geq c(\mu_B)$.

Простыми случаями дополнения являются приведенная выше $c(\mu_A) = 1 - \mu_A$, а также $c_r(\mu_A) = \sqrt{1 - \mu_A^2}$,

$$c_\lambda(\mu_A) = \frac{1 - \mu_A}{1 + \lambda \mu_A} \quad (-1 \leq \lambda \leq \infty),$$

$$c_\alpha(\mu_A) = \begin{cases} 1 & \text{при } \mu_A \leq \alpha \\ 0 & \text{при } \mu_A > \alpha. \end{cases}$$

Кроме перечисленных понятий нечеткого множества и теоретико-множественных операций для множества всех нечетких подмножеств универсального множества нам понадобится понятие лингвистической переменной, введенное Заде, и связанное с ним понятие полных ортогональных семантических пространств.

Нечеткой переменной называется тройка (α, U, G) , где α — наименование (имя) нечеткой переменной; U — универсальное множество (область определения); G — ограничения на возможные значения (смысл) переменной представленные в виде $\mu_A(U)$.

Таким образом, нечеткая переменная — это поименованное нечеткое множество.

Лингвистическая переменная представляет собой пятерку $(A, T(A), U, V, M)$, где A — наименование (имя) лингвистической переменной; $T(A)$ — терм-множество лингвистической переменной A , то есть множество ее лингвистических значений; U — универсальное множество, в котором определяются значения лингвистической переменной A ; V — синтаксическое правило, порождающее значения лингвистической переменной A (часто имеет форму грамматики); M — семантическое правило, которое ставит в соответствие каждому элементу $T(A)$ его «смысл» как нечеткое подмножество U .

Пример лингвистической переменной приведен на рис. 12.

Для задания правила V обычно задаются *базовые значения* (например, маленький, средний, большой), *модификаторы* (не, очень, слегка и т. п.) и *правила образования элементов терм-множества* (немаленький, очень маленький, не очень маленький и т. п.).

При задании правила M предполагается, что функции принадлежности базовых терминов заданы, а также известно, каким образом модификаторы воздействуют на функции

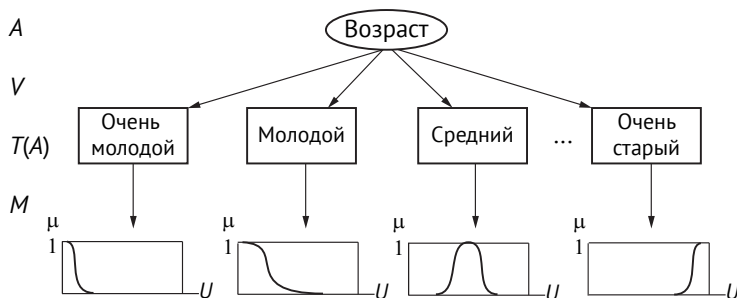


Рис. 12. Пример лингвистической переменной «Возраст»

принадлежности базовых значений. Например, предполагается, что если известна функция принадлежности $\mu_A(u)$ некоторого базового значения A , то функции принадлежности лингвистических значений «не A », «очень A » могут вычисляться следующим образом:

$$\mu_{\text{не } A}(u) = 1 - \mu_A(u), \quad \mu_{\text{очень } A}(u) = \mu_A^2(u).$$

Предполагалось, что, задав функции принадлежности нескольких базовых значений и выбрав формализации логических операций «и», «или», «не», мы сможем вычислять функции принадлежности любых элементов терм-множества $T(A)$. Более того, мы сможем вычислять «смысл» любого высказывания, производя соответствующие операции над функциями принадлежности понятий различных лингвистических переменных. Этот «смысл» может быть представлен либо в виде функции принадлежности, либо в виде лингвистического значения, наилучшим способом аппроксимирующей результирующую функцию принадлежности.

Однако реализовать данную программу в полном объеме не удалось. Исследования, проведенные в этой области, составили базу для большого числа теоретических и практических результатов. Они получили новый импульс в свете развиваемой Заде в течение последних лет идеи грануляционных вычислений (granularity computing) и вычислений на словах (computing with words), а также упоминавшихся выше дополненного или гибридного интеллекта (augmented intelligence) и automation of knowledge work.

Одна из проблем разработки теории лингвистической переменной в полном объеме заключалась в слишком широком ее определении. Известно, что чем меньше ограничений в некоторой теории, тем труднее получить для нее значимые результаты. Кроме этого семантика (функция принадлежности) термина зависит от контекста. Рассмотрим терм-множество T_2 лингвистической переменной «возраст», содержащее два значения — «молодой» и «старый» — с функциями принадлежности $\mu_1(u)$ и $\mu_2(u)$ соответственно (рис. 13). Теперь рассмотрим терм-множество T_3 с тремя значениями, отличающееся от T_2 значением «средний» с функцией принадлежности $\mu_3(u)$. Такое из-

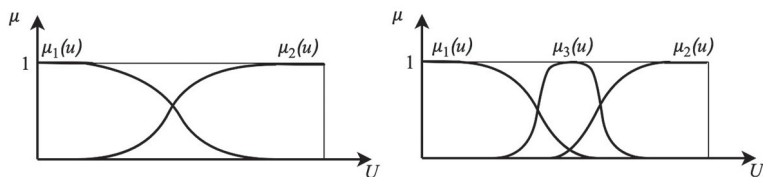


Рис. 13. Модификация функций принадлежности

менение T_2 приводит к модификации функций принадлежности (см. рис. 13). Таким образом, функции принадлежности $\mu_1(u)$ и $\mu_2(u)$ в контексте T_2 отличаются от функции принадлежности $\mu_1(u)$ и $\mu_2(u)$ в контексте T_3 . Получается, что понятие функции принадлежности без указания контекста вообще не имеет смысла. Поэтому зафиксируем контекст.

Частным случаем лингвистической переменной, свободным от этого недостатка и имеющим широкий спектр практических приложений, может быть набор нечетких переменных, описывающих некоторое понятие (то есть лингвистическая переменная с фиксированным терм-множеством). Такие структуры можно интерпретировать, в частности, как набор лингвистических значений, описывающих степень противоречивости полученной информации или имеющейся возможности реализации некоторых процессов. В теории распознавания образов это можно интерпретировать как множество значений качественного признака.

Рассмотрим t нечетких переменных с именами $a_1, a_2, \dots, a_t$, заданных на одном универсальном множестве (рис. 14). Назовем такую совокупность *семантическим пространством* S_t .

Таким образом, семантическое пространство — это набор нечетких переменных $S_t = \{(\alpha_1, U, G_1), \dots, (\alpha_t, U, G_t)\}$, определенных в одном универсальном множестве. Для одной и той же лингвистической переменной A могут существовать различные семантические пространства:

$$\begin{aligned} S_2 &= \{\langle \alpha_1, U, G_1 \rangle, \dots, \langle \alpha_t, U, G_t \rangle, \dots, S_k = \\ &= \{\langle \alpha_1, U, G_1 \rangle, \dots, \langle \alpha_k, U, G_k \rangle\}. \end{aligned}$$

Теперь мы сможем перейти к решению проблем 1 и 2 (раздел 2.2).

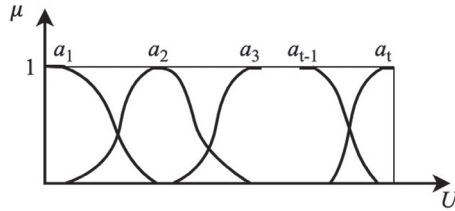


Рис. 14. Пример семантического пространства

2.4. ПРОБЛЕМА 1 — МОДЕЛИРОВАНИЕ ОПИСАНИЯ ЧЕЛОВЕКОМ ОБЪЕКТОВ

Можно ли сравнивать такие семантические пространства, выбирать в каком-то смысле лучшие из них?

Рассмотрим пример.

Пример 1. Пусть нас просят описать рост человека. Рассмотрим две крайние ситуации.

Ситуация 1. Разрешено использовать только два значения: «низкий» и «высокий».

Ситуация 2. Разрешено использовать много значений: «очень низкий», «не очень высокий», ... «не высокий и не низкий», ... «очень-очень высокий».

Ситуация 1 является неудобной. Действительно, для многих людей оба разрешенных значения могут не подходить, и, описывая их, мы выбираем между двумя «плохими» значениями.

Ситуация 2 также является неудобной. Действительно, при описании конкретного человека могут подходить несколько из разрешенных значений. Мы опять испытываем дискомфорт, но теперь от того, что вынуждены выбирать между двумя или более «хорошими» значениями. Может ли быть оптимальное в этом смысле множество лингвистических значений?

Приведенный пример позволяет сформулировать уточнение проблемы 1 как проблему выбора множества значений качественного признака: можно ли, учитывая некоторые особенности восприятия человеком объектов реального мира и их описание, сформулировать правило выбора оптималь-

ного множества значений признаков, по которым описываются эти объекты? Возможны два критерия оптимальности.

Критерий 1. Под оптимальными понимаются такие множества значений, используя которые человек испытывает минимальную неопределенность при описании объектов.

Критерий 2. Если объект описывается некоторым количеством экспертов, то под оптимальными понимаются такие множества значений, которые обеспечивают минимальную степень рассогласования описаний.

Эта проблема носит фундаментальный характер для человеко-компьютерных систем и касается соотношения реального объекта или процесса и его модели (раздел 2.2). Мы всегда хотим, чтобы модель как можно более точно соответствовала объекту или процессу: то, что мы получим в модели, переносим в реальный мир (рис. 15а). Для естественных наук (физика, химия) были разработаны измерительные приборы, обеспечивающие адекватность модели объекту (рис. 15б). Для социальных процессов, социальных сетей, дополненного интеллекта и пр. таким измерительным прибором является человек. Таким образом, проблема 1 — это проблема разработки оптимального измерительного прибора для человеко-компьютерных систем. Такой прибор обеспечивает максимальную адекватность модели реальному объекту или процессу: количество ситуаций, когда одному реальному объекту соответствуют несколько объектов в модели или, наоборот, разным реальным объектам соответствует один объект в модели, минимально.

Проблема 1 имеет решение не для произвольных семантических пространств, а для их подмножества — полных ортогональных семантических пространств, определяемых далее.

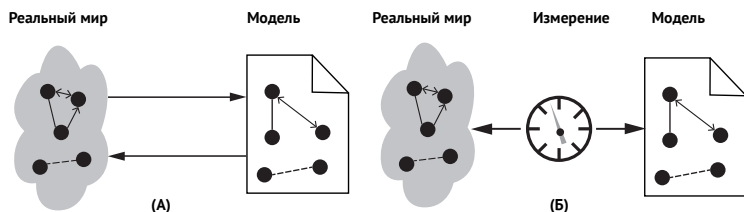


Рис. 15. Модель и объект

Введем систему ограничений для функций принадлежности нечетких переменных, составляющих s_t . Для простоты обозначим функцию принадлежности a_j через μ_j . Будем считать, что:

1. $\forall \mu_j (1 \leq j \leq t) \exists U_j^1 \neq \emptyset$, где $U_j^1 = \{u \in U : \mu_j = 1\}$, U_j^1 есть отрезок или точка;
2. $\forall j (1 \leq j \leq t) \mu_j$ не убывает слева от U_j^1 и не возрастает справа от U_j^1 (так как, согласно 1, U_j^1 является отрезком или точкой, понятия «слева» и «справа» определяются однозначно).

Требования 1 и 2 естественны для функций принадлежности понятий, образующих семантическое пространство. Действительно, первое означает, что для любого используемого понятия в универсальном множестве существует хотя бы один объект, который является эталонным для данного понятия. Если таких эталонов много, то они расположены рядом, а не «раскиданы» по универсуму. Второе требование означает, что если объекты близки в смысле метрики в универсальном множестве, то они также близки в смысле принадлежности к некоторому понятию.

Наряду с функциями принадлежности будем использовать и характеристические функции, поэтому потребуем выполнения следующего технического условия:

3. $\forall j (1 \leq j \leq t) \mu_j(u)$ имеет не более двух точек разрыва первого рода.

Для простоты обозначим требования 1–3 через L .

Введем также систему ограничений для совокупностей функций принадлежности нечетких переменных, образующих s_t , то есть будем считать, что:

4. $\forall u \in U \exists j (1 \leq j \leq t) \mu_j(u) > 0$;
5. $\forall u \in U \sum_{j=1}^t \mu_j(u) = 1$.

Требования 4 и 5 также имеют довольно естественную интерпретацию. Требование 4, называемое *полнотой*, означает, что для любого объекта из универсального множества существует хотя бы одно понятие, к которому он может принадлежать.

Это означает, что в нашем семантическом пространстве нет «дыр». Требование 5, называемое *ортогональностью*, означает, что мы не допускаем использования семантически близких понятий или синонимов, требуем достаточной различимости используемых понятий. Отметим также, что это требование часто выполняется или нет в зависимости от используемого метода построения функций принадлежности понятий, образующих семантическое пространство. Так, например, если у нас есть некоторое количество экспертов, и мы им предъявляем объект $u \in U$ и допускаем только ответы «Да, $u \in a_j$ » и «Нет, $u \notin a_j$ » (ответа «Не знаю» не допускается), а в качестве значения функции принадлежности $\mu_j(u)$ берем отношение числа экспертов, ответивших положительно, к общему числу экспертов, то данное требование выполняется автоматически. Заметим также, что все приводимые ниже результаты справедливы при некотором ослаблении требования ортогональности, но для его описания необходимо введение ряда дополнительных понятий. Поэтому остановимся на этом требовании.

Для простоты обозначим требования 4, 5 через G .

Назовем семантическое пространство, состоящее из нечетких переменных, функции принадлежности которых удовлетворяют требованиям 1–3, а их совокупности — требованиям 4 и 5, *полным ортогональным семантическим пространством* и обозначим его $G(L)$.

Для $G(L)$ можно определить понятие меры неопределенности (степени нечеткости).

Пусть есть некоторая совокупность из t функций принадлежности $s_i \in G(L)$ и пусть $s_i = \{\mu_1(u), \dots, \mu_t(u)\}$. Будем называть совокупность из t характеристических функций $\hat{s}_i = \{h_1(u), \dots, h_t(u)\}$ *ближайшей совокупностью характеристических функций*, если

$$h_j(u) = \begin{cases} 1, & \text{если } \mu_j(u) = \max_{1 \leq i \leq t} \mu_i(u) \\ 0, & \text{иначе } (1 \leq j \leq t). \end{cases}$$

Нетрудно заметить, что если полное ортогональное семантическое пространство состоит не из функций принадлежности, а из характеристических функций (рис. 5), то при описании объектов в нем не возникает никакой неопределенности.

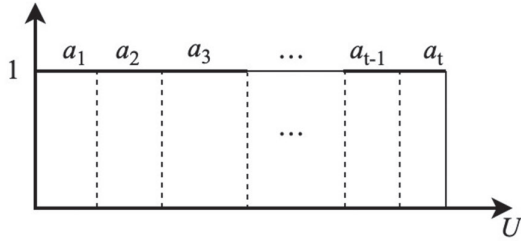


Рис. 16. Полное ортогональное семантическое пространство характеристических функций

Эксперт однозначно выбирает термин a_j , если объект находится в соответствующей области универсального множества. Несколько экспертов описывают один и тот же объект одним и тем же термином. Эту ситуацию можно проиллюстрировать следующим образом. Допустим, есть весы достаточной точности, и мы имеем возможность взвешивать некоторый материал. Кроме того, мы договорились, что если вес материала попадает в определенный интервал, то он принадлежит одной из категорий. Тогда мы в точности будем иметь описанную ситуацию. Проблема заключается в том, что для нашей задачи нет таких весов, а также возможности «взвешивать» на них интересующие нас объекты.

Однако мы можем предположить, что из двух семантических пространств то имеет меньшую неопределенность, которое более похоже на пространство из совокупностей характеристических функций.

В математике степенью сходства может быть расстояние. Можно ли ввести расстояние между семантическими пространствами? Для полных ортогональных семантических пространств это возможно.

Так как мы рассматриваем полные ортогональные семантические пространства, функции принадлежности используемых понятий удовлетворяют условиям L . Это подмножество известного пространства функций, интегрируемых на некотором отрезке U . Известно, что в нем можно ввести расстояние, например, следующим образом:

$$d(f, g) = \int_U |f(u) - g(u)| du.$$

Расстояние в $G(L)$ можно ввести с помощью леммы 1.

Лемма 1. Пусть $s_t = \{\mu_1(u), \dots, \mu_t(u)\}$, $s'_t = \{\mu'_1(u), \dots, \mu'_t(u)\}$; $s_t, s'_t \in G(L)$; $\rho(\mu, \mu')$ — метрика в L . Тогда $d(s_t, s'_t) = \sum_{j=1}^t \rho(\mu_j, \mu'_j)$ и есть метрика в $G(L)$.

Доказательство леммы приведено в [8].

Содержательные замечания при анализе рис. 16 можно сформулировать следующим образом.

Пусть $s_t \in G(L)$. Под мерой неопределенности s_t будем понимать значение функционала $\xi[0,1]$, определенного на элементах $G(L)$ и принимающего значения в $[0,1]$ (то есть $\xi: G(L) \rightarrow [0,1]$), удовлетворяющего следующим условиям (аксиомам):

A1. $\xi(s_t) = 0$, если s_t представляет собой совокупность характеристических функций.

A2. Пусть $s_p, s'_p \in G(L)$, t и t' могут быть равны или неравны друг другу. Тогда $\xi(s_p) \leq \xi(s'_p)$, если $d(s_p, \hat{s}_t) \leq d(s'_p, \hat{s}'_t)$. (Напомним, что $\hat{s}_t$ — ближайшая к s_t совокупность характеристических функций.)

Существуют ли такие функционалы? Ответ на этот вопрос дает теорема существования.

