

БИОНИКА ИНТЕЛЛЕКТА

ИНФОРМАЦИЯ, ЯЗЫК, ИНТЕЛЛЕКТ

№ 2 (87)

2016

НАУЧНО-ТЕХНИЧЕСКИЙ ЖУРНАЛ

Основан в октябре 1967 г.

Учредитель и издатель

Харьковский национальный университет радиоэлектроники

Периодичность издания — 2 раза в год

СОДЕРЖАНИЕ

СТРУКТУРНАЯ, ПРИКЛАДНАЯ И МАТЕМАТИЧЕСКАЯ ЛИНГВИСТИКА

<i>Широков В.А.</i> Лингвистика и системный подход. Часть 2.....	3
<i>Лазаренко О.В., Панченко Д.И., Айвас Е.Ю.</i> Моделирование базовых составляющих процесса понимания текста в системе автоматического реферирования	12
<i>Чалая Л.Э., Удовенко С.Г., Кушвид Е.В.</i> Метод двухэтапной классификации электронных текстов	16
<i>Пузик А.С.</i> Лексикографическая система электронного терминологического словаря.....	24

ТЕОРЕТИЧЕСКИЕ ОСНОВАНИЯ ИНФОРМАТИКИ И КИБЕРНЕТИКИ. ТЕОРИЯ ИНТЕЛЛЕКТА

<i>Ерохин А.Л., Нечипоренко А.С., Бабий А.С., Турута А.П.</i> Применение глубоких сверточных нейронных сетей для классификации риноманометрических данных	30
<i>Свиридов А.С., Завизиступ Ю.Ю., Михаль О.Ф.</i> Методологические аспекты сегментной обработки Кирлиан-объектов.....	35
<i>Михаль О.Ф.</i> Эволюционирование мультиагентной системы как аналог формирования индивидуального человеческого интеллекта	42

МАТЕМАТИЧЕСКОЕ МОДЕЛИРОВАНИЕ. СИСТЕМНЫЙ АНАЛИЗ. ПРИНЯТИЕ РЕШЕНИЙ

<i>Каграманян А.Г., Четвериков Г.Г., Шляхов В.В.</i> Мультиалгебраические дискретные системы при наличии алгебраической структуры носителя. Сообщение 1	48
<i>Каграманян А.Г., Четвериков Г.Г., Шляхов В.В.</i> Мультиалгебраические дискретные системы при наличии алгебраической структуры носителя. Сообщение 2	53
<i>Чумаченко Д.И., Яковлев С.В.</i> О нечетких рекуррентных отображениях при мультиагентном моделировании популяционной динамики	59
<i>Левыкин И.В.</i> Разработка обобщенной модели процесса решения задачи с интервальным представлением времени	64
<i>Михаль О.Ф., Лебедев О.Г.</i> Мультиагентная интерпретация генерации событий в модели системы массового обслуживания	69
<i>Литвин О.М., Ярмош О.В., Чорна Т.І.</i> Метод сплайн-інтерлінації при знаходженні найбільших (найменших) значень функції двох змінних в замкнутій області	77

ИНТЕЛЛЕКТУАЛЬНАЯ ОБРАБОТКА ИНФОРМАЦИИ. РАСПОЗНАВАНИЕ ОБРАЗОВ

<i>Гороховатский В.А., Столяров В.С.</i> Классификация изображений на основе кластерного представления структурных описаний	83
<i>Литвин О.М., Литвин О.О., Лісний Г.Д., Славік О.В.</i> Відновлення зображень в зонах відсутності піксельної інформації з використанням інтерстріпації функцій.....	88
<i>Ляшенко В.В., Кобылин О.А., Томич С.Н.</i> Основные этапы обработки изображений цитологических препаратов с использованием идеологии вейвлетов.....	94
<i>Чалая О.В.</i> Метод обобщения представления знание-емкого бизнес-процесса	101
<i>Шелехов І.В., Прилепа Д.В.</i> Визначення лінії симетрії на зображенні обличчя людини при діагностуванні емоційно-психічного стану	106
<i>Васильєв А.В., Довбиш А.С., Кулік Є.С., Козлов З.В.</i> Інформаційно-аналітична система адаптації навчального контенту випускової кафедри до вимог ринку праці	111