Теорема 1 (теорема существования). Пусть $s_t \in G(L)$. Тогда функционал

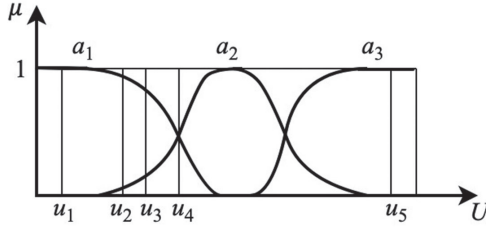
$$\xi(s_t) = \frac{1}{|U|} \int_U f(\mu_{i_1}^*(u) - \mu_{i_2}^*(u)) du,$$

где $\mu_{i_1}^*(u) = \max_{1 \leq j \leq t} \mu_j(u)$, $\mu_{i_2}^*(u) = \max_{1 \leq j \leq t, j \neq i_1} \mu_j(u)$, f удовлетворяет следующим условиям:

F1: $f(0) = 1$, $f(1) = 0$; F2: f не возрастает — является мерой неопределенности s_p , то есть удовлетворяет аксиомам A1 и A2.

С доказательством можно ознакомиться в [8].

Достаточно очевидно, что существует только одна линейная функция, удовлетворяющая F1 и F2 — это функция $f(x) = 1 - x$. В классе полиномов второй степени — это параметрическое семейство $f_a(x) = ax^2 - (1+a)x + 1$. Подмножества функций других типов (логарифмических, тригонометрических и т. п.) могут быть описаны аналогичным образом. Подставляя эти функции в формулу $\xi(s_p)$, мы получаем различные степени нечеткости.

Рис. 17. Семантическое пространство s_3

Рассмотрим простейший класс таких функционалов, когда подынтегральная функция является линейной:

$$\xi(s_t) = \frac{1}{|U|} \int_U f(1 - (\mu_{i1^*}(u) - \mu_{i2^*}(u))) du.$$

Интерпретация. Рассмотрим процесс описания объектов в рамках семантического пространства s_3 (рис. 17).

Для объектов u_1 и u_5 человек без всяких колебаний выбирает один из термов (a_1 и a_3 соответственно). При описании объекта u_2 человек начинает выбирать между термами a_1 и a_2 . Такого рода колебания возрастают и достигают своего пика при описании объекта u_4 : для этой точки термы a_1 и a_2 неразличимы. Если вспомнить процедуру построения функций принадлежности, описанную при анализе свойства ортогональности (лекция 9), мы можем также утверждать, что все эксперты будут единодушны при описании объектов u_1 и u_5 и иметь некоторую рассогласованность при описании u_2 , которая будет возрастать и достигнет своего пика для объекта u_4 .

Теперь рассмотрим подынтегральную функцию. Нетрудно увидеть, что

$$\begin{aligned} 0 = \eta(s_3, u_1) = \eta(s_3, u_5) &< \eta(s_3, u_2) < \\ &< \eta(s_3, u_3) < \eta(s_3, u_4) = 1. \end{aligned}$$

Таким образом, $\eta(s_3, u_i)$ действительно отражает степень неопределенности, которую испытывает человек при описании объектов в рамках семантического пространства, или степень рассогласования мнения экспертов при таком описании, а $\xi(s_t)$ — усредненная степень трудности описания ре-

альных объектов (выбора того или иного значения) в рамках s_t .

Справедлива следующая теорема [8].

Теорема 2. Пусть $s_t \in G(L)$. Тогда $\xi(s_t) = c \frac{d}{|\bar{U}|}$, где $d = |\bar{U}|$; $c = \text{Const}$, $c > 1$; $\bar{U} = \{u \in U : \exists j(1 \leq j \leq t) 0 < \mu_j(u) < 1\}$.

$\bar{U}$ — это мощность подмножества универсального множества U , где имеется неопределенность (функции принадлежности не равны нулю или единице). Таким образом, степень нечеткости $\xi(s_t)$ пропорциональна отношению размера области неопределенности к размеру всего универсального множества.

Рассмотрим еще одно важное свойство степени нечеткости — ее инвариантность к линейным преобразованиям универсального множества. Пусть g — некоторая взаимно-однозначная функция, определенная на U . Эта функция индуцирует преобразование некоторого полного ортогонального семантического пространства $s_t \in G(L)$, определенное на универсуме U , в полное ортогональное семантическое пространство $g(s_t)$, определенное на универсуме U' , где $U' = g(U) = \{u' = g(u), u \in U\}$.

Данное преобразование может быть определено следующим образом: $g(s_t)$ есть множество функций принадлежности $\mu'_j(u') = \mu'_j(g(u)) = \mu_j(g^{-1}(u')) = \mu_j(u)$, где $\mu_j(u) \in s_t, 1 \leq j \leq t$.

Это определение можно проиллюстрировать следующим примером.

Пример 2. Пусть $s_t \in G(L)$, U — универсум s_t и g есть растяжение (сжатие) универсума U . Тогда $g(s_t)$ есть множество функций, получаемых из s_t таким же растяжением (сжатием) — рис. 18.

Теорема 3 (о линейных преобразованиях). Пусть $s_t \in G(L)$, U — универсум s_t , g — некоторая линейная функция, определенная на U и $\xi(s_t) \neq 0$. Тогда $\xi(s_t) = \xi(g(s_t))$.

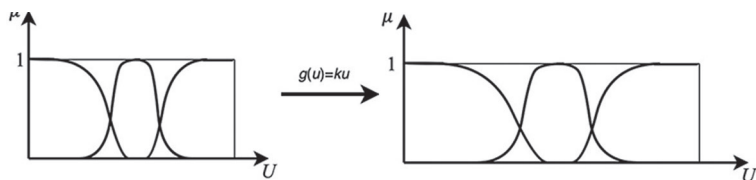


Рис. 18. Линейное преобразование семантического пространства

Доказательство приведено в [8].

Указанное свойство означает, что человек описывает различные объекты, используя некоторое множество лингвистических значений, с одинаковыми трудностями, если физические параметры объектов одного типа могут быть получены из параметров объектов другого типа некоторым линейным преобразованием. Например, используя термины «высокий», «средний», «низкий», мы описываем людей, деревья, здания и т. п. с одинаковыми трудностями; используя множество термов «очень близко», «близко», «не близко», «очень далеко», мы описываем расстояние между молекулами, расстояние между улицами в городе, расстояние между городами на карте с одинаковыми трудностями. Это подтверждает известную в лингвистике гипотезу лингвистической относительности Сепира–Уорфа [8].

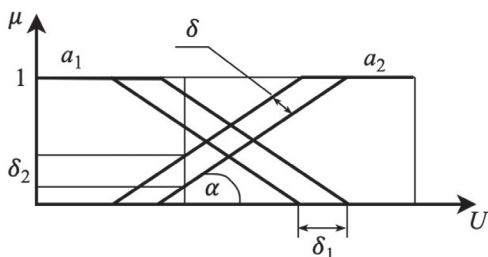
Степень нечеткости одного множества, индуцированная $\xi(s_i)$, может быть определена как степень нечеткости тривиального полного ортогонального семантического пространства, определяемого одним нечетким множеством $\mu(u)$:

$$\xi(\mu) = \frac{1}{|U|} \int_U (1 - |2\mu(u) - 1|) du.$$

Это — известная степень нечеткости (одного) нечеткого множества, что позволяет утверждать, что описанное здесь обобщение степени нечеткости множества $\xi(s_i)$ определено корректно.

Важным вопросом любого практического применения любой модели является ее устойчивость к естественным девиациям параметров (устойчивость по Ляпунову или робастность в англоязычной терминологии). Рассмотрим этот вопрос подробнее.

Функции принадлежности (семантика понятий) различных людей могут отличаться. Например, 50 лет для юноши — это старость, для пожилого человека — средний возраст; то же самое для оценки роста низким и высоким человеком — функции принадлежности сдвинуты относительно друг друга. Подобные ситуации давно подмечены и послужили основой многих детективных романов.


 Рис. 19. δ -модель

Более сложно взаимодействие между семантикой одинаковых терминов у людей, проживающий в разных геоклиматических зонах: чукчи и эскимосы имеют несколько названий белого (соответствующих разным состояниям снега), англичане имеют одно слово для синего цвета (blue), русские — два (синий и голубой) и т. п. Это называется принцип лингвистической дополнительнойности — один из базовых принципов психоллингвистики. Здесь мы не будем изучать такие сложные взаимодействия, ограничимся более простыми, связанными с возрастом, ростом и другими физическими параметрами наблюдателя, описанными выше.

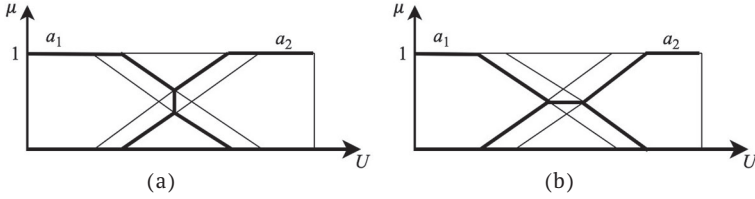
Для учета перечисленных факторов будем считать, что функции принадлежности могут находиться в некоторой «трубе» шириной δ (рис. 19). Назовем эту ситуацию δ -моделью и обозначим $G^\delta(L)$.

Для простоты рассмотрим не весь класс L функций принадлежности, а его подмножество $\bar{L}$, содержащее только линейные функции принадлежности (как изображено на рис. 19).

Для $s_i \in G(\bar{L})$ мы можем непосредственно посчитать $\xi(s_i)$.

В этом случае $\xi(s_i) = \frac{d}{2|U|}$, где d — мощность области неопределенности, как и в теореме 2 (то есть константа $c = 1/2$)).

В этой ситуации можно вычислить нижние ($\underline{\xi}(s_i)$) и верхние ($\bar{\xi}(s_i)$) оценки степени неопределенности $\xi(s_i)$. Совокупности нечетких множеств из $G^\delta(L)$, имеющие минимальную ($\underline{\eta}(s_i, u)$) и максимальную ($\bar{\eta}(s_i, u)$) степени неопределенности, представлены на рис. 20(a) и 20(b) соответственно.

Рис. 20. Максимально четкие и максимально нечеткие $s_2 \in G^*(L)$

Теорема 4. Пусть $s_t \in G^{\delta}(\bar{L})$.

$$\text{Тогда } \underline{\xi}(s_t) = \frac{d(1-\delta_2)^2}{2|U|}, \quad \bar{\xi}(s_t) = \frac{d(1+\delta_2)}{2|U|}.$$

С доказательством теоремы можно ознакомиться в [8].

Сравнивая теорему 2 для $s_t \in G(\bar{L})$ ($\xi(s_t) = \frac{d}{2|U|}$) и теорему 4, получаем, что при малых значениях δ закономерность модели сохраняется, то есть она является устойчивой (робастной).

Суммируя интерпретацию и свойства $\xi(s_t)$ (теоремы 2–4), ее связь с общепринятой степенью нечеткости множества, мы можем сформулировать следующий метод выбора оптимального (в смысле формулировки проблемы 1) множества значений качественного признака:

- 1) генерируются все «разумные» множества значений лингвистической переменной;
- 2) каждое из таких множеств представляется в форме полного ортогонального семантического пространства;
- 3) для каждого из них вычисляется мера неопределенности $\xi(s_t)$;
- 4) в качестве оптимального множества значений как с точки зрения минимизации неопределенности описания объектов, так и с точки зрения минимизации степени рассогласования мнений экспертов выбирается то множество, мера неопределенности которого минимальна — это $s_t^* : \xi(s_t^*) = \min_{s_t \in G(L)} \xi(s_t)$.

Следуя этому методу, мы можем описывать объекты с минимально возможной неопределенностью, то есть *гарантировать оптимальность* свойств систем гибридного интеллекта

с этой точки зрения. Устойчивость $\xi(s_i)$ (теорема 4) позволяет использовать этот метод в практических задачах.

Возвращаясь к анализу проблемы 1, можно сказать, что нами разработан оптимальный измерительный прибор для человеко-компьютерных систем, обеспечивающий максимальную адекватность модели реальному объекту или процессу: количество ситуаций, когда одному реальному объекту соответствуют несколько объектов в модели или, наоборот, разным реальным объектам соответствует один объект в модели, минимально.

2.5. ПРОБЛЕМА 2 — ОБРАБОТКА НЕЧЕТКИХ ОПИСАНИЙ ОБЪЕКТОВ

Итак, мы можем описывать объекты оптимальным в смысле критериев 1 и 2 образом, что обеспечивает максимальную адекватность таких описаний реальным объектам и процессам среды. Что это дает для обработки информации? Далее рассмотрим простейшую задачу — поиск информации.

Разберем ситуацию, когда описания человеком объектов / процессов реального мира образуют базу данных системы. Для пользователя этой базы данных она является информационной моделью реального мира — результаты работы он переносит на реальный мир (рис. 21).

Распространенным примером такой ситуации является социальная сеть (рис. 22).

Что нам дает описание объектов по предложенному выше методу с точки зрения поиска информации?

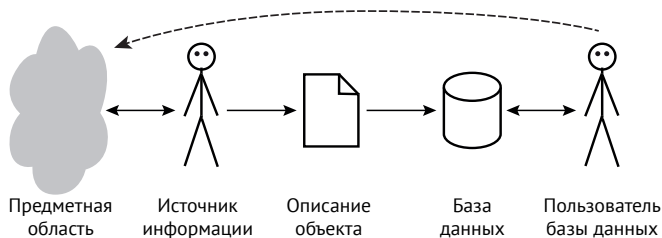


Рис. 21. База данных как информационная модель среды

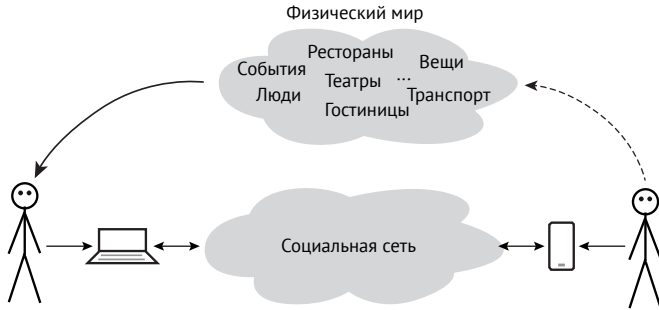


Рис. 22. Социальная сеть как информационная модель реального мира

Методы и алгоритмы поиска информации в нечетких (лингвистических) базах данных анализируются в работе [7]. Вводятся понятия потерь информации ($P_x(U)$) и информационных шумов ($H_x(U)$), возникающих при поиске информации в лингвистических базах данных. Смысл этих понятий следующий. При общении с системой пользователь формулирует запрос, содержащий определенные значения лингвистических признаков, и получает ответ на запрос. Если бы он мог знать физические (не лингвистические) значения признаков, он, возможно, не принял бы некоторые записи из ответа на запрос (такие записи составляют информационный шум); если бы он имел возможность при этом «видеть» всю базу данных, он, возможно, дополнил бы некоторыми записями ответ на свой запрос (такие записи составляют потери информации). Такого рода потери информации и шумы порождаются нечеткостью лингвистических описаний объектов.

Мы можем формализовать эти понятия следующим образом. Так же, как и в разделе 2.4, будем считать, что множество лингвистических значений задано в виде $s_t \in G(L)$. Рассмотрим случай $t=2$ (рис. 23). Зафиксируем значение $u^* \in U$ и введем следующие обозначения:

- $N(u^*)$ — число объектов, описания которых хранятся в базе данных, имеющих реальные (физические, не лингвистические) значения признака, равные u^* ;
- N^{user} — число пользователей системы.

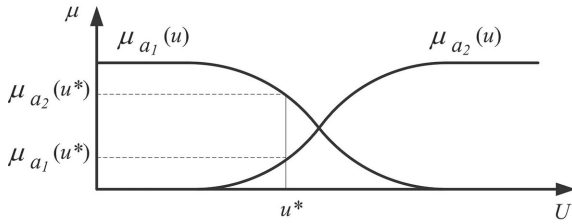


Рис. 23. $s_2 \in G(L)$

Тогда:

- $N_{a_1}(u^*) = \mu_{a_1}(u^*)N(u^*)$ — число объектов, информация о которых хранится в базе данных, имеющих реальные значения признака, равные u^* , и описанных источником информации как a_1 ;
- $N_{a_2}(u^*) = \mu_{a_2}(u^*)N(u^*)$ — число таких объектов, описанных как a_2 ;
- $N_{a_1}^{user}(u^*) = \mu_{a_1}(u^*)N^{user}$ — число пользователей системы, считающих, что u^* есть a_1 ;
- $N_{a_2}^{user}(u^*) = \mu_{a_2}(u^*)N^{user}$ — число пользователей, считающих, что u^* есть a_2 .

Заметим, что для $s_i \in G(L)$ в силу свойства ортогональности имеют место следующие соотношения:

$$N_{a_1}(u^*) + N_{a_2}(u^*) = N(u^*), \quad N_{a_1}^{user}(u^*) + N_{a_2}^{user}(u^*) = N^{user}.$$

В этих обозначениях после запроса «Выдать все объекты, имеющие значение признака, равное a_1 » (обозначим его как $(I(O) = a_1)$ пользователь получает $N_{a_1}(u^*)$ описаний объектов, имеющих реальное значение признака, равное u^* . При этом $N_{a_2}^{user}(u^*)$ пользователи недополучают $N_{a_2}(u^*)$ описаний объектов (несут потери). Это те описания объектов, которые имеют реальное значение признака, равное u^* , но описанные источником как a_2 . По аналогии, оставшиеся $N_{a_2}^{user}(u^*)$ пользователей имеют шум («лишние» с их точки зрения описания объектов в объеме $N_{a_1}(u^*)$ описаний).

Усредненные индивидуальные потери пользователя в точке u^* при анализируемом запросе равны:

$$\pi_{a_1}(u^*) = \frac{1}{N_{a_1}^{user}} N_{a_1}^{user}(u^*) N_{a_2}(u^*) = \mu_{a_1}(u^*) \mu_{a_2}(u^*) N(u^*).$$

По аналогии средние индивидуальные шумы информации

$$\text{в точке } u^* \quad h_{a_1}(u^*) = \frac{1}{N_{a_1}^{user}} N_{a_1}^{user}(u^*) N_{a_2}(u^*) = \mu_{a_2}(u^*) \mu_{a_1}(u^*) N(u^*).$$

Средние индивидуальные потери информации и шумы при анализируемом запросе ($\Pi_{a_1}(U)$ и $H_{a_1}(U)$ соответственно) естественно определить как:

$$\Pi_{a_1}(U) = \frac{1}{|U|} \int_U \pi_{a_1}(u) du, \quad H_{a_1}(U) = \frac{1}{|U|} \int_U h_{a_1}(u) du.$$

$$\text{Таким образом, } \Pi_{a_1}(U) = H_{a_1}(U) = \frac{1}{|U|} \int_U \mu_{a_1}(u) \mu_{a_2}(u) N(u) du.$$

По аналогии для запроса ($I(O) = a_2$) или из соображений симметрии мы получаем, что в этом случае средние индивидуальные потери информации и информационные шумы равны друг другу: ($\Pi_{a_2}(U) = H_{a_2}(U)$) и равны правой части формулы выше.

Под информационными потерями и шумами при поиске информации по признаку, имеющему множество значений $X = \{a_1, a_2\}$ ($\Pi_X(U)$ и $H_X(U)$), естественно понимать $\Pi_X(U) = p_1 \Pi_{a_1}(U) + p_2 \Pi_{a_2}(U)$, $H_X(U) = p_1 H_{a_1}(U) + p_2 H_{a_2}(U)$, где p_i ($i = 1, 2$) — вероятность запроса по значению признака. Так как $p_1 + p_2 = 1$, то

$$\Pi_X(U) = H_X(U) = \frac{1}{|U|} \int_U \mu_{a_1}(u) \mu_{a_2}(u) N(u) du.$$

Рассмотрим общий случай $t > 2$. В этом случае область интегрирования U может быть представлена как $U = U_1 \cup U_{12} \cup U_2 \cup \dots \cup U_{t-1,t} \cup U_t$ где:

- $U_j = [\bar{u}_{j-1,R}, \bar{u}_{j+1,L}]$ ($j = 2, \dots, t$; $\bar{u}_{0,R} = \bar{u}_{1,L}$, $\bar{u}_{t+1,L} = \bar{u}_{t,R}$) — подмножество U , на котором $\mu_{aj}(u) = 1$, и в силу ортогональности $G(L)$;
- $U_{j-1,j} = [\bar{u}_{j,L}, \bar{u}_{j-1,R}]$ — подмножество U , на котором $0 < \mu_{aj-1}(u)$, $\mu_{aj}(u) < 1$; $\mu_{ai}(u) = 0$ для $i \neq j-1, i \neq j$.

Рассмотрим запрос ($I(O) = a_j$ ($1 < j < t$)). В этом случае на качество поиска необходимых объектов оказывают влияние со-

средние значения признака: левое $(j-1)$ и правое $(j+1)$. Для средних потерь информации и информационных шумов, таким образом, справедливо:

$$\Pi_{a_j}(U) = \Pi_{a_j}^{a_{j-1}}(U) + \Pi_{a_j}^{a_{j+1}}(U), H_{a_j}^{a_{j-1}}(U) + H_{a_j}^{a_{j+1}}(U),$$

где:

$$\begin{aligned} \Pi_{a_j}^{a_{j-1}}(U) &= H_{a_j}^{a_{j-1}}(U) = \frac{1}{|U|} \int_U \mu_{a_{j-1}}(u) \mu_{a_j}(u) N(u) du = \\ &= \frac{1}{|U|} \int_{U_{j-1,j}} \mu_{a_{j-1}}(u) \mu_{a_j}(u) N(u) du, \\ \Pi_{a_j}^{a_{j+1}}(U) &= H_{a_j}^{a_{j+1}}(U) = \frac{1}{|U|} \int_U \mu_{a_j}(u) \mu_{a_{j+1}}(u) N(u) du = \\ &= \frac{1}{|U|} \int_{U_{j+1,j}} \mu_{a_j}(u) \mu_{a_{j+1}}(u) N(u) du. \end{aligned}$$

Первое и второе соотношения есть повторение рассуждений вывода формулы для случая $t=2$, последние соотношения следуют из определений множеств $U_{j-1,j}$ и $U_{j,j+1}$.

В этом случае средние индивидуальные потери информации и информационные шумы, возникающие при поиске информации по данному атрибуту, будут равны соответственно $\Pi_X(U) = \sum_{j=1}^t p_j \Pi_{a_j}(U)$, $H_X(U) = \sum_{j=1}^t p_j H_{a_j}(U)$, где p_j – вероятность запроса по j -значению признака X , $\sum_{j=1}^t p_j = 1$.

Заметим, что для крайних значений $j=1$ и $j=t$ справедливо:

$$\begin{aligned} \Pi_{a_1}(U) &= \Pi_{a_1}^{a_2}(U) = H_{a_1}(U) = H_{a_1}^{a_2}(U) = \\ &= \frac{1}{|U|} \int_{U_{12}} \mu_{a_1}(u) \mu_{a_2}(u) N(u) du, \\ \Pi_{a_t}(U) &= \Pi_{a_t}^{a_t}(U) = H_{a_t}(U) = H_{a_t}^{a_t}(U) = \\ &= \frac{1}{|U|} \int_{U_{t-1,t}} \mu_{a_{t-1}}(u) \mu_{a_t}(u) N(u) du. \end{aligned}$$

Подставляя $\Pi_{a_j}(U)$ и $H_{a_j}(U)$ в формулы для $\Pi_X(U)$, $H_X(U)$ и помня сделанное замечание, мы получаем:

$$\begin{aligned}
\Pi_X(U) &= H_X(U) = p_1 \frac{1}{|U|} \int_{U_{12}} \mu_{a_1}(u) \mu_{a_2}(u) N(u) du + \\
&+ \sum_{j=1}^{t-1} p_j \frac{1}{|U|} \left(\int_{U_{j-1}} \mu_{a_{j-1}}(u) \mu_{a_j}(u) N(u) du + \right. \\
&+ \left. \int_{U_{j,j+1}} \mu_{a_j}(u) \mu_{a_{j+1}}(u) N(u) du \right) + \\
&+ p_t \frac{1}{|U|} \int_{U_{t-1,t}} \mu_{a_{t-1}}(u) \mu_{a_t}(u) N(u) du = \\
&= \frac{1}{|U|} \sum_{j=1}^{t-1} (p_j + p_{j+1}) \int_{U_{j,j+1}} \mu_{a_j}(u) \mu_{a_{j+1}}(u) N(u) du.
\end{aligned}$$

Таким образом, мы можем обобщить формулу для случая $t=2$:

$$\Pi_X(U) = H_X(U) = \frac{1}{|U|} \sum_{j=1}^{t-1} (p_j + p_{j+1}) \int_{U_{j,j+1}} \mu_{a_j}(u) \mu_{a_{j+1}}(u) N(u) du, \text{ где}$$

$X = \{a_1, \dots, a_t\}$, $p_j (j=1, 2, \dots, t)$ — вероятность запроса по j -му значению признака.

Для установления искомой связи качества описания объектов и качества поиска информации надо сравнить последнее выражение с формулой для $\xi(s_t)$ (раздел 2.4). Справедлива следующая теорема [7]:

Теорема 5. Пусть $X = \{a_1, a_2\}$, $s_2 \in G(L)$, $N(u) = N = \text{Const}$. Тогда $\Pi_X(U) = H_X(U) = c \xi(s_2)$, где c есть константа, зависящая только от N .

Справедливо также ее обобщение [7].

Теорема 6. Пусть $X = \{a_1, \dots, a_t\}$, $s_t \in G(L)$, $N(u) = N = \text{Const}$ и $p_j = \frac{1}{t} (1 \leq j \leq t)$. Тогда $\Pi_X(U) = H_X(U) = \frac{c}{t} \xi(s_t)$, где c есть константа, зависящая только от N .

Таким образом (теоремы 5 и 6), мы можем утверждать, что объемы потерь информации и информационных шумов, возникающие при поиске информации в нечетких (лингвистических) базах данных, согласуются со степенью неопределенности описания объектов.

Потери и шумы являются робастными. Для $s_t \in G(L)$ (напомним, что $\bar{L}$ — подмножество L , содержащее только линейные

функции принадлежности) мы можем, как и для степени нечеткости, посчитать потери информации и шумы. Для $N(u) = N = \text{Const}$ и $p_j = \frac{1}{t} (1 \leq j \leq t)$ они будут равны: $\Pi_X(U) = H_X(U) = \frac{Nd}{3t|U|}$, где d — мощность области неопределенности. Как следствие, при сформулированных условиях получаем: $\Pi_X(U) = H_X(U) = \frac{2N}{3t} \xi(s_t)$.

Справедлива следующая теорема [7].

Теорема 7. Пусть $X = \{a_1, \dots, a_t\}$, $s_t \in G^8(L)$, $N(u) = N = \text{Const}$ и $p_j = \frac{1}{t} (1 \leq j \leq t)$. Тогда

$$\underline{\Pi}_X(U) = \underline{H}_X(U) = \frac{ND(1-\delta_2)^3}{3t|U|},$$

$$\underline{\Pi}_X(U) = \underline{H}_X(U) = \frac{ND(1-\delta_2)^3}{3t|U|} + \frac{2ND\delta_2}{t|U|}.$$

Вспоминая теорему 4, получаем следующее следствие теоремы 7:

$$\begin{aligned} \underline{\Pi}_X(U) = \underline{H}_X(U) &= \frac{2N}{3t} (1-\delta_2) \underline{\xi}(s_t), \quad \bar{\Pi}_X(U) = \bar{H}_X(U) = \\ &= \frac{2N}{t(1+\delta_2)} \left[\frac{(1-\delta_2)^3}{3} + 2\delta_2 \right] \bar{\xi}(s_t). \end{aligned}$$

Это означает, что и связь качества поиска информации и качества описания объектов является робастной.

Учитывая перечисленные свойства потерь информации и шумов, метод выбора оптимального для поиска множества значений качественного признака можно сформулировать так:

- 1) генерируются все «разумные» множества значений лингвистической переменной;
- 2) каждое из таких множеств представляется в форме полного ортогонального семантического пространства;
- 3) для каждого из них вычисляется мера неопределенности $\xi(s_t)$;

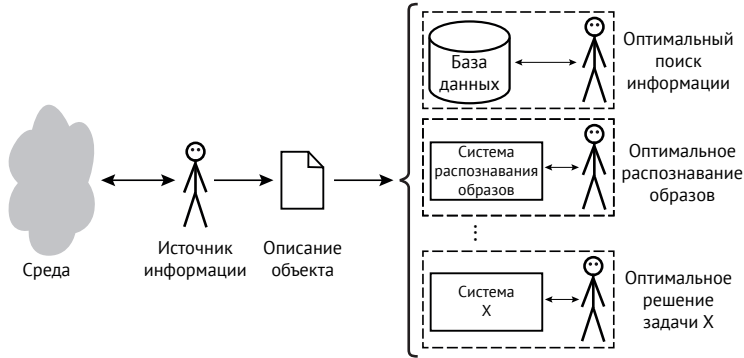


Рис. 24. Обработка описаний человеком объектов среды

- 4) в качестве оптимального множества значений как с точки зрения поиска информации выбирается то множество, отношение меры неопределенности которого к числу значений минимально: $\xi(s_t^*) = \min_{s_t \in G(L)} \frac{\xi(s_t)}{t}$.

Следуя этому методу, можно описать объекты таким образом, чтобы обеспечить *максимально возможное* качество поиска информации, то есть *гарантировать оптимальность* свойств систем гибридного интеллекта с этой точки зрения. Устойчивость потерь информации и шумов (теорема 7) позволяет использовать этот метод в практических задачах.

Аналогичные соотношения справедливы и для задач классификации [9]. Мы не будем так же подробно, как и для поиска информации, описывать соответствующую модель (это требует определения и описания ряда дополнительных понятий), интересующийся читатель сможет самостоятельно изучить работу по приведенной ссылке.

Рассмотрим следующее обобщение (рис. 24).

Обобщение метода выбора оптимального множества значений качественных признаков на общий случай любой задачи может звучать так:

- 1) формализуется качество решения задачи;
- 2) находится (если существует) или численно моделируется связь функционала качества решения задачи и $\xi(s_t)$;

- 3) в качестве оптимального множества значений выбирается такое, которое обеспечивает максимальное значение функционала качества решения задачи.

Приведенные в этой главе результаты позволяют утверждать, что базовые проблемы гибридного интеллекта имеют свое решение. Мы можем строить оптимальные системы такого класса. Более того, решения являются робастными, то есть мы можем использовать их при разработке реальных систем. Сценарии использования систем гибридного интеллекта рассмотрим в следующей главе.

3. Сценарии использования гибридного интеллекта

Развитие микроэлектроники, информационно-телекоммуникационных технологий, а также появление компаний, способных привлечь миллионы пользователей, привели к возникновению нового явления—цифровой модели реального мира (см. раздел 1.2.3). Гибридный интеллект—это способ соединить реальный и цифровой миры.

Есть несколько сценариев использования гибридного интеллекта. Одним из них является персонализация общения человека с цифровым миром. Человечество оказалось способным сделать мир физический очень персонализированным: у нас разная одежда, нет двух одинаковых машин (после того, как их кто-то купил), квартир (после того, как в них поселился хотя бы один человек), смартфонов и много всего, что нас окружает. Если вдуматься, то это было очень нелегко сделать, нужно было многому научиться, но это сделано. Общаясь с миром цифровым, мы понимаем, что в лучшем случае мы элементы какого-то сегмента, в котором находятся еще сотни тысяч людей, воспринимаемых алгоритмами как совершенно одинаковые. Магазины, банки, телекоммуникационные и страховые компании постоянно шлют нам какие-то ненужные рекламы и предложения как элементам своих сегментов, что раздражает. Можно ли сделать цифровой мир таким же удобным и персонализированным, как

и мир физический, для всех и каждого? Эти вопросы обсудим в разделе 3.1.

Второй сценарий связан с совместным (человека и компьютера) способом решения широкого класса задач оценки и мониторинга процессов в бизнесе, экономике, социальной жизни. Может ли компания достигнуть определенной цели (например, производить автомобили дешевле и лучше конкурента)? Что для этого надо сделать в нашей конкретной компании? Сколько это стоит? Что и где нужно улучшить для достижения поставленной цели? Аналогичные вопросы справедливы для разных субъектов экономики и бизнеса — от стартапов до стран, над этим постоянно работают многие тысячи специалистов. Ясно, что нет волшебной формулы для ответа на подобные вопросы, ясно, что нет приборов, которые могут измерить нужные нам параметры и получить ответ. В этом случае системы гибридного интеллекта могут оказаться полезными. Примеры таких систем описаны в разделе 3.2.

3.1. Персонализация

Для эффективной работы гибридного интеллекта алгоритмы искусственного интеллекта (компьютерная компонента) должны понимать, кто сидит «по ту сторону монитора». Простейший пример: кто-то ищет «недорогой ресторан». Кто ищет: студент или олигарх? То, что для первого очень дорого, для второго — просто ошибка округления. У этого сценария есть две ветки: когда активным элементом является человеческая компонента и когда активной является компьютерная компонента. Первый случай будет рассмотрен на примере задачи поиска информации (раздел 3.1.1). Во втором случае мы должны иметь набор алгоритмов, позволяющих получить ответ на сформулированный вопрос, наблюдая за человеком «по ту сторону монитора». Далее мы приведем два примера такого наблюдения: системы компьютерного обучения и персонализация новостной ленты (раздел 3.1.2).

Отметим, что персонализация — это совсем не модельная или учебная задача, потребность в персонализации велика

сейчас и будет только увеличиваться. За последние несколько лет произошло качественное изменение использования информации людьми в решении повседневных задач. Количество накопленной информации в доступном для электронной обработки виде измеряется зеттабайтами (триллион гигабайт). Так, в недавно проведенном исследовании ведущих компаний *Seagate* (www.seagate.com) и *IDC* (<https://www.idc.com>) [11] подсчитано, что в 2016 году объем данных измерялся 16 Збайт, а к 2025 году этот показатель увеличится до 163 Збайт. Там же утверждается, что к 2025 году около 20% всей информации будут играть критически важную роль в повседневной жизни, а примерно 10% этих данных будут «сверхкритичными». Кроме того, прогнозируется, что в 2025 году почти 20% генерируемых данных будут представлять собой информацию, получаемую в режиме реального времени. Количество активных SIM-карт, используемых в телефонах, смартфонах и планшетах, впервые в истории превысило число живущих на Земле людей в 2014 году и продолжает расти [10]. Количество сервисов, используемых в повседневной жизни, также растет: трудно представить себе современного человека, не использующего поиск в интернете, новостные агрегаторы, социальные сети, электронные карты, заказ билетов и многие другие ставшие привычными сервисы.

В сложившейся ситуации способы взаимодействия пользователя с информационными ресурсами также меняются. Одним из признанных трендов таких изменений является персонализация [12–24]. Компании видят существенное влияние персонализации на многие сферы бизнеса от маркетинга и продаж до оптимизации операционной деятельности и принятия стратегических решений. Это справедливо для многих индустрий — от массовой торговли до индустрии моды. В связи с этим разработка и исследование моделей и алгоритмов персонализации видится востребованной задачей.