Рефераты	117
Об авторах	125
Правила оформления рукописів для авторів науково-технічного журналу «Біоніка інтелекту».....	127
Instructions for authors of manuscripts of the scientific journal «Bionics of intelligence»	128
Приклад оформлення статті.....	129

**ЛИНГВИСТИКА И СИСТЕМНЫЙ ПОДХОД.
ЧАСТЬ 2**

В статье получили дальнейшее развитие принципы системного подхода, основанные на рассмотрении отношений внутри триады «структура – субстанция – субъект». Определено понятие субстанции мысли и естественного языка как инструмента для формализации мыслительного процесса. Дается определение искусственного интеллекта как формы индивидуализации технических систем, обладающих языковым статусом. Приведены десять системных теорем.

**СИСТЕМНЫЙ ПОДХОД, ЛИНГВИСТИЧЕСКИЕ ПРЕЗУМПЦИИ, ФЛЕКТИВНЫЕ ЯЗЫКИ,
СИСТЕМНЫЕ ТЕОРЕМЫ, ИНФОРМАЦИЯ****Введение**

Данная статья представляет собою непосредственное продолжение работы, опубликованной в № 1(84) «Бионики интеллекта», где на основе феноменологического анализа свойств языка были сформулированы принципы системного подхода, основанного на рассмотрении отношений внутри триады «структура – субстанция – субъект». В этой статье изложенные представления детализируются и иллюстрируются на конкретных примерах. Предлагаются также некоторые обобщения, выходящие за рамки собственно языковой системы и, по мнению автора, имеющие общенаучное значение.

О «лингвистических презумпциях».

В ходе исследования автор пришел к неким утверждениям, которые были положены им в основу изложения. Данные соглашения представляют собой комплекс постулатов, своеобразных «презумпций», которым автор – явно или неявно – предполагает следовать в данной статье.

Список отмеченных презумпций предполагается открытым. На данный момент в списке лингвистических презумпций автор насчитывает семь предложений, а именно, следующие:

- 1) любая мысль может быть выражена в естественной языковой форме;
- 2) естественный язык является инструментом, который обеспечивает разнообразные преобразования информации;
- 3) преобразование информации в языке осуществляется в двух основных аспектах – когнитивном и коммуникативном;
- 4) между двумя любыми языками можно установить эквивалентность в смысле предложений 1 – 3;
- 5) математика является разновидностью языка;
- 6) язык представляет собой систему в смысле, который будет прояснен ниже;
- 7) основными системообразующими отношениями языка есть отношения «субъект – объект» и «форма – содержание».

Следует отметить, что сам подход к формулированию изложенных презумпций выкристаллизовался, когда автор пытался ответить для себя на следующие вопросы:

Что такое «система» в лингвистике и как она соотносится с понятием «системы» вообще?

Как соотносятся между собой информация и язык?

Как следует трактовать «Форму» и «Содержание» в языке. Что такое лингвистическая (внелингвистическая) форма и лингвистическое (внелингвистическое) содержание?

Каковы границы формализации языковой системы?

Каково возможное расширение границ концептуального описания языковой системы?

Автор отмечает, что не имеет полных, достаточно внятных ответов на поставленные вопросы, которые могли бы полностью удовлетворить читателя. Более того, автору не удалось преодолеть даже ряд внутренних противоречий в предмете его исследования. Этим объясняется не полная и даже не достаточная связность текста и фрагментарность статьи. Сознательно относясь к отмеченной внутренней противоречивости текста, автор обращается к читателю с просьбой о снисхождении, тем более, что она (противоречивость) при желании вполне может быть отнесена к факторам, стимулирующим развитие научных представлений о природе языка.