Впервые идея алгоритма персонализации была предложена в работе [25], в [26] дана ее детализация и рассмотрены варианты оптимизации. Примеры использования таких моделей в компьютерном обучении, оптимизации новостных лент, социальных сетях представлены в работах [27–32].

Базовой для многих моделей персонализации является задача поиска информации. В общем виде задача поиска информации для человеко-компьютерных систем рассмотрена в [7].

3.1.1. Персонализация поиска информации

Первая ветка (когда активным элементом гибридного интеллекта является человеческая компонента) подразумевает реакцию пользователя на результат работы компьютерной компоненты и корректировку ее алгоритмов в соответствии с этой реакцией. Ясно, что реакция не должна быть нудной и обременительной для пользователя, это должен быть простой диалог, похожий на «лайк», к чему многие уже привыкли. Самый простой пример: мы что-то ищем (ресторан, телефон, автомобиль) в соответствующей базе данных; мы хотим «не очень дорогой, комфортабельный...». Ознакомившись с результатом поиска, мы можем сказать: «можно чуть подороже» или «можно менее комфортабельный». Алгоритмы должны понять, как перестроить параметры поиска.

Стандартные инструменты поиска (например, SQL запросы) позволяют нам осуществлять поиск информации только с использованием понятий (слов, символов, цифр), которые присутствуют в базе данных. Это приводит к трудностям в тех ситуациях, когда нужная нам информация не выражается однозначно на языке значений атрибутов.

Одним из важных свойств информации, которую мы ищем в сети, является нечеткость понятий: пользователь, как и любой человек, думает в качественных категориях [6]. В то же время в базе данных информация обычно представлена в виде четких значений. Это несоответствие является одной из главных проблем в «переводе» потребности пользователя в запрос к базе данных. Мы можем проиллюстрировать это на следующем примере.

Пример 3. Выбор автомобиля

Рассмотрим ситуацию покупателя при выборе автомобиля из базы данных (электронного каталога), которая содержит следующую информацию:

- цена;
- модель;
- год выпуска;
- расход топлива.

Для формулирования запроса к такой базе данных пользователь вынужден представить свою мечту на языке конкретных, однозначных чисел (интервалов) и моделей. Если наш пользователь имеет в виду конкретный автомобиль, для которого вышеперечисленные атрибуты имеют определенное значение, и если он решил, что другие машины (похожие или, например, немного более экономичные) не подходят, то стандартные технологии баз данных решают проблему. Но если он хочет купить «не очень дорогой», «престижный», «достаточно экономичный» и «не очень старый» автомобиль, то сформулировать такой запрос к стандартной базе данных невозможно.

Аналогичную ситуацию мы имеем при выборе квартиры или другой недвижимости с использованием базы данных, содержащей описание конкретных квартир в городе. Обработка запросов типа «удобная, не очень дорогая квартира, рядом с парком и не очень далеко от станции метро» невозможна в рамках стандартных SQL технологий.

Это ограничение было давно подмечено, и для его разрешения появились обобщения SQL, позволяющие оперировать с нечеткими понятиями—Fuzzy SQL или FSQL. Читатель без труда найдет множество ссылок, описывающих технические решения такого расширения. Достаточно детальные обзоры этого направления представлены, например, в работе [34]; с конкретными техническими решениями можно ознакомиться в [33]. Поэтому мы не будем их описывать, считая, что такой инструмент существует, достаточно детально описан в доступной литературе и широко используется.

Адаптивный семантический интерфейс (АСИ), позволяющий разрешить указанное противоречие, можно рассматривать как настройку FSQL. Его структура представлена на рис. 25.

Идея АСИ заключается в предоставлении пользователю интерфейса, позволяющего:

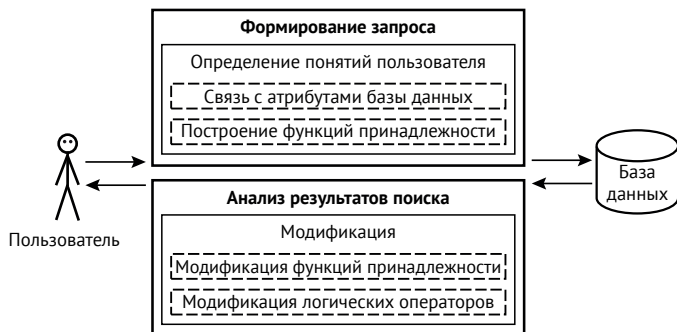


Рис. 25. Структура адаптивного семантического интерфейса

- определять понятия пользователя;
- осуществлять поиск информации по этим понятиям;
- настраивать семантику понятий на основе результатов поиска.

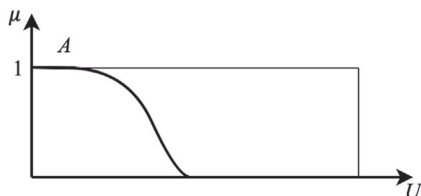
Первые два пункта реализуются в рамках FSQL, последний — это его настройка для конкретного пользователя.

Работа пользователя начинается с лингвистического описания объектов, которые он хочет найти в базе данных (например, цена = «не очень дорогой» в примере 2). Если система не знает этого лингвистического значения, управление передается программе построения функций принадлежности. Для этого надо указать атрибут базы данных, на котором определяется данное значение (в нашем примере — цена), и собственно задать функцию принадлежности.

Мы можем выделить две ситуации: непрерывное универсальное множество ($U \subseteq R^1$) и дискретное ($U \not\subseteq R^1$). В обоих случаях мы имеем методы построения функций принадлежности, которые были широко апробированы при использовании нечетких моделей. Мы можем предоставить следующее продолжение примера 3.

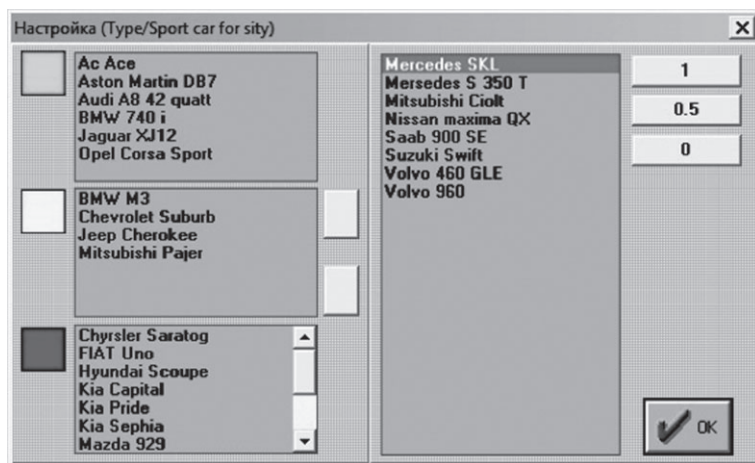
Пример 3 (продолжение):

- а) непрерывное универсальное множество: $U = [\$5000, \$50\,000]$; лингвистическое значение $A =$ «не очень дорогой автомобиль» представлено на рис. 26;

Рис. 26. Пример формализации A = «не очень дорогой автомобиль»

- б) дискретное универсальное множество: U – множество всех моделей автомобилей базы данных; лингвистическое значение B = «спортивный автомобиль для города» (рис. 27).

В правой колонке представлены все автомобили из базы данных; серый ящик — это множество автомобилей, которые безусловно принадлежит определяемому понятию; черный ящик — это автомобили, которые безусловно не принадлежат определяемому понятию; белый ящик — это множество автомобилей, частично принадлежащих определяемому понятию. Процесс построения функции принадлежности начинается с пустых ящиков (все модели есть в правой колонке

Рис. 27. Пример формализации B = «спортивный автомобиль для города»

ке) и заканчивается их заполнением (правая колонка становится пустой). При этом всем автомобилям из серого ящика присваивается значение 1, из черного — 0, из белого вычисляется в виде линейной функции (считается, что объекты упорядочены по степени их «спортивности»). Таким же способом мы можем определить «автомобили для охоты», «автомобили для фермеров», «автомобили для молодых девушек» и др.

Теперь можно приступать к поиску информации с помощью определенных таким образом лингвистических значений. Смысл алгоритма состоит в вычислении для каждой записи базы данных степени ее соответствия сформулированному запросу: от 1 (полное соответствие) до 0 (полное несоответствие). Результат поиска — это упорядочение записей в базе данных по степени соответствия запросу. Схему алгоритма можно представить следующим образом.

Обозначения:

- $i(1 \leq i \leq N)$ — номер атрибута базы данных;
- U_i — домен атрибута;
- $A^i = \{A_1^i, \dots, A_{n_i}^i\}$ — лингвистические значения пользователя, определенные на i -атрибуте ($n_i > 0$);
- $a_{n_k}^i(u_i) = \mu_{A_{n_k}^i}(u_i)$ — функция принадлежности n_k -лингвистического значения i -атрибута ($n_k \leq n_i, u_i \in U_i$);
- $Q = \langle A_{k_1}^1 \circ A_{k_2}^2 \circ \dots \circ A_{k_N}^N \rangle$, где $k_i \leq n_i, A_{k_i}^i \in A^i \cup \emptyset, (1 \leq i \leq N)$; — запрос в форме лингвистических значений.

Алгоритм:

- 1) $r = 0$ (r — номер записи в базе данных; $r \leq R$);
- 2) $r = r + 1$;
- 3) $i = 0$ (i — номер атрибута базы данных, входящего в Q);
- 4) $i = i + 1$;
- 5) если $A^i \in Q$, вычисляем $a_{n_k}^i(u_i)$;
- 6) если $i < N$, возвращаемся на шаг 4;
- 7) вычисляем $\mu_Q(r) = a_{k_1}^1 * a_{k_2}^2 * \dots * a_{k_N}^N$, где $*$ — t -норма, если “о” в Q есть “и”; -конорма, если “о” в Q есть “или”; $1 - a_{k_i}^i$, если “о” в Q есть “не”;
- 8) если $r > R$, возвращаемся на шаг 2.

Результат: $\mu_Q(1), \dots, \mu_Q(R)$ – вектор принадлежностей каждой записи из базы данных запросу пользователя.

В рамках этой общей схемы возникает вопрос: как понять, насколько формализация лингвистических значений соответствует конкретному пользователю? Например, понятие «не очень дорогой автомобиль» может иметь разный смысл для людей с разным достатком (пример 3).

Есть достаточно много различных методов построения функций принадлежности, есть их классификация [7], однако понимание, насколько конкретная функция принадлежности соответствует представлению конкретного пользователя без привлечения самого пользователя, невозможно. При этом достаточно очевидно, что такое взаимодействие с пользователем не должно быть обременительным.

Допустим, мы каким-либо образом задали функции принадлежности интересующих понятий¹. По результатам поиска возможны две ситуации: 1) пользователь удовлетворен результатами и 2) пользователь не удовлетворен результатами поиска. В первом случае задача решена: мы «угадали» функции принадлежности. Дальше мы можем зафиксировать эти формализации семантики пользователя в его профиле и использовать в дальнейшем при его общении с цифровым миром. Во втором случае мы должны настроить семантику понятий. Такая настройка возможна при предположении, что пользователь может выразить свое отношение к результатам поиска (неудовлетворенность ими) в виде выражений «значительно дешевле», «можно чуть дороже» и подобных им.

Обозначим такого типа высказывания в виде пары {<сила модификации>, <направление модификации>}. Тогда возможно изменение положения текущей функции принадлежности, как это показано на рис. 28.

Направление модификации однозначно определяет направление сдвига функции принадлежности. Сила модификации определяет шаг, на который сдвигается функция принадлежности, и здесь возможна неоднозначность. Эта неоднозначность в значительной степени определяется рас-

¹ Мы вернемся к этому допущению позже.

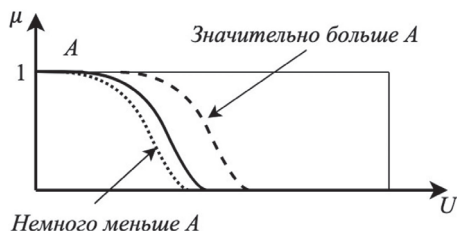


Рис. 28. Модификация функции принадлежности

пределением объектов в базе данных (если оно плотное в окрестности функции принадлежности, то шаг может быть не таким большим, как для разреженного, для одного и того же модификатора). Каких-то общих рекомендаций или теорем, позволяющих выбрать шаг в зависимости от статистики базы данных, в настоящее время не сформулировано. Однако заметим, что если мы будем использовать итерационную процедуру, уменьшая шаг, ассоциированный с силой модификации, на каждой итерации, то мы либо найдем нужный объект, если он существует, либо увидим заикливание с минимальной силой модификации, если такого объекта в нашей базе данных нет (выбор между двумя соседними объектами базы данных).

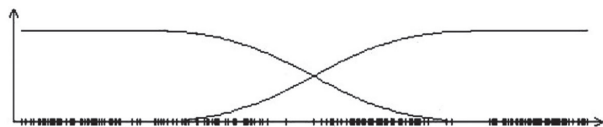
Первая ситуация означает удовлетворенность пользователя результатами, и, соответственно, мы можем зафиксировать эти функции принадлежности в его профиле и использовать в дальнейшем при его общении с цифровым миром. Сходимость (напомним, при условии существования объекта в базе данных) достигается, как и в любом алгоритме, основанном на методе дихотомии и его вариациях (например, алгоритм поиска льва в пустыне).

Вторая ситуация (объекта в нашей базе данных нет, происходит выбор между двумя соседними объектами) означает, что мы нашли некоторое приближение функции принадлежности пользователя, которое может быть зафиксировано в его профиле и, возможно, уточняться в дальнейшем (до появления «идеального объекта» в нашей базе данных).

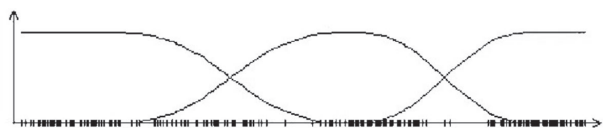
Таким образом, в результате серии естественных интеграций в формате {<сила модификации>, <направление моди-

фикации>} мы получаем формализацию семантики конкретного пользователя, которая позволяет персонализировать его общение с цифровым миром. Получившиеся в результате такой процедуры функции принадлежности отличаются для разных пользователей и представляют собой окно в цифровой мир, позволяющее сделать его адаптированным для конкретного пользователя.

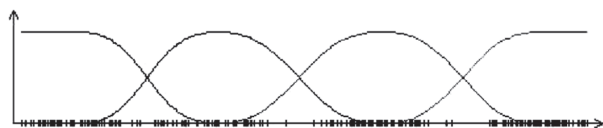
Вернемся к допущению выше (мы каким-либо образом задали функции принадлежности интересующих понятий). Сейчас достаточно очевидно, что описанная итерационная процедура позволяет нам не заботиться о точности угадывания функции принадлежности для конкретного пользователя. Время ее сходимости, правда, может различаться в зависимости от начального приближения—здесь полная аналогия с итерационными методами вычислительной математики. Наверное, соображения о назначении такого приближения исходя из близости профилей пользователей (уже имеющих настроенные функции принадлежности и нового) могут быть достаточно разумны, но эти вопросы относятся к конкретной реализации для конкретной задачи (например, социальной сети), поэтому здесь мы не будем их рассматривать. Однако заметим, что существуют надежные и проверенные в многочисленных приложениях методы нечеткой кластеризации, позволяющие строить функции принадлежности понятий исходя из статистических свойств базы данных. Мы задаем лишь число необходимых нам понятий (классов), например три—«дешевый», «средней цены» и «дорогой», а алгоритм разбивает нам область определения (значения соответствующего атрибута базы данных) на три класса с примерно одинаковым количеством объектов в каждом классе. В качестве примера такого алгоритма можно привести алгоритм нечеткой кластеризации (FCM-алгоритм), который был предложен Данном в [35] и развит Бездеком в [36]. FCM-алгоритм является обобщением известного алгоритма кластеризации ISODATA на случай нечетких множеств. Этот (и подобные) алгоритм включен в состав многих пакетов по кластеризации, читатель без труда найдет его open source код в интернете, поэтому мы не будем его описывать, приведем лишь результаты его работы (рис. 29).



(а)



(б)



(в)

Рис. 29. Функции принадлежности, построенные процедурой нечеткой кластеризации для двух (а), трех (б) и четырех (в) классов

Применяя данный алгоритм, мы автоматически получаем выполнение требований полного ортогонального семантического пространства $G(L)$ (раздел 2.4). Это позволяет оптимизировать АСИ как с точки зрения описания объектов (раздел 2.4), так и с точки зрения поиска информации (раздел 2.5).

При таком построении функций принадлежности мы можем осуществлять модификацию понятий не на основе сдвигов функции (рис. 28), а на основе изменения самого универсального множества. Например, получив реакцию «значительно дешевле», мы просто убираем из универсального множества «дорогие» объекты и строим функции принадлежности заново. Такая процедура была предложена Н. М. Огородниковым [37]. В рамках ее можно строго доказать, что если в базе есть интересующий нас объект (наша «мечта»), то, имея лишь направление модификации (дороже — дешевле, более комфортабельный — менее комфортабельный) и ее

силу (чуть более, значительно менее), мы его найдем за разумное количество шагов или в противном случае поймем, что его нет (при возникновении заикливания).

Описанная технология адаптивного семантического интерфейса позволяет персонализировать социальные сети, выбор товаров в интернет-магазинах, планирование путешествий и другие активности в цифровом мире. Набор таких понятий и их функций принадлежностей составляет уникальный портрет пользователя, его персональный ключ к цифровым ресурсам. Человек становится не субъектом, которого постоянно сегментируют владельцы цифровых ресурсов, а активным объектом при взаимодействии с ними. Это позволяет совершенно по-новому организовать такое взаимодействие, включая маркетинг, продажи и многое другое. Цифровой мир становится по-настоящему персонализированным для каждого. Возможности оптимизации АСИ (разделы 2.4 и 2.5) позволяют сделать взаимодействие с цифровым миром максимально комфортным.

3.1.2. Персонализация обучения в компьютерных обучающих системах

Здесь рассматривается другой вариант взаимодействия человека и компьютера — когда активной является компьютерная компонента. Наблюдая за человеком «по ту сторону монитора», компьютерная компонента пытается восстановить его портрет для оптимизации совместного решения задач. Мы рассмотрим механизмы такого взаимодействия на примере обучения в рамках компьютерных обучающих систем. Эта задача выбрана нами не только потому, что она всем понятна и не требует дополнительных разъяснений, но и потому, что она является чрезвычайно актуальной.

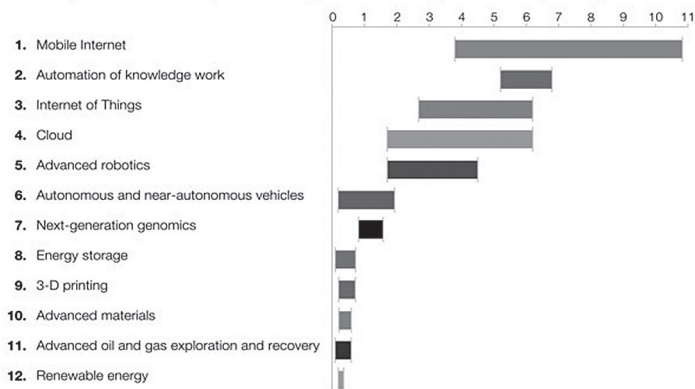
Понятие компьютерных обучающих систем возникло практически одновременно с понятием искусственного интеллекта (про это, в частности, пишет А. Тьюринг в знаменитой работе [1]) и претерпело значительную эволюцию одновременно с развитием компьютерных технологий. Мы не будем проводить такой исторический экскурс; заинтересованному читателю

лю можно порекомендовать работу [38], где эти вопросы детально обсуждаются. Однако нельзя не отметить глобальных изменений, произошедших в последнее время и происходящих сегодня. Наиболее ярко они представлены в работах McKinsey Global Institute (MGI), упоминавшегося выше [39], [40]. Авторы одного из наиболее авторитетных аналитических институтов, изучая влияние более 100 технологий на развитие экономики и общества, пришли к следующим заключениям:

- 1) наиболее важными являются технологии, представленные на рис. 30 [39]. Там же оценки экономического вклада технологии в 2025 году; с методикой расчета эффективности можно ознакомиться в [40];
- 2) технологии автоматизации интеллектуальной деятельности (годовой вклад в экономику 2025 года от 5,2 до 6,7 триллиона долларов в текущих ценах) в первую очередь обеспечат новые способы обучения на базе интеллектуальных систем и big data (рис. 31).

A gallery of disruptive technologies

Estimated potential economic impact of technologies across sized applications in 2025, \$ trillion, annual



SOURCE: McKinsey Global Institute

Notes on sizing: These economic impact estimates are not comprehensive and include potential direct impact of sized applications only. They do not represent GDP or market size (revenue), but rather economic potential, including consumer surplus. The relative sizes of technology categories shown do not constitute a "ranking," since our sizing is not comprehensive. We do not quantify the split or transfer of surplus among or across companies or consumers, since this would depend on emerging competitive dynamics and business models. Moreover, the estimates are not directly additive, since some applications and/or value drivers are overlapping across technologies. Finally, they are not fully risk- or probability-adjusted.

Рис. 30. Список прорывных технологий по версии MGI

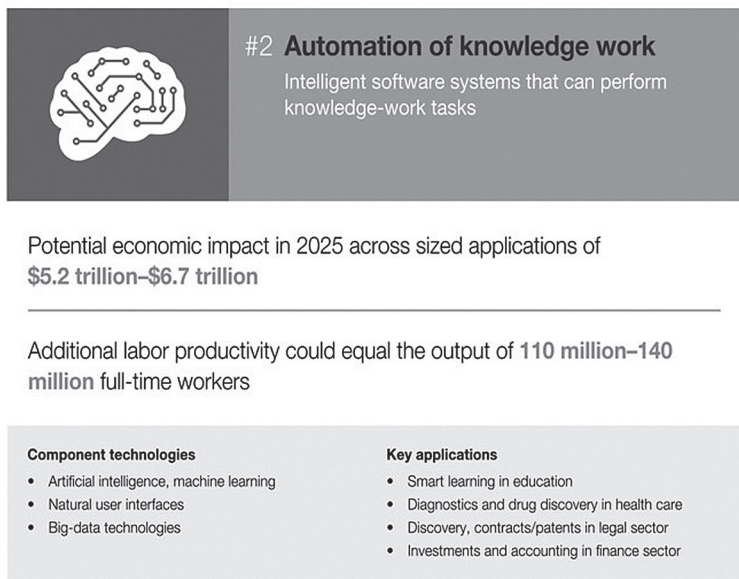


Рис. 31. Основные параметры технологии автоматизации интеллектуальной деятельности по версии MGI

Основным аргументом в пользу такого влияния обучающих систем на экономику авторы видят всеми нами осознаваемое противоречие: исчезновение массовых специальностей (машинистки, операторы ЭВМ и пр.) и появление новых (потребность — сотни тысяч / миллионы рабочих мест), с одной стороны, и практическое отсутствие развития технологий обучения — с другой. Имеется в виду не замена мела и деревянной доски на компьютер и интерактивную доску, а принципиальные новшества, соразмерные вызовам современной экономики.

Первые компьютерные обучающие системы появились в 90-х годах прошлого века, представляли собой учебные материалы (тексты, картинки, клипы и пр.) в электронной форме с гиперссылками, распространяемые на наиболее удобном для того времени носителе — CD-ROM-диске. Со временем возникала специализация по контенту (математика, физика, иностранные языки, музыка и т. п.), добавлялись

мультимедийные эффекты, возникали клиент-серверные решения (с развитием доступа в интернет), но мышление в стиле «контент с гиперссылками» было преобладающим. Потребности участников образовательного процесса — студентов, образовательных учреждений, производителей образовательного контента — удовлетворялись разработкой специальных информационных систем или адаптацией известных решений для нужд участников, что привело к формированию определенного направления в современных информационных технологиях. Сложившийся к настоящему времени ландшафт направления компьютерных обучающих систем (Ed Tech в англоязычной терминологии) наиболее полно представлен в [41]. Инструменты обучения практически не менялись со времени начала массового использования книг, то есть на протяжении нескольких веков. Массовая доступность компьютеров и развитие ИТ-технологий (в основном — сетевых) создают новую инфраструктуру, сравнимую по своему влиянию на повседневную жизнь с началом массового использования книг и печатной продукции. Поэтому естественным является вопрос: каковы необходимые условия для использования особенностей новой технологической среды в обучении?

Если рассматривать процесс обучения как управление в рамках стандартной кибернетической системы [42] (рис. 32), можно сказать, что для организации такого процесса как минимум необходимо:

- разнообразие контента (прямая связь);
- возможность измерения действий ученика (обратная связь).

Кроме перечисленного надо также иметь критерий качества управления (задается извне; может быть несколько критериев) и собственно алгоритм управления (тоже задается извне). Однако без наличия прямой и обратной связи — собственного контента и собственного измерительного механизма — никакого отличия от книжной технологии обучения не возникает. В рамках стандартного процесса обучения такие связи также присутствуют, но реализуются без привлечения информаци-

онных технологий. Использование современных ИТ дает дополнительные возможности и имеет потенциал качественного изменения процесса обучения.

Для эффективности компьютерного обучения важно так организовать подачу контента, чтобы ученику было интересно:

не было бы очень простых и очень сложных заданий; если ученик быстро устает, то вначале давать более сложные задания, если медленно разгоняется — наоборот, и т. п. Ясно, что этого нельзя сделать для «усредненного ученика» или элемента какого-либо «сегмента» (такой сегмент может содержать тысячи учеников, и все они для алгоритма являются одинаковыми как шарики). Однако развитая система трекинга (время выполнения задания, количество ошибок и т. п.) позволяет восстановить портрет конкретного ученика. Такая персонализация обучения имеет огромный спрос¹, ее потенциал — обучение на уровне репетитора [27].

Архитектура системы, обеспечивающей персонализацию обучения, представлена на рис. 33. Анализ параметров выполнения заданий (скорость, количество правильных/неправильных ответов — первичные измерения) для нескольких десятков тысяч учеников позволяет их разбить на классы «сильный»... «слабый», причем сделать это можно оптимальным (с минимальной степенью нечеткости) образом [27]. То же самое можно сделать и для заданий (классы «простой»... «сложный»). Это позволяет по первичным измерениям понять скорость выполнения заданий («быстро»... «медленно»), выносливость (например, делает много сложных заданий за среднее время с малым числом ошибок) и другие параметры ученика. Сравнение с его историей позволяет понять, штатно или нет идет процесс обучения (например, устал он

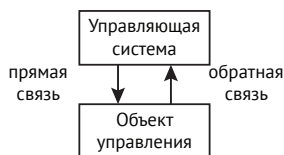


Рис. 32. Схема простейшей кибернетической системы

¹ An education startup you've never heard of is now in 70% of US school districts: <https://www.businessinsider.com.au/kiddom-startup-using-personalized-learning-in-most-school-districts-2018-1>.

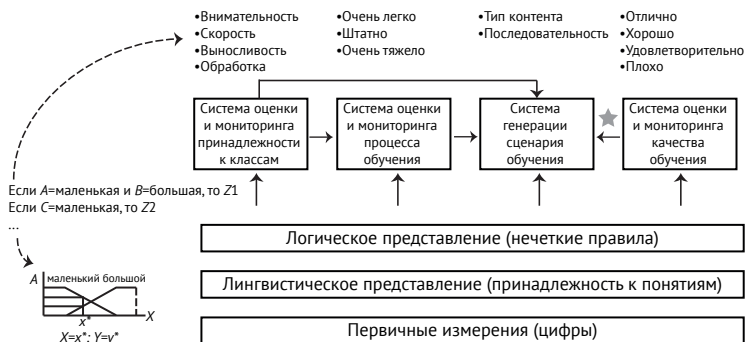


Рис. 33. Архитектура системы персонализации обучения

или нет). На основе этого мы можем давать ему задания, максимально соответствующие его профилю и состоянию (например, видя, что ученик быстро устает, давать сначала сложные задания; если же он, наоборот, медленно «разгоняется» — давать сначала более простые задания). За это отвечает блок генерации сценария обучения.

Такая платформа обеспечивает детальный контроль каждого ученика, персонализацию и оптимизацию процесса обучения с учетом всех параметров, которые доступны для системы трекинга.

Развитие подобных систем позволит преодолеть одно из важных противоречий современной экономики: необходимость переобучения миллионов людей (что связано с развитием технологий) и отсутствие адекватных механизмов для этого (традиционное обучение практически ничем не отличается от времени появления массовых и дешевых учебников).

3.1.3. Цифровые привычки и персонализация новостных лент

В описанной выше задаче персонализации обучения компьютер наблюдает за ходом ее выполнения учеником и в зависимости от значений доступных параметров корректирует процесс, подстраиваясь под индивидуальные особенности ученика. Подобные системы могут повысить эффективность любой систематической деятельности, если определена цель (мы зна-

ем, куда хотим попасть); есть система параметров, характеризующих процесс, и мы умеем их измерять (система трекинга — как телеметрия у ракеты); есть возможность управления (для обучающих систем — подача материала, для самолета — управление рулями высоты); есть модель управления (для человеко-машинных систем — «если — то» правила).

Однако мы можем также наблюдать за особенностями конкретного человека в потреблении информации (то есть без всякой связанной с этим задачи) и на этой основе формировать удобное для него информационное пространство. Прибором, позволяющим проводить такие наблюдения, является, например, смартфон (а также планшет, ноутбук — любое персональное устройство, через которое мы коммуницируем с цифровым миром, или их комбинация). Смартфон является тем физическим носителем, который соединяет физический и цифровой миры для конкретного человека, его персональным ключом в цифровой мир. Есть возможность фиксировать почти все действия: когда, где, куда зашел, сколько там был, что делал и т. п. Значительная часть таких действий фиксируется в логах операционных систем, логах операторов и других «дверей», которые мы открывали с помощью такого персонального ключа. Имея набор таких логов, можно понять, что было до определенного действия, что стало после, можно попытаться найти закономерности и здесь. Рассмотрим следующий пример.

Многие из нас пользуются новостными лентами, но у каждого есть свои привычки. Я, например, пользуюсь общественным транспортом, поэтому утром мне удобно смотреть короткие текстовые новости; вечером, дома, я часто смотрю аналитику (длинные тексты и картинками) и клипы. Выяснив это, можно организовывать новостной поток соответствующим образом [30]. Для этого поток надо охарактеризовать в аналогичных терминах, чтобы была возможность управления им. Обычно известна длина файла, наличие в нем картинок, клипов и пр. элементов. Этого достаточно, чтобы потребление информации было комфортным. Архитектура такой системы персонализации представлена на рис. 34. Поясним ее основные элементы.

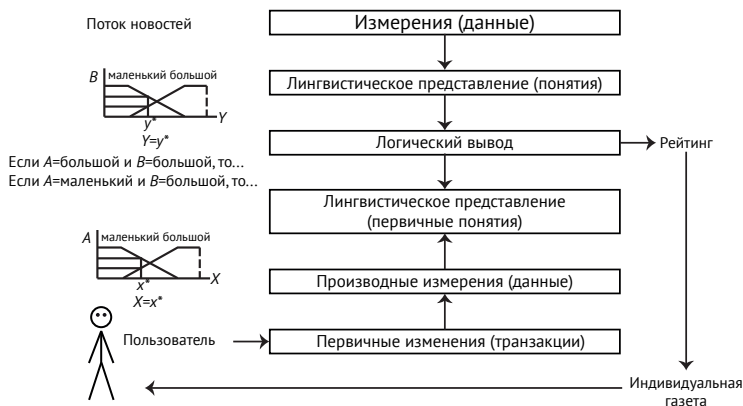


Рис. 34. Архитектура системы персонализации новостных лент

Набор данных (dataset) включал в себя запись всех действий со смартфоном (логов) 800 пользователей за четыре недели, имеющих формат <идентификатор, транзакция>, где транзакция имеет формат <время начала транзакции, приложение, время окончания транзакции>, где приложение — имя одного из приложений, установленных на смартфоне пользователя, с которым происходила работа в указанное время. Это — *первичные измерения*.

Из такого dataset мы можем извлекать различные данные, например среднее количество звонков, среднее количество использования телефона в сутки / за неделю и т. п. Это — *производные измерения* или данные.

Заметим, что выделение производных измерений требует определенной аккуратности. Так, например, разбиение времени суток с точностью до 1 секунды (то, что логируется в смартфоне) оказалось не самым удачным: человек если и делает что-то регулярное, то не с точностью до секунды. Проведенные нами эксперименты показали, что наилучшие результаты достигаются для интервала времени в 15 минут, что вполне разумно. Именно такая агрегация времени использовалась в качестве универсального множества (95 интервалов) для работы со временем суток.

Используя такие данные как универсальные множества, мы можем строить лингвистические значения, описывающие интересующие нас признаки (звонит утром часто/редко, читает новости вечером часто/редко, смотрит клипы утром/вечером и пр.).

Эксперименты, проведенные на реальных данных, показали хорошее качество кластеризации. Использование алгоритмов нечеткой кластеризации гарантирует выполнение аксиом полного ортогонального семитического пространства (раздел 2.4), а следовательно, и возможность оптимизации процесса персонализации новостной ленты (максимальный комфорт, минимальные потери информации и шумы).

Такая программа-«наблюдатель» создает индивидуальный портрет пользователя как потребителя информации и, имея его, организует максимально адаптированное под конкретного пользователя взаимодействие с цифровым миром (не только новостную ленту, но и любое потребление информации).

3.2. ОЦЕНКА И МОНИТОРИНГ ПРОЦЕССОВ

Здесь мы рассмотрим второй сценарий использования гибридного интеллекта—совместное (человека и компьютера) решение широкого класса задач оценки и мониторинга процессов в бизнесе, экономике, социальной жизни.

3.2.1. Понятие систем оценки и мониторинга сложных процессов

Многие процессы в бизнесе, экономике, политике и других областях, называемых слабо (или плохо) формализуемыми, невозможно представить в виде набора дифференциальных уравнений, автоматов и других математических средств анализа динамических систем, однако специалисты как-то работают с такими процессами. В общем виде такая работа заключается в оценке текущего состояния процесса на основе всей доступной информации, построении прогнозов его развития и выработке рекомендаций по управлению им исходя

из целей, стоящих перед специалистом. Можно привести следующие примеры процессов из различных областей:

- поведение клиента (маркетинг);
- продвижение кандидата на выборах (политология);
- диагностика (медицина);
- оценка кредитного риска (финансовый анализ);
- оценка страховых рисков (страховое дело);
- развитие технологий/научных исследований (наукометрия, технологии).

Свойство перечисленных проблем — слабая формализация (наличие количественных и качественных признаков, отсутствие математических моделей), структурная организация (процесс имеет некоторую структуру) и наличие человека как активного элемента системы.

В качестве примера процессов, не являющихся таковыми, можно привести взаимодействие двух тел или распространение колебаний в однородной среде. Имеются математические модели таких процессов, информация измерима и доступна, результат можно вычислить для любого момента времени.

Будем называть задачу оценки текущего состояния процесса и построения прогнозов его развития *задачей оценки и мониторинга*, а человеко-компьютерные системы, обеспечивающие информационную поддержку подобного рода информационных задач, *системами оценки и мониторинга*.

Основные элементы систем оценки и мониторинга — информационное пространство и аналитик. Существенными являются следующие их свойства.