1. Информация и язык

Первую из лингвистических презумпций мы сформулировали следующим образом: *«Любая мысль может быть выражена в естественной языковой форме»*. Это неоспоримое, а возможно и сомнительное с общеполитической точки зрения утверждение мы призываем воспринимать как приглашение к рассуждению на темы: что такое мысль?, что такое язык (главным образом, естественный)?, что такое форма языка и как в ней отображается содержание (собственно, «мысль»)?

Можно сформулировать поставленные вопросы немного по-другому, а именно. «Субстанция» мысли, как мы полагаем, *может* быть отображена на «субстанцию» языка, то есть «субстанция» мысли имеет достаточно адекватное выражение в языковых формах. Это означает, что язык обладает потенциями и соответствующими внутренними системными ресурсами для того чтобы достаточно надёжно и уверенно порождать формы, в которых может быть выражена мысль.

С другой стороны, он (язык) сам находит своё содержание: а) во внеязыковой реальности, как форма отображения ментальных (или психоментальных, а также и сенсорно-перцептивных) процессов и б) в собственно языковой реальности или в языковой *системе*, о чем речь пойдет ниже.

Мы следуем информационному подходу к теории языка. Более того, в Презумпции № 2 мы позволили себе утверждать, что *«Естественный язык является инструментом, который обеспечивает разнообразные преобразования информации»*. В этой связи считаем необходимым объяснить, какой информационной доктрине мы предполагаем следовать в своем изложении. Ведь, как известно, существует очень большое число взглядов и подходов к такому предмету как информация, которые приобретают всё большую остроту в связи с особенностями современного этапа развития цивилизации, трактуемого антропологами культуры как «информационная эпоха».

Действительно, мир эволюционирует. Он вступил в сетевую или, как ее еще называют, сетцентрическую фазу эволюции. Всемирная Сеть стала именно той средой, где разыгрывается когнитивно-коммуникативный сценарий развития цивилизации. Экономика, основанная на знаниях, ежечасно и ежесекундно требует колоссальных объемов новой научно-технической, бизнес-, оперативной и иной информации, производство которой также включено в общий когнитивно-коммуникативный процесс.

По данным известной фирмы IDC еще в 2011 году общий мировой объем созданных и реплицированных человечеством данных составил более 1,8 зеттабайт (18 триллионов Гб). По прогнозам этой же фирмы, объем данных на планете будет, как минимум, удваиваться каждые два года до 2020 года. Если этот тренд считать справедливым — а на данный момент нет никаких оснований сомневаться в нем, то уже сейчас объем мировых, адаптированных Сетью данных превысил величину 10^{14} , а то и 10^{15} Гб. Если, следуя А.Колмогорову¹, считать информацию перифразой понятия

сложности, то приведенный факт ярко подтверждает основное эмпирическое правило (необходимое условие) общей теории эволюции: *система эволюционирует только тогда, когда ее сложность возрастает*.

Однако возникают вполне обоснованные сомнения, что такие большие объемы информации могут быть эффективно обработаны, а главное — адекватно восприняты и осознаны их реципиентами-адресатами в реальном времени современного бытия. Ведь психо-ментально-физиологическая природа человека, в том числе, ее способность к восприятию и обработке информации вряд ли существенно изменилась со времен Адама. Основной когнитивный тракт человека, а именно: *«восприятие → ощущение → переживание → осознание → понимание → рефлексия → реакция»*² продолжает так же действовать, как и десятки тысяч лет назад, хотя, так сказать, информационный антураж цивилизации изменился коренным образом. Экспоненциальный рост информации не сопровождается адекватным ростом человеческих возможностей по ее усвоению, а тем более эффективному использованию.