Информационное пространство представляет собой совокупность различных информационных элементов:

- разнородность носителей информации, то есть фиксация информации в виде статей, газетных заметок, компьютерном виде, аудио- и видеоинформация и т. п.;
- фрагментарность: информация чаще всего относится к какому-либо фрагменту процесса, причем разные фрагменты могут быть по-разному «покрыты» информацией;

- **разноуровневость:** информация может относиться ко всему процессу в целом, к некоторой его части, к конкретному элементу;
- **различная степень надежности:** информация может содержать конкретные данные различной степени надежности, косвенные данные, результаты выводов на основе надежной информации или косвенные выводы;
- **возможная противоречивость:** информация из различных источников может совпадать, слегка различаться или вообще противоречить друг другу;
- **изменяемость во времени:** процесс развивается во времени, поэтому и информация в разные моменты времени об одном и том же его элементе может и должна различаться;
- **возможная тенденциозность:** информация отражает определенные интересы источника информации, поэтому может носить тенденциозный характер. В частном случае она может являться намеренной дезинформацией (например, для политических проблем или для проблем, связанных с конкуренцией).

Аналитики являются активным элементом систем оценки и мониторинга и, наблюдая и изучая элементы информационного пространства, делают выводы о состоянии проблемы и перспективах ее развития с учетом перечисленных выше свойств информационного пространства. Обычно аналитики образуют некоторую структуру (министерство, агентство, консультационную службу и т. п.). В этом случае каждый аналитик «нижнего уровня» имеет дело с некоторой частью процесса и работает с элементами информационного пространства, аналитики «более высокого уровня» имеют дело с более крупными фрагментами процесса или процессом в целом и работают уже с выводами предыдущих аналитиков. При этом они могут ознакомиться с выводами более низкого уровня вплоть до элементов информационного пространства.

Взаимодействие основных элементов систем оценки и мониторинга представлено на рис. 35.

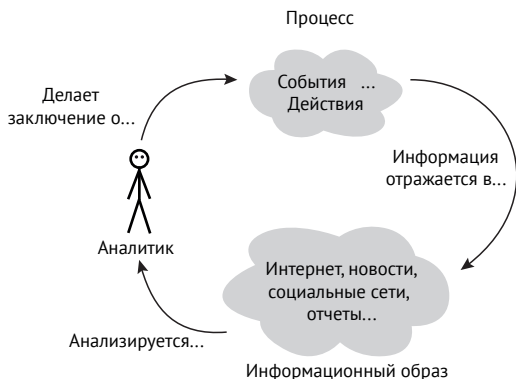


Рис. 35. Взаимодействие основных элементов систем оценки и мониторинга

Учитывая особенности информации и конкретных способов ее обработки, основные черты технологии оценки и мониторинга можно изложить следующим образом.

Эта технология базируется на использовании ряда приемов, позволяющих обрабатывать подобного рода информацию:

- для реализации возможности обработки информации из разнородных источников в базе данных системы хранятся как сами документы, так и ссылки на них с оценкой содержащейся в них информации, данной аналитиком;
- для возможности обработки фрагментарной информации используется структурная модель процесса в виде графа;
- обработка разноуровневой информации достигается за счет предоставления пользователю возможности отнестись к оценке конкретного информационного материала к разным вершинам модели;
- обработка информации различной степени надежности и обладающей возможной противоречивостью или тенденциозностью достигается за счет использования лингвистических оценок данной информации;
- изменяемость во времени учитывается фиксацией даты поступления информации при оценке конкретного ма-

териала, то есть время является одним из элементов описания объектов системы.

Таким образом, системы, построенные на базе этой технологии, позволяют иметь развивающуюся во времени модель процесса на основе оценок аналитиков, подкрепленную ссылками на все информационные материалы, выбранные ими, с общими и частными оценками состояния процесса или его аспектов. Использование времени как параметра системы позволяет проводить ретроспективные анализы и строить прогнозы развития процесса (отвечать на вопросы типа «Что будет, если...?»). В последнем случае возникает возможность выделения «критических путей», то есть таких элементов модели, малое изменение которых может вызвать значительные изменения в состоянии всего процесса. Знание таких элементов имеет большое практическое значение и позволяет выявить слабые места в процессе на текущий момент времени, разработать мероприятия по блокированию нежелательных ситуаций или провоцированию желательных, то есть в некоторой степени управлять развитием процесса в интересах организации, его отслеживающей.

Концептуальная схема систем оценки и мониторинга представлена на рис. 36.

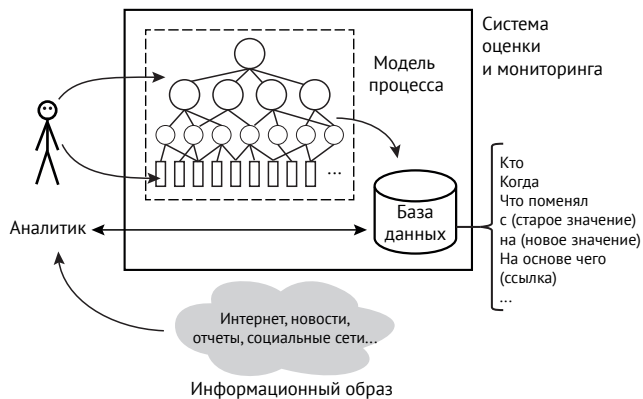


Рис. 36. Концептуальная схема систем информационного мониторинга

Для эффективного практического применения предложенных технологических решений необходима проработка ряда теоретических проблем. Проблема описания человеком объектов среды (нижняя стрелка), как и проблема поиска аналогов в базе данных (горизонтальная стрелка) нами решены — см. разделы 2.4 и 2.5 соответственно. Эти проблемы имеют свои оптимальные решения, которые являются робастными, то есть могут уверенно использоваться в практических задачах. Третья проблема разработки таких систем — проблема выбора адекватных операторов агрегирования информации в нечетких иерархических системах. Остановимся на ней подробнее.

3.2.2. Проблема агрегирования информации в системах оценки и мониторинга

Эта проблема возникает в силу того, что системы оценки и мониторинга ориентированы на обработку разноуровневой фрагментарной информации. Это означает, что при вводе информации в систему пользователь может осуществлять привязку информационных материалов к узлам различных уровней иерархии в модели проблемы (и соответственно подтверждать / изменять их оценки). Это коренным образом отличает системы оценки и мониторинга от аналогичных иерархических систем (например, поддержки принятия решений), в которых разрешается менять только оценки нижнего уровня иерархии. Указанное допущение позволяет использовать системы оценки и мониторинга при решении значительно более широкого класса практических задач, однако платой за это является необходимость разработки соответствующей теории и создание инструментария выбора адекватных операторов агрегирования информации. Поясним суть проблемы на следующем примере.

Пример 4. Пусть текущее состояние модели имеет вид, изображенный на рис. 37.

Для простоты будем считать, что оценки пользователя представляются в виде числа из отрезка $[0,1]$. Пусть далее в качестве операторов агрегирования информации используется операция взятия максимума. Как не трудно видеть,

в этом случае модель является согласованной.

Рассмотрим две ситуации.

Ситуация 1. Пользователь изменил значение в узле a_{211} с текущего на 0.3. В этом случае согласованность остается.

Ситуация 2. Пользователь изменил значение в узле a_{22} с текущего на 0.5. В этом случае введенная оценка входит

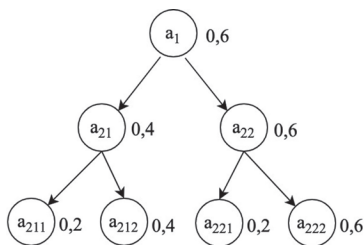


Рис. 37. Упрощенный пример состояния модели

в противоречие как с оценками нижнего уровня, так и с оценкой верхнего уровня. Как быть в этом случае? Если оценки узлов a_1 и a_{212} являются не результатом вычислений в системе, а введены пользователем на основе предыдущей информации, мы должны для сохранения согласованности модели либо заставить пользователя изменить их, либо выбрать другие операторы агрегирования информации. Первый путь, несмотря на насильственный характер, иногда оказывается применимым при условии развитой подсистемы работы с архивом (например, в ситуации, когда оценки, которые необходимо менять, являются достаточно старыми и/или поставлены пользователем на основе недостаточно надежной информации). Однако если пользователь одинаково уверен в правильности всех оценок, необходимо менять оператор агрегирования информации.

Проблема 3. Можно ли предложить процедуру выбора операторов агрегирования информации в нечетких иерархических динамических системах, минимизирующих противоречивость модели процесса в системах оценки и мониторинга?

Пусть для простоты моделью процесса является дерево D с вершинами $d_j (j = 0, \dots, N_D)$, каждой из которых поставлено в соответствие некоторое множество лингвистических значений $X_j = \{a_j^1, a_j^2, \dots, a_j^{N_j}\}$, характеризующих состояние вершины. Каждой неконцевой вершине также приписан некоторый оператор агрегирования информации, позволяющий на основе оценок состояния подчиненных вершин вычислять ее со-

стояние (то есть выбирать один из элементов соответствующего множества значений). Часто выбор такого оператора определяется свойствами модели. Например, если мы имеем оценки наличия на складе колес, корпусов, двигателей и т. п. и хотим оценить возможность сборки автомобиля, в качестве естественного оператора агрегирования мы должны использовать минимум (при отсутствии хотя бы одного компонента мы не сможем собрать автомобиль). Однако часто этот выбор не является настолько очевидным (например, для задач из области политологии, социологии или медицины). В этом случае необходима разработка методов выбора адекватных операторов агрегирования на основе доступной информации от экспертов и анализа функционирования модели.

В общем виде схема выбора оператора агрегирования информации выглядит следующим образом. Рассмотрим для простоты некоторое дерево D и некоторую неконцевую его вершину d_{j_0} с подчиненными ей (в смысле рассматриваемого дерева D) вершинами $d_{j_1}, d_{j_2}, \dots, d_{j_{N_0}}$. Тогда оператор агрегирования информации (ОАИ) O_{j_0} есть функция, определенная на множестве всех возможных значений подчиненных вершин и принимающая значения в множестве X_{j_0} : $O_{j_0} = X_{j_1} \times X_{j_2} \times \dots \times X_{j_{N_0}} \rightarrow X_{j_0}$. Обозначим множество операторов агрегирования информации для вершины d_{j_0} через $M[O_{j_0}]$. Ясно, что для конкретного элемента модели проблемы d_{j_0} число возможных ОАИ является большим: $|M[O_{j_0}]| = |X_{j_0}|^{|X_{j_1} \times X_{j_2} \times \dots \times X_{j_{N_0}}|}$. Нашей задачей является выбор конкретного оператора $o_j \in M[O_j]$ для всех неконцевых вершин d_j дерева-модели D . Этот выбор базируется на некоторой информации I_j об «идеальном» ОАИ $\hat{o}_j \in M[O_j]$. Эта информация представляет собой два множества: $I_j = \Gamma_j \cup \Pi_j$, где Γ_j — множество высказываний экспертов о «правильном поведении» $\hat{o}_j$; Π_j — множество результатов работы выбранного ОАИ. Примерами элементов Γ_j могут служить высказывания: «При $d_{j_1} = a_{j_1}^1$ и $d_{j_2} = a_{j_2}^1$ и ... и $d_{j_{N_0}} = a_{j_{N_0}}^1$ значение $d_{j_0} = a_{j_0}^1$ », «При сильном возрастании d_{j_1} значение d_{j_0} убывает», «Значение d_{j_0} монотонно изменяется по всем аргументам» и т. п. Множество Π_j представляет собой таблицу вида:

$d_{j_1}, d_{j_2}, \dots, d_{j_{N_0}}$	d_{j_0}
$a_{j_1}^1, a_{j_2}^1, \dots, a_{j_{N_0}}^1$	$\tilde{o}(a_{j_1}^1, a_{j_2}^1, \dots, a_{j_{N_0}}^1)$
$a_{j_1}^1, a_{j_2}^1, \dots, a_{j_{N_0}}^2$	$\tilde{o}(a_{j_1}^1, a_{j_2}^1, \dots, a_{j_{N_0}}^2)$
...	...
$a_{j_1}^1, a_{j_2}^1, \dots, a_{j_{N_0}}^{N_{N_0}}$	$\tilde{o}(a_{j_1}^1, a_{j_2}^1, \dots, a_{j_{N_0}}^{N_{N_0}})$
...	...
$a_{j_1}^{N_{j_1}}, a_{j_2}^{N_{j_2}}, \dots, a_{j_{N_0}}^{N_{N_0}}$	$\tilde{o}(a_{j_1}^{N_{j_1}}, a_{j_2}^{N_{j_2}}, \dots, a_{j_{N_0}}^{N_{N_0}})$

В левом столбце таблицы попарно расположены различные значения $d_{j_1}, d_{j_2}, \dots, d_{j_{N_0}}$, в правом — значения d_{j_0} , полученные на основе информации $I_{j_i}^1$ (то есть для этих строк $\tilde{o}_{j_i} = \hat{o}_{j_i}$), или значения d_{j_0} , полученные в ходе работы с системой, или пустые значения.

В начале работы с системой таблица содержит только значения d_{j_0} первого типа (полученные на основе информации $I_{j_i}^1$). По мере получения и оценки пользователем информации таблица заполняется на основе вычислений для выбранного оператора $\hat{o}_{j_i}$ до момента, когда пользователь не согласен с «теоретическим» значением d_{j_0} . Если такое противоречие не возникает, значит, оператор $\hat{o}_{j_i}$ выбран удачно и является адекватным ОАИ для данного узла дерева-модели. Если противоречие возникает, необходимо повторить процедуру выбора адекватного ОАИ, но на основе дополненной и, быть может, уточненной с экспертом информации $I_{j_i}^1$ и $I_{j_i}^2$. Этот процесс повторяется до тех пор, пока вся таблица не будет заполнена. Полученная заполненная таблица и есть адекватный ОАИ для данной вершины d_{j_0} .

Описанная укрупненная схемы выбора ОАИ допускает несколько реализаций в зависимости от природы множества X_j .

Мы можем интерпретировать данное множество как набор дискретных или как набор нечетких значений элемента d_j . Эта интерпретация зависит от свойств задачи. Так, если при оценке информации набора значений $a_j^1, a_j^2, \dots, a_j^{N_j}$ достаточно (ситуации, когда оценка находится между соседними значениями, либо отсутствуют, либо их мало), то мы можем говорить о дискретной модели выбора ОАИ. Если же при оценке поступающей в систему информации возникают ситуации, когда оценка находится между соседними значениями a_j^i и a_j^{i+1} , причем пользователь может говорить, что она более близка, например, к a_j^i , чем к a_j^{i+1} , то мы должны использовать нечеткую модель выбора ОАИ. Заметим, что дискретная модель является частным случаем нечеткой и может быть получена при замене соответствующих нечетких множеств на множества уровня 0,5 для полных ортогональных семантических пространств. Мы ее выделяем, однако потому, что в этом случае возможна разработка специальных алгоритмов выбора ОАИ.

В зависимости от доступной информации I_j в рамках дискретной модели можно выделить два подхода к выбору ОАИ: геометрический и логический. Первый применим тогда, когда возможно только определить значение $\hat{o}_j$ на некоторых наборах значений $d_{j_1}, d_{j_2}, \dots, d_{j_{N_0}}$. Второй — когда кроме этого возможна формулировка некоторых условий на «поведение» $\hat{o}_j$.

Геометрический подход основан на представлении оператора агрегирования информации в узле, имеющем n потомков, как некоторой поверхности в $n+1$ -мерном пространстве. Предполагается, что мы можем определить значение в узле-родителе на некоторых наборах значений узлов-потомков. Эти наборы значений и значение в узле-родителе интерпретируются как точки в $n+1$ -пространстве. Далее исходя из этих точек и некоторых предположений строится поверхность. Примерами таких предположений могут быть:

- поскольку мы не имеем информации о значениях оператора на не определенных нами наборах, будем считать искомую поверхность кусочно-линейной;
- желательно удобное хранение оператора агрегирования, поэтому будем считать, что искомая поверхность явля-

ется гладкой и описывается полиномом $k+1$ степени, где k — число известных точек в нашем пространстве.

Несомненные достоинства данного подхода — его простота и наглядность. Недостатками являются необходимость знать значение оператора агрегирования хотя бы на $n+1$ -наборе значений и невозможность в его рамках использования дополнительной информации о поведении оператора агрегирования. Такое представление работает только для достаточно хорошо определенных задач.

Логический подход основан на следующей интерпретации ОАИ. Если мы имеем конечное множество значений оценок в каждом узле дерева модели k , то мы можем представить оператора агрегирования информации как некоторую функцию k -значной логики. Если количество входов узла равно n , то в нем может быть использована в качестве оператора агрегирования одна из функций k -значной логики от n переменных. Обозначим множество функций через P_n^k . Таких функций очень много (k^n). Как выбрать одну из них для использования в качестве оператора агрегирования информации? Обычно эксперт может определить значение функции на некоторых наборах ее аргументов. В этом случае мы говорим о частично определенной функции. Пусть известны значения функции на t наборах. Обозначим класс таких функций через $P_{n,t}^k$. Число таких функций при большой разнице $k^n - t$ также является необозримым. Если эксперт может сформулировать содержательные условия на поведение искомой функции типа «При сильном возрастании первого аргумента значение функции слегка убывает», «При совместном возрастании аргументов 3 и 5 значение функции сильно возрастает» и т. п., мы можем говорить о нечетко заданных подклассах k -значной логики. Рассмотрим $P_{n,t}^k$ и обозначим через S множество нечетких условий на поведение функций из $P_{n,t}^k$. Можно сформулировать следующие задачи.

Задача 1. Являются ли условия S на поведение конкретной функции $f \in P_{n,t}^k$ совместимыми или противоречивыми?

Задача 2. Если условия непротиворечивы, можно ли каким-то образом описать класс функций $S(P_{n,t}^k)$, им удовлетворяю-

щих? В частности, можно ли предложить процедуру вычисления степени принадлежности любой $f \in P_{n,t}^k$ нечетким условиям S ?

Задача 3. Если условия S противоречивы, можно ли сформулировать условия S^- ($S^- \subset S$), максимально похожие на S и являющиеся непротиворечивыми?

Эти задачи детально рассмотрены в [43]. Они имеют положительное решение. Их описание требует введения специфических дополнительных понятий из теории нечетких множеств, k -значных логик и других разделов математики, поэтому мы не будем здесь его приводить; указанную выше работу можно найти в интернете. Обобщению этих результатов на случай более сложной структуры модели процесса — графа без циклов — посвящены работы [44] и [45], составившие основу кандидатской диссертации. В частности, были получены следующие результаты:

- для решения задачи выбора операторов агрегирования по экспертным описаниям введена и изучена новая управляющая система — граф нечеткого условия;
- сформулирована и решена задача устойчивости для иерархических систем в дискретном и нечетком случаях:
 - доказана NP-полнота задачи в общем случае;
 - выделены частные случаи задачи, допускаемые разрешимость за полиномиальное время;
 - получены эффективные алгоритмы выделения этих частных случаев;
- рассмотрена задача оптимального распределения ресурсов для иерархических систем в различных метриках:
 - доказана NP-полнота задачи в общем случае;
 - выделены частные случаи задачи, допускаемые разрешимость за полиномиальное время;
 - получены эффективные алгоритмы выделения этих частных случаев;
- для решения последних двух задач введено и изучено новое понятие для ориентированных ациклических графов — вершина-сверхдоминатор.

В рамках нечеткой модели происходит заполнение таблицы выше с использованием технологий обучения. Можно выделить два класса таких алгоритмов: обучение на основе генетических алгоритмов и обучение на основе нейронных сетей.

Генетические алгоритмы — достаточно стандартные процедуры в обучении любых нечетких моделей, не только нечетких операторов агрегирования информации. В частности, они способны быстро подбирать параметры модели, чтобы она удовлетворяла некоторому критерию качества. Генетические алгоритмы в задаче выбора оператора агрегирования информации применяются тогда, когда мы можем определить «смысл» операции агрегирования — это операция «и» или это операция «или». Для многих задач такая интерпретация агрегирования является естественной. Как мы говорили в разделе 2.3, таких операций бесконечно много: это t -нормы и t -конормы. Какую из них взять в качестве оператора агрегирования информации? У нас, как и в дискретном подходе, есть частично заполненная таблица истинности — берем в качестве ОАИ t -норму (если его «смысл» — логическое «и») или t -конорму (если его «смысл» — логическое «или»), которая обеспечивает необходимые значения для всех известных строк; если таких t -норм (t -конорм) несколько — берем любую. При возникновении конфликта (пользователь не согласен с тем значением, которые выдает оператор агрегирования для некоторой строки) надо заменить оператора. Мы ищем новую t -норму (или t -конорму), которая на всех известных значениях (строках) выдает необходимые значения оператора (включая бывшую «конфликтную строку»). Продолжаем использовать последнюю t -норму (или t -конорму) до возникновения нового конфликта. Если наше предположение о природе ОАИ (это операция «и» или это операция «или») верно, то, следуя описанной процедуре, мы найдем необходимого оператора в виде t -нормы или t -конормы.

Для выбора правильной t -нормы в описанной выше процедуре мы использовали строгие архимедовы t -нормы, задаваемые формулой

$$H_{\lambda} = \frac{\mu_A \times \mu_B}{\lambda + (1 - \lambda)(\mu_A + \mu_B - \mu_A \times \mu_B)}, \text{ завися-$$

щей от единственного параметра $\lambda (0 \leq \lambda \leq \infty)$. Такое представление описывает широкий класс t -норм; в частных

$$\text{случаях — это известные } t\text{-нормы: } H_0 = \frac{\mu_A \times \mu_B}{\mu_A + \mu_B - \mu_A \times \mu_B},$$

$$H_1 = \mu_A \times \mu_B, H_\infty = \begin{cases} \mu_A, & \text{если } \mu_B = 1 \\ \mu_B, & \text{если } \mu_A = 1 \\ 0 & \text{в других случаях} \end{cases}$$

(сравните с примерами t -норм из раздела 2.3). Строгая архимедова t -конорма выглядит следующим образом:

$$H_\lambda^* = \frac{\mu_A + \mu_B - (2 - \lambda) \times \mu_A \times \mu_B}{1 - (1 - \lambda) \times \mu_A \times \mu_B}, \text{ и для нее справедливо все, ска-}$$

занное выше для t -норм. Генетический алгоритм подбирал значения λ , при которых разрешался описанный выше конфликт между «теоретическим» значением ОАИ и мнением пользователя. Эта процедура впервые была кратко представлена в работе [46] и описана более детально в [43].

Если мы не знаем даже смысла операторов агрегирования информации в модели процесса, то все равно возможен подбор адекватного ОАИ на основе технологии искусственных нейронных сетей. Основой для такого оптимизма служит, как, впрочем, и для всей технологии искусственных нейронных сетей, теорема А. Н. Колмогорова. Это блестящее решение 13-й проблемы Гильберта представлено в 1957 году. Согласно этой теореме, любая непрерывная функция от n переменных может быть представлена в виде суперпозиции непрерывных функций от одной переменной и операции сложения следующим образом: $f(x_1, \dots, x_n) = \sum_{i=1}^{2n+1} g_i(\sum_{j=1}^n \varphi_{ji}(x_j))$, где g_i, φ_{ji} — непрерывные функции, причем φ_{ji} не зависят от выбора f . Эту формулу можно интерпретировать как нейронную сеть с одним скрытым слоем из $(2n+1)$ -нейронов. А это для нас означает, что если оператор агрегирования информации существует (существует $f(x_1, \dots, x_n)$), то можно (теоретически) построить нейронную сеть, которая его воспроизводит. Структура такой нейронной сети представлена на рис. 38. Более подробное описание дано в [43].

Для автоматического построения модели процесса нами также использовались методы углубленного анализа данных, которые показали свою эффективность [47].

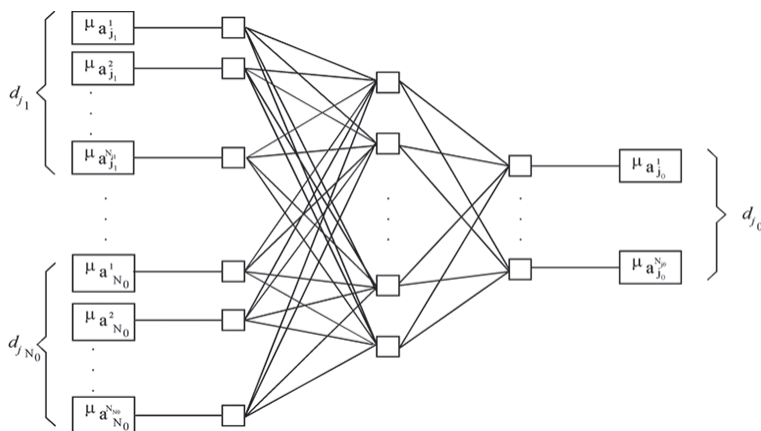


Рис. 38. Структура нейронной сети для ОАИ вершины d_{j_0} модели процесса

Суммируя вышесказанное, а также опыт разработки систем оценки и мониторинга, можно предложить следующие рекомендации для выбора адекватных операторов агрегирования информации в системах оценки и мониторинга сложных процессов:

1. Если есть уверенность в конкретной функции — используйте ее. Это справедливо для многих «простых», близких к инженерным или производственным, процессов (например, оператор *min* — см. анализ проблемы 3).
2. Если такой уверенности нет, но проблема мониторинга тем не менее достаточно понятна и возможны формулировки достаточно большого набора четких и нечетких условий на поведение оператора агрегирования информации — используйте формализацию на основе нечетко заданных подклассов k -значных логик. Это наиболее надежный способ.
3. Если формулировка таких условий невозможна, но есть хотя бы идеи о природе ОАИ (это логическое «и» или логическое «или») — используйте обучение на основе генетических алгоритмов. Если логика верна, то оператор агрегирования найдется: это будет одна из бесконечно-

го числа t -норм (если «смысл» ОАИ—логическое «и») или t -конорм (если его «смысл»—логическое «или»).

4. Если невозможны даже рассуждения п. 3, используйте обучение на основе искусственных нейронных сетей. Теоретически (теорема А. Н. Колмогорова) вы сможете найти неизвестную функцию агрегирования, если она объективно существует. Если такой функции нет в решаемой задаче, вы находитесь в ситуации, когда задача не решается на систематической основе даже человеком (каждый решает ее по-своему, нет общей логики). Такую систему построить нельзя—невозможно воспроизвести Леонардо да Винчи...
5. При наличии достаточно большого набора транзакционных данных возможно применение методов интеллектуального анализа данных (в частности, извлечения ассоциативных зависимостей или деревьев-решений) для построения прототипа модели процесса. Агрегирование будет определяться соответствующими параметрами используемых алгоритмов. Есть положительный пример такого использования, но нет массового опыта применения.

3.2.3. Аналитические возможности систем оценки и мониторинга

На основе описанных методов разработаны системы оценки и мониторинга процессов создания странами специальных видов технологий и материалов в интересах управления гарантиями Международного агентства по атомной энергии [48], риска развития сердечно-сосудистых заболеваний в Российской Федерации [49], создания изделий микроэлектроники (оценка технологических стартапов) в интересах компании *Cadence Design Systems* [50]. Возможности применения подобных систем в оценке и мониторинге кризисных ситуаций и в стратегическом управлении компаниями описаны в [51], [52] и [53] соответственно.

Мы не будем описывать упомянутые выше приложения. Это потребует описания соответствующих предметных обла-

стей и особенностей процессов, усилий читателя на понимание деталей, которые потом не пригодятся. Доступная статья [49] о применении описанной технологии в достаточно значимой всем области (здравоохранение) может быть рекомендована читателю, интересующемуся деталями. Вместо этого мы обсудим уникальные аналитические возможности систем оценки и мониторинга.

Итак, пусть у нас есть рабочая версия системы оценки и мониторинга (стабильная структура модели и настроенные операторы агрегирования информации). В этом случае возможно решение прямой и обратной задач.

Прямая задача заключается в нахождении критических путей — таких элементов модели, малое изменение которых приводит к изменению всех вышележащих узлов, включая верхний. Практическая значимость критических путей обсуждена в разделе 3.2.1. Для большого класса операторов агрегирования информации возможно вычисление степени критичности каждого элемента модели. Задача также может решаться перебором для достаточно компактной модели.

Обратная задача позволяет оптимизировать бюджет на достижение определенного уровня целевого показателя. Если задан бюджет и мы знаем стоимость изменения состояния узла модели, то возможно нахождение тех узлов, изменение которых находится в рамках заданного бюджета и дает максимальной эффект повышения корневого узла модели. Формализовать такую постановку можно следующим образом. Пусть модель содержит N узлов, к i -му узлу приписано m_i значений: $a_i = \{a_i^1, a_i^2, \dots, a_i^{m_i}\}$ ($i = 0, 1, \dots, N-1$) (рис. 39). Пусть задано: X — бюджет, c_i — стоимость изменения состояния i -го узла (перевода его из состояния a_i^k в состояние a_i^{k+1} , $k = 1, \dots, m_i - 1$; для простоты записи будем считать такую стоимость одинаковой для всех k). Обозначим через Δa_i силу изменения узла a_i . Эта величина равна q для ситуации, когда состояние узла поменялось с a_i^k на a_i^{k+q} . Для множества номеров узлов $I = \{1, 2, \dots, N-1\}$ обозначим $S(I)$ множество всех его непустых подмножеств.

Задача 1. Найти множество номеров узлов $s \in S(I)$: $\Delta a_0 \rightarrow \max$, $\sum_{i \in s} c_i \leq X$.

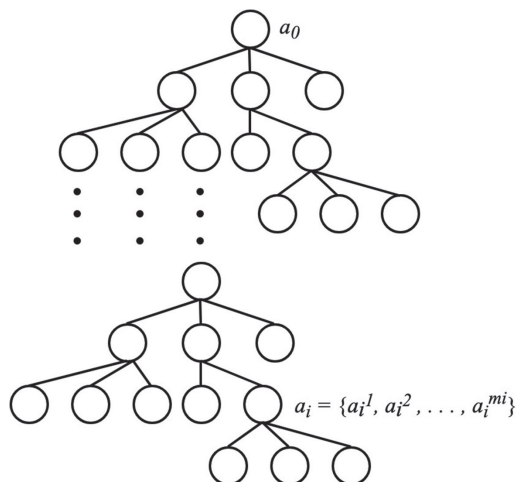


Рис. 39. Пример модели процесса

Это задача нахождения максимального эффекта в рамках заданного бюджета.

При заданных условиях возможно и решение сопряженной задачи.

Задача 2. Найти множество номеров узлов $s \in S(I): \sum_{i \in s} c_i \rightarrow \min, \Delta a_i \geq q$.

Это задача нахождения минимального бюджета, позволяющего достичь необходимого эффекта.

Примерами таких задач при наличии соответствующих моделей могут быть оценка и повышение капитализации технологических стартапов [54]; повышение рейтинга банков, университетов, спортивных команд (обычно структура рейтинга известна — например [55]); повышение инвестиционной привлекательности компаний, регионов, стран (обычно рейтинговые агентства публикуют структуру оценки); снижение уровня технологических рисков опасных производств (при наличии карты технологических рисков); повышение устойчивости развития бизнеса [56] и многие другие задачи.

Перечисленные задачи можно решать как в рамках фрагментов модели (например, оптимизация работы соответ-

вующих подразделений компании), так и модели в целом (например, оптимизация работы компании в целом). Такая возможность определяется тем, что модель в виде дерева или графа без циклов можно «распилить» как по горизонтали (рис. 40), так и по вертикали (рис. 41). Тогда каждый аналитик будет работать со своим фрагментом модели (выполнять свои функциональные обязанности), но «выход» его модели может быть одним из «входов» модели более высокого уровня, как представлено на рис. 40, 41.

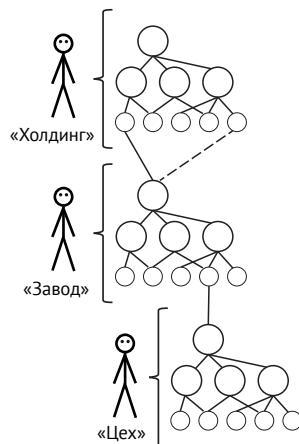


Рис. 40. Пример декомпозиции модели по горизонтали

На рисунках приведен пример производственного холдинга. На самом деле сказанное справедливо для любых иерархических систем управления. Например, заменяя слова «Цех» на «Управа», «Завод» на «Город», «Холдинг» на «Регион», мы получаем представление для системы госуправления; решение социальных и политических проблем на основе систем их оцен-

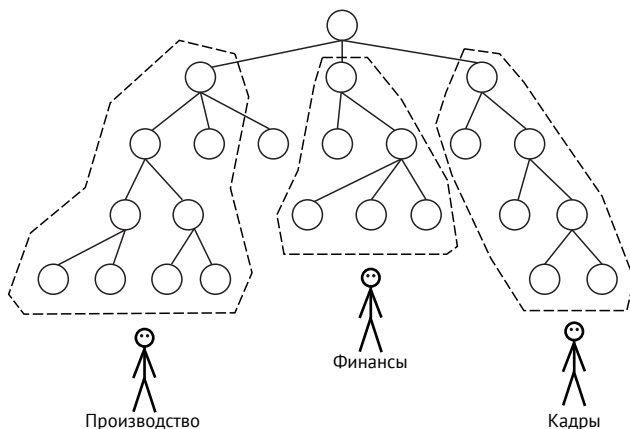


Рис. 41. Пример декомпозиции модели по вертикали

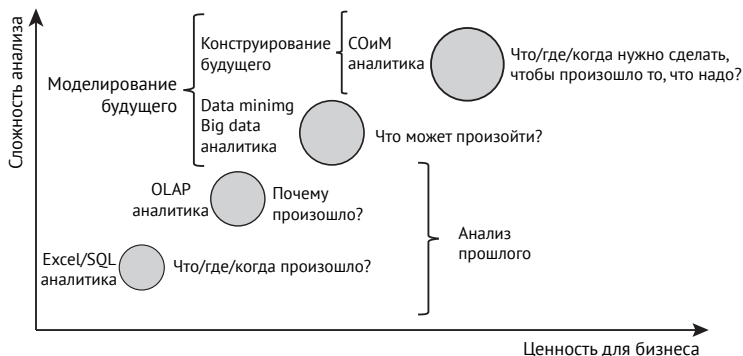


Рис. 42. Эволюция аналитических инструментов

ки и мониторинга позволит иметь объективную актуальную картину на любой момент времени, историю развития процессов, «вычислять» (то есть иметь количественную оценку) их сильные и слабые стороны для текущего состояния (критические пути), оптимизировать расходование бюджетных средств для достижения нужного результата (обратные задачи).

Место систем оценки и мониторинга в контексте аналитических инструментов представлено на рис. 42.

Текущий уровень почти всех компаний — OLAP аналитика; лишь немногие компании используют Big Data аналитику; аналитические инструменты систем оценки и мониторинга используются в единичных международных и правительственных организациях, а также в крупном международном высокотехнологическом бизнесе. Использование аналитических возможностей систем оценки и мониторинга — это естественный эволюционный этап развития аналитических инструментов.

3.2.4. Возможности и ограничения систем оценки и мониторинга

Суммируя функционал систем оценки и мониторинга, можно утверждать, что они:

- решают задачу оценки и мониторинга (раздел 3.2.1);

3. Сценарии использования гибридного интеллекта

- позволяют аналитику вводить информацию из всех доступных источников естественным образом;
- хранят историю развития процесса, оценивают его текущее состояние, прогнозируют и моделируют будущее;
- позволяют решать оптимизационные задачи достижения максимального эффекта в рамках заданного бюджета и поиска минимального бюджета для достижения заданного эффекта;
- допускают декомпозицию модели как по горизонтали, так и по вертикали, создавая возможность наращивания экспертного и аналитического потенциалов, практически ограниченных только масштабом владельца системы («страна», «холдинг»).