Таким образом, возникает главное противоречие современной эпохи: *закон эволюции требует роста объемов информации, производимой и воспринимаемой человечеством, а общество, рассматриваемое как сумма индивидов и отношений между ними, в силу особенностей человеческой природы не в состоянии должным образом воспользоваться этими объемами и вынуждено ограничивать их производство, обработку и использование*. Эволюция как будто сама «включает» механизмы своего торможения, которые иногда проявляются в весьма драматических формах.

Насущная необходимость ответа на эти вызовы стимулировала новые подходы к оперированию сверхбольшими массивами информации и привела к концепции «*Big Data*» или «*Больших данных*», которые символически характеризуются как «три V», а именно: *Volume* («объем» — петабайты хранимых данных), *Velocity* («скорость» — получение данных, преобразование, загрузка, анализ, обработка и реакция в реальном времени) и *Variety* («разнообразие» — обработка структурированных, полуструктурированных и неструктурированных данных различных типов). В последнее время к приведенным «трем V» добавляют еще два: *Veracity* — достоверность данных, которой пользователи начали придавать все большее значение, и *Value* — ценность добытой и накопленной информации. Сейчас «*Большие Данные*» признаются вторым по

¹ А.Н.Колмогоров. Три подхода к определению понятия «количество информации». В кн. «Теория информации и теория алгоритмов». — М.: Наука, 1987. — С.220.

² Palagin A.V., Shyrokov V.A. Principles of cognitive lexicography. //Informational theories & application. — 2000. — Vol.9. — № 2. — P. 43–51.

значимости (после виртуализации) трендом в информационно-технологической инфраструктуре.

Указанные особенности современного этапа развития Сети поставили проблему создания инструментов, способных «взять» на себя значительную часть функций основного когнитивного тракта человека. Таким образом проблема интеллектуализации Сети и ее инструментов начала уверенно передвигаться на передний план актуальности.

Среди многих аспектов интеллектуализации здесь мы выделяем язык. И действительно: среди огромного количества определений искусственного интеллекта фигурирует и такое: *«Искусственный интеллект — это форма индивидуализации технических систем, обладающая языковым статусом»*³, а значит интеллект вообще можно считать формой индивидуализации систем, обладающей языковым статусом.

Лингвистическое обеспечение Сети, на самом деле, сейчас играет роль ведущего фактора и основного интерфейса, обеспечивающего взаимодействие Человека с Сетью и Человека с Человеком через Сеть. Видимо, такая ситуация будет иметь место и в ближайшем будущем. Итак, приобретает все большую актуальность проблема объединения («интеграции») идей и технологий виртуализации, Больших данных и интеллектуализации (главным образом, через механизмы естественного языка).

Таким образом, лингвистическое обеспечение Сети и ее компонентов также приобрело новые черты актуальности. Но особенность лингвистических средств заключается в том, что языки существуют лишь в форме отдельных национальных языков, носители которых, собственно, и способны разрабатывать лингвистические компоненты программного обеспечения на должном уровне и приемлемого качества на основе использования современных моделей и инструментов лингвистического исследования. Основой такого рода разработок являются хорошо кодифицированные, аннотированные и репрезентативные модели и массивы лингвистических данных, представляющих все (в идеале) аспекты функционирования того или иного языка — как в когнитивном, так и в коммуникативном плане. Такими средствами служат лексикографические системы, обобщающие понятия словаря⁴, и лингвистические корпуса, представляющие большие цифровые массивы лингвистически квалифицированных и аннотированных текстов⁵.

³ В. А. Широков. Феноменология лексикографических систем. — К. — «Наукова думка». — 2004. С. 20.

⁴ В. А. Широков. Комп'ютерна лексикографія. — К.: Наук. думка, 2011. — 356 с.

⁵ В. А. Широков та ін. Корпусна лінгвістика: Моногр. / Широков В. А., Бугаков О. В., Грязнухіна Т. О., Костишин О. М., Кригін М. Ю., Сидорчук Н. М.; НАН України, Укр. мов.-інформ. фонд. — К.: Довіра, 2005. 350 С.

Следует отметить, что создание лексикографических систем и лингвистических корпусов, которые являются базой для дальнейшей разработки высококачественных средств лингвистического обеспечения, само по себе является достаточно сложной и чрезвычайно трудоемкой задачей, требующей привлечения значительного числа высококвалифицированных специалистов. При этом многие аспекты этих разработок на данном этапе не подлежат автоматизации и должны выполняться, так сказать, вручную. Тем более актуализируются проблемы создания формальных моделей, способных составить необходимый фундамент для формализации интеллектуальных свойств языка и создания эффективных моделей интеллектуализации. Разделяя убеждение об информационной природе естественного языка, мы считаем необходимым высказаться по вопросу об основных представлениях о природе самой информации, которые они считают релевантными рассматриваемому предмету.

Информация (от лат. *informatio* — осведомление, разъяснение, изложение) — в широком смысле является абстрактным понятием, имеющим множество значений в зависимости от контекста. В более узком смысле — это сведения (сообщения, данные) независимо от формы их представления.

Разумеется, что понятие «информация» по-разному толкуется и используется в различных научных дисциплинах. При этом в каждой дисциплине понятие «информация» связано с различными представлениями о мире и различными концептуальными парадигмами. Одновременно это понятие исключительно широко используется и в традиционном, обыденном смысле⁶ — это сведения, знания, сообщения о положении дел, которые человек воспринимает из окружающего мира с помощью органов чувств (зрения, слуха, вкуса, обоняния, осязания, ...), ну и, конечно же, приборов, посредством которых он производит наблюдения и мониторинг внешнего мира.

Информация может храниться, передаваться и обрабатываться в различных формах, в том числе и символической (знаковой) форме. Одна и та же информация может иметь различные формы представления, которые осуществляются с помощью определенных знаковых систем, конструируемых из неких базовых элементов («алфавитов»), имеющих правила для выполнения операций над ними.

Для обеспечения информационного процесса необходимы как минимум три элемента: *источник информации*, ее *получатель* (реципиент) и

⁶ Google на запрос «information» 24 февраля 2016 года, когда писался этот текст, выдал ни много, ни мало 8 560 000 000 ссылок (!) — число, превышающее число жителей Земли на тот момент.

устройство связи, обеспечивающее доставку информации от источника к реципиенту.

Таким образом, информация всегда есть информацией о «чём-то», то есть о соответствующем объекте, который и является её источником. В силу этого совершенно естественно предположить, что информация является общим, универсальным свойством всех вещей и процессов — не существует «сущностей», не обладающих информацией. По мнению такого авторитета как А. Колмогоров⁷ информация существует объективно, независимо от того, воспринимает её кто-нибудь или нет, хотя проявляется она только в процессе её восприятия. Информация отражает такие универсальные свойства вещей и процессов как *строение, структура, устройство, сложность, неоднородность, закономерности изменения, состояния объекта, его процессуальные характеристики* и т. д.

Отметим, что в наших работах⁸ изложен обзор различных представлений, подходов, концепций и теорий информации. Здесь мы не будем разбирать множество мнений и взглядов на данный предмет. Постараемся очертить некие общие естественные параметры и закономерности, регламентирующие информационные свойства систем, в связи с такими фундаментальными понятиями как материя, энергия и знания. Оказывается, что информация весьма тесно связана с данными концептами. Для этого нам понадобятся более или менее ясные, корректные и свободные от метафор определения отмеченных понятий. Это необходимо, так как в последнее время, в результате общей тенденции к свободному обращению с научными знаниями, данными понятиями оперируют весьма произвольно, особенно, люди, далекие от науки. Поэтому мы сначала приведем элементарные определения, которым будем в дальнейшем следовать. Сделаем несколько замечаний, так сказать, физического и метафизического характера, которые помогут войти в круг понятий, касающихся предмета рассмотрения.