Когда такие системы эффективны? Они эффективны:

- когда нет / нельзя построить математическую модель процесса в виде уравнений, автоматов и пр. (то есть когда не работают классические методы);
- есть аналитики, решающие задачу оценки и мониторинга на систематической основе (то есть решение существует; деятельность ближе к «инженерии», чем к «искусству»).

При каких условиях разработка систем оценки и мониторинга возможна? Она возможна:

- когда можно построить семантическую модель процесса в виде набора понятий и их взаимосвязей («что из чего состоит»);
- поступает и анализируется реальная информация — возможны обучение и настройка.

Отметим, что возможна разработка оптимальных систем оценки и мониторинга точки зрения:

- удобства ввода информации;
- согласованности оценок;
- информационной поддержки ввода и моделирования.

К близким технологиям можно отнести системную динамику и нечеткие когнитивные карты. Мы не будем их описы-

вать—есть много популярной и специальной литературы по данным направлениям. Отметим, однако, отличие систем оценки и мониторинга от указанных типов систем: у них нет возможностей обработки информационного потока, они имеют только подсистему моделирования в терминах систем оценки и мониторинга.

Когда разработка систем оценки и мониторинга имеет смысл, а когда она бессмысленна? Она осмысленна, если мы с их помощью автоматизируем наблюдение за процессами, которые уже делаются людьми на систематической основе. Это означает, что функция агрегирования информации существует (люди как-то решают задачу). Если такое условие не выполняется (например, мы пытаемся решать задачи, с которыми никто ранее не сталкивался), то разработка системы оценки и мониторинга невозможна.

Возвращаясь к взглядам McKinsey и IBM, мы можем отметить, что системы оценки и мониторинга позволяют «взаимодействовать с машиной так, как это было бы с коллегой» ([40], с. 41), и являются одним из вариантов «практических приложений ИИ, которые помогают людям с четко определенными задачами».

Мы видим в таких системах также выход из кризиса современного искусственного интеллекта—это не системы вместо человека, это системы-партнеры, усиливающие аналитические возможности специалиста при выполнении им своих задач.

Список источников

1. *Тьюринг А.* Может ли машина мыслить? М.: Государственное изд-во физико-матем. литературы, 1960; http://www.etheroneph.com/files/can_the_machine_think.pdf.
2. *Мак-Коллок У.С., Пуммс В.* Логическое исчисление идей, относящихся к нервной активности. Автоматы / пер. с англ. М.: Изд-во иностр. литературы, 1956.
3. *Newell A., Shaw J.C., Simon H.A.* Report on a general problem-solving program. Proceedings of the International Conference on Information Processing, 1959; <http://digitalcollections.library.cmu.edu/awweb/awarchive?type=file&item=356096>.
4. *Розенблатт Ф.* Принципы нейродинамики / пер. с англ. М.: Мир, 1966.
5. *Минский М., Пейперт С.* Перцептроны / пер. с англ. М.: Мир, 1971.
6. *Заде Л.А.* Понятие лингвистической переменной и его применение к принятию приближенных решений. М.: Мир, 1976.
7. *Рыжов А.П.* Модели поиска информации в нечеткой среде. М.: Изд-во Центра прикладных исследований при механико-математическом факультете МГУ, 2004; <http://www.intsys.msu.ru/staff/ryzhov/FuzzyRetrieval2010.htm>.
8. *Рыжов А.П.* Элементы теории нечетких множеств и измерения нечеткости. М.: Диалог-МГУ, 1998; <http://www.intsys.msu.ru/staff/ryzhov/FuzzySetsTheoryApplications.htm>.
9. *Рыжов А.П.* О качестве классификации объектов на основе нечетких правил // Интеллектуальные системы. Т. 9. Вып. 1–4. М.: МНЦ КИТ, 2005; <http://www.intsys.msu.ru/staff/ryzhov/>.

10. *Елистратов В.* Количество активных смартфонов и планшетов превысило население Земли. 7 октября 2014: <https://tjournal.ru/52596-more-phones-than-people>.
11. *Reinsel D., Gantz J., Rydning J.* The digitization of the world: from edge to core. November 2018: <https://www.seagate.com/files/www-content/our-story/trends/files/idc-seagate-dataage-whitepaper.pdf>.
12. The Biggest Marketing Trend of 2017: Personalization. January 31, 2017: <https://fusiononemarketing.com/biggest-marketing-trend-2017-personalization/>.
13. *Baker D.* How Personalization Is Changing Content Marketing. March 31st, 2017: <https://contently.com/2017/03/31/personalization-changing-content-marketing/>.
14. *Rund J.* HOW PERSONALIZATION WILL EVOLVE IN 2018, January 29, 2018: <http://adage.com/article/epsilon/personalization-evolve-2018/312048/>.
15. *Hyken S.* Recommended Just For You: The Power Of Personalization. May 13, 2017: <https://www.forbes.com/sites/shep-hyken/2017/05/13/recommended-just-for-you-the-power-of-personalization/#613481826087>.
16. Discover the Power of Personalization: <https://www.demandware.com/resources/power-of-personalized-shopping>.
17. *Young H.* Welcome to 2018, The Year of Personalization. Here's 3 Things You Need to Know. December 21, 2017: <https://www.demandware.com/blog/retail-intelligence/welcome-2018-year-personalization-heres-3-things-need-know>.
18. *Wixey N.* Made-to-order: The rise of mass personalisation. The Deloitte Consumer Review: <https://www2.deloitte.com/tr/en/pages/consumer-business/articles/made-to-order-the-rise-of-mass-personalisation.html>.
19. *Betts A.* A new era of personalization: The hyperconnected customer experience. January 23, 2018: <https://martechtoday.com/new-era-personalization-hyper-connected-customer-experience-209529>.
20. *Robinson S.* Personalization: The Next Big E-Learning Trend: <https://www.td.org/magazines/td-magazine/personalization-the-next-big-e-learning-trend>.

21. Fashion in 2018. Getting Personal. BY BOF TEAM AND MCKINSEY & COMPANY, JANUARY2, 2018: <https://www.businessoffashion.com/articles/intelligence/top-industry-trends-2018-4-getting-personal>.
22. SPECIAL REPORT: PERSONALIZATION TECH TRENDS: <https://www.chiefmarketer.com/special-reports/special-report-personalization-tech-trends/>.
23. *Kuin T.* How to Start with Digital Marketing Personalization. 23 February, 2018: <https://www.accenture-insights.nl/en-us/articles/digital-market-personalization>.
24. *Mouradian N.* The Dieline's 2018 Trend Report: Brands Become Hyper-Personalized. February 27, 2018: <http://www.thedieline.com/blog/2018/2/26/brands-become-hyper-personalized>
25. *Lyapin B., Ryjov A.* A Fuzzy Linguistic Interface for Data Bases in Nuclear Safety Problems. Fuzzy Logic and Intelligent Technologies in Nuclear Science. Proceedings of the 1st International FLINS Workshop, Mol, Belgium, September 14–16, 1994. Edited by Da Ruan, Pierre D'hondt, Paul Govaerts, Etienne E. Kerre, World Scientific.
26. *Ryjov A.* Personalization and Optimization of Information Retrieval: Adaptive Semantic Layer Approach. In: Zadeh L., Yager R., Shahbazova S., Reformat M., Kreinovich V. (eds) Recent Developments and the New Direction in Soft-Computing Foundations and Applications. Studies in Fuzziness and Soft Computing. Vol. 361. Springer, Cham, 2018.
27. *Рыжов А.П., Журавлев А.Д., Вахов А.Н., Кривцов В.В.* Об одном подходе к персонификации обучения в рамках компьютерных обучающих систем // Интеллектуальные системы. Теория и приложения. Т. 20. Вып. 3. М., 2016; http://intsysjournal.ru/articles/is2003/31_Ryzhov.pdf;
28. *Ryjov A., Vakhov A., Krivtsov V. and Zhuravlev A.* Personalization and optimization of learning based on technology for evaluation and monitoring of complex processes: Uchi.ru case study. The 2016 International Conference on Computational Science & Computational Intelligence. Ed. by: Hamid R. Arabnia, Leonidas

- Deligiannidis, and Mary Yang. Las Vegas, Nevada, USA, 15–17 December, 2016.
29. Рыжов А.П., Кривцов В.В., Журавлев А.Д. Некоторые задачи кластеризации и ранжирования для персонификации компьютерных обучающих систем. Дистанционные образовательные технологии: Материалы I Всероссийской научно-практической интернет-конференции (г. Ялта, 19–23 сентября 2016 года). Симферополь: ИТ «АРИАЛ», 2016.
30. Рыжов А.П., Новиков П.А. Об одной модели цифровых привычек // Интеллектуальные системы. Теория и приложения. Т. 21. Вып. 3. М., 2017; <http://intsysjournal.ru/pdfs/21-4/91-100-RyzovNovikov.pdf>.
31. Рыжов А.П. Некоторые задачи оптимизации и персонификации социальных сетей. Saarbrucken, LAP, 2015.
32. Ryjov A. Personalization of Social Networks: Adaptive Semantic Layer Approach. In: Social Networks: A Framework of Computational Intelligence. Witold Pedrycz and Shyi-Ming Chen (Eds.). Springer International Publishing Switzerland 2014.
33. FSQL (Fuzzy SQL): <http://www.lcc.uma.es/~ppgg/FSQL.html#Ref>.
34. Galindo J. (ed.). Handbook of Research on Fuzzy Information Processing in Databases (2 Vol.). IGI Global, 2008.
35. Dunn J.C. A fuzzy relative of the ISODATA process and its use in detecting compact well separated clusters // Journal of Cybernetics. 1974. Vol. 3. № 3.
36. Bezdek J.C. Pattern recognition with fuzzy objective function algorithm. Kluwer Academic Publishers Norwell, MA, USA, 1981.
37. Рыжов А.П., Огородников Н.М. Об одном методе персонализации поиска информации // Интеллектуальные системы. Теория и приложения. Т. 22. Вып. 4. М., 2018.
38. Алисейчик П.А., Вашик К., Кнап Ж., Кудрявцев В.Б., Строгалов А.С., Шеховцов С.Г. Компьютерные обучающие системы // Интеллектуальные системы. Т. 8. 2004.
39. A gallery of disruptive technologies http://www.mckinsey.com/assets/dotcom/mgi/slideshows/disruptive_tech/index.html#.

40. *Manyika J., Chui M., Bughin J., Dobbs R., Bisson P., Marrs A.* Disruptive technologies: Advances that will transform life, business, and the global economy. McKinsey Global Institute (MGI), May 2013; http://www.mckinsey.com/insights/business_technology/disruptive_technologies.
41. Ed Tech Market Map by Flybridge Capital Partners: <https://prezi.com/xguiky7u7aur6/ed-tech-market-map/>.
42. *Глушков В.М.* Кибернетика: <http://victor-safronov.ru/systems-analysis/glossary/cybernetics.html>.
43. *Рыжов А.П.* Об агрегировании информации в нечетких иерархических системах // Интеллектуальные системы. Т. 6. Вып. 1–4. 2001; [http://www.intsys.msu.ru/magazine/archive/v6\(1-4\)/ryzhov.pdf](http://www.intsys.msu.ru/magazine/archive/v6(1-4)/ryzhov.pdf).
44. *Лебедев А.А.* О задачах оптимального распределения ресурсов и проверки устойчивости для схем функциональных элементов в k-значной логике // Интеллектуальные системы. Теория и приложения. Т. 18. Вып. 4. 2014; http://intsysjournal.ru/articles/is04/08_lebedev.pdf.
45. *Лебедев А.А.* Синтез операторов агрегирования информации по экспертным описаниям // Интеллектуальные системы. Теория и приложения. Т. 14. Вып. 1–4. 2012; [http://www.intsys.msu.ru/magazine/archive/v16\(1-4\)/lebedev-025-048.pdf](http://www.intsys.msu.ru/magazine/archive/v16(1-4)/lebedev-025-048.pdf).
46. *Рыжов А.П., Федорова М.С.* Генетические алгоритмы в задаче выбора операторов агрегирования информации в системах информационного мониторинга // V Всероссийская конференция «Нейрокомпьютеры и их применение»: сб. докладов. Москва, 17–19 февраля 1999 года. М., 1999.
47. *Рыжов А.П., Распорзиев В.В.* Методы извлечения нечетких ассоциативных правил в системах информационного мониторинга // Труды международных научно-технических конференций «Интеллектуальные системы» и «Интеллектуальные САПР», 3–10 сентября 2006 г., Дивноморское, Россия. Т. 1. М.: Физматлит, 2006.

48. *Rylov A., Belenki A., Hooper R., Pouchkarev V., Fattah A. and Zadeh L.A.* Development of an Intelligent System for Monitoring and Evaluation of Peaceful Nuclear Activities (DISNA), IAEA, STR-310. Vienna, 1998.
49. *Ахмеджанов Н.М., Жукоцкий А.В., Кудрявцев В.Б., Оганов Р.Г., Расторгуев В.В., Рыжов А.П., Строгалов А.С.* Информационный мониторинг в задаче прогнозирования риска развития сердечно-сосудистых заболеваний // Интеллектуальные системы. Т. 7. Вып. 1–4. 2003; [http://www.intsys.msu.ru/magazine/archive/v7\(1-4\)/ryzhov-005-038.pdf](http://www.intsys.msu.ru/magazine/archive/v7(1-4)/ryzhov-005-038.pdf).
50. *Лебедев А.А., Рыжов А.П.* Оценка и мониторинг проектов разработки высокотехнологических изделий микроэлектроники // Известия ТРТУ. Тематический выпуск «Интеллектуальные САПР». 2006. № 8.
51. *Rylov A.* Basic principles and foundations of information monitoring systems // Monitoring, Security, and Rescue Techniques in Multi-Agent Systems. Springer, 2005.
52. *Sunkpho J., Wipulanusat W., Kokkaew N., Rylov A.* Disaster Management based on Information Monitoring Technology // Creative Construction Conference 2014 Proceeding, June 21–24, 2014, Prague.
53. *Rylov A.* Information Monitoring Systems as a Tool for Strategic Analysis and Simulation in Business // International Conference on Fuzzy Sets and Soft Computing in Economics and Finance FSSCEF 2004 Proceedings, Saint-Petersburg, Russia, June 17–20, 2004.
54. *Рыжов А.П., Анохина Е.С.* Об одном подходе к оценке технологических стартапов средствами технологии оценки и мониторинга сложных процессов. Материалы VII ежегодной международной научно-практической конференции «Инвестиции. Инновации. Информационные технологии». Москва, 27–28 февраля 2015 г., РАНХиГС. М., 2015.
55. *WASEDA — IAC10th International E-Government Ranking 2014:* http://www.e-gov.waseda.ac.jp/pdf/2014_E-Gov_Press_Release.pdf.
56. *Рыжов А.П.* Оценка и мониторинг бизнеса на основе сбалансированной системы показателей // Материалы V еже-

Список источников

годной международной научно-практической конференции «Инвестиции. Инновации. Информационные технологии». Москва, 1–2 марта 2013 г., РАНХиГС. М., 2013.

Научное издание

Серия «Библиотека Школы IT-менеджмента»

Заказное издание

Александр Павлович Рыжов

**Гибридный интеллект.
Сценарии использования в бизнесе**

Выпускающий редактор *Е. В. Попова*

Корректор *Г. А. Лакеева*

Оригинал-макет *О. З. Элоева*

Верстка *Е. В. Немешаевой*

Издательский дом «Дело» РАНХиГС

119571, Москва, пр-т Вернадского, 82

Дизайн обложки *К. С. Артёмов*

Издательство «Академиздат»

630132, Новосибирск, пр-т Димитрова, 7, оф. 514

Тел. +7 (383) 263-24-88, 8 (800) 201-70-63

E-mail: knigi@academizdat.ru

Сайт: www.academizdat.ru

Подписано в печать 14.08.2019. Формат 60×90/16

Гарнитура PT Serif Pro. Усл. печ. л. 7,25.

Тираж 500 (1 завод — 100) экз. Заказ № 320

Отпечатано в типографии издательства «Академиздат»

630058, Новосибирск, ул. Плотинная, 7

Данная книга открывает серию «Библиотека Школы IT-менеджмента», в рамках которой планируются к публикации результаты исследований ведущих профессоров и преподавателей Школы IT-менеджмента Института экономики, математики и информационных технологий Российской академии народного хозяйства и государственной службы при Президенте РФ.

Книга написана на основе части курса по аналитическим информационным технологиям, читаемом автором в течение 10 лет в Школе IT-менеджмента и части курса по теории нечётких множеств, читаемом автором более 20 лет на механико-математическом факультете МГУ имени М. В. Ломоносова.

Книга ориентирована на IT-менеджеров, специалистов по разработке и внедрению интеллектуальных цифровых технологий, а также на их реальных и потенциальных пользователей.



Рыков Александр Павлович — к. ф.-м. н., д. т. н., профессор, MBA International Executive Development Center — Bled School of Management. Заведующий кафедрой Школы IT-менеджмента Института ЭМИТ РАНХиГС при Президенте РФ, профессор кафедры

МаТИС Механико-математического факультета МГУ имени М. В. Ломоносова, профессор Национального медицинского исследовательского центра профилактической медицины Министерства здравоохранения РФ. Автор более 100 научных работ, член программных и организационных комитетов более 90 международных научных конференций.

> ISBN 978-5-6043351-0-9



9 785604 335109



Новосибирск, 2019

→
ПОДРОБНЕЕ
ОБ АВТОРЕ