Во-первых, интуитивно мы осознаем, что с понятием информации корреспондируется понятие знания, но, по сути, они совершенно различны.

Информация, как отмечалось выше, является объективной характеристикой объективных явлений и процессов. По нашему мнению, существует глубокая аналогия между определениями понятия

информации, с одной стороны, и энергии, также являющейся некоей объективной характеристикой вещей, с другой. И хотя лингвисты непосредственно не занимаются исследованием энергетических свойств сущего, привести данную аналогию мы полагаем весьма полезным, исходя из общеметодологических и общекультурных предпосылок. Отмеченную аналогию весьма легко проследить, обратившись к таблице «Сопоставление свойств понятий «ЭНЕРГИЯ» и «ИНФОРМАЦИЯ» (см. ниже).

Итак, мы наблюдаем определенный параллелизм в определении свойств понятий энергии и информации. В то же время подчеркнем, что данные понятия манифестируют совершенно различные свойства и характеристики объектов, что видно из п 2 приведенного в таблице сопоставления. Однако, несмотря на отмеченные сущностные различия, между ними существуют и весьма знаменательные связи. Например, замечательным свойством информации является то, что для её получения обязательно необходимо затратить какое-то количество энергии. Оказывается, что минимальные затраты энергии, необходимые для получения одного бита информации, вычисляются по формуле⁹:

$$kT \cdot \ln 2 \text{ эрг}, \quad (1)$$

где $k = 1,38 \cdot 10^{-16} \text{ эрг/}^\circ\text{K}$ — постоянная Больцмана, T — абсолютная температура системы, в которой происходит выработка информации. При 300 K° , то есть при нормальной температуре, эта величина равняется $4 \cdot 10^{-14} \text{ эрг}$, или приблизительно 10^{-2} электронвольт. Это очень малая величина.

Например, минимальное количество энергии, необходимой для получения 100 терабайт информации (такое примерно количество содержится в 50 миллиардах страниц печатного текста), равно примерно 1 эрг. Нам пока что неизвестны реальные информационные процессы, работающие с такой колоссальной эффективностью!

Существуют в определенном смысле обратные процессы, а именно такие, где, образно говоря, уже информацию можно преобразовать в энергию. Идея о возможности таких процессов возникла довольно давно в связи с обсуждением так называемого парадокса с «демоном Максвелла»¹⁰,

⁹ Волькенштейн М.В. Теория информации и эволюция. // «Кибернетика живого: биология и информация». — М.: Наука, 1984. — с. 45-53. Волькенштейн М. В. Энтропия и информация. — М.: Наука, 1986. — 190 с.

¹⁰ The thought experiment first appeared in a letter Maxwell wrote to Peter Guthrie Tait on 11 December 1867. It appeared again in a letter to John William Strutt in 1871, before it was presented to the public in Maxwell's 1872 book on thermodynamics titled Theory of Heat. (Theory of Heat. J.Clerk Maxwell, M.A. LL.D.Edin., F.R.SS. L.&E. LONDON: LONGMANS, GREEN, AND CO. 1871.; Cargill Gilston Knott (1911). "Quote from undated letter from Maxwell

⁷ Колмогоров А.Н. Теория передачи информации//Сессия академии наук СССР по научным проблемам автоматизации производства, 15-20 октября 1956 г.: Пленар. Заседания. —М.: Изд-во АН СССР, 1957. С.66-99.

⁸ Широков В. А. Інформаційна теорія лексикографічних систем. — К.: Довіра, 1998. Сс. 19-50. Широков В. А. та ін. Корпусна лінгвістика: Моногр. / Широков В. А., Бугаков О. В., Грязнухіна Т. О., Костишин О. М., Кригін М. Ю.; НАН України, Укр. мов.-інформ. фонд. — К.: Довіра, 2005. Сс. 105-120.

Сопоставление свойств понятий «ЭНЕРГИЯ» и «ИНФОРМАЦИЯ»

ЭНЕРГИЯ	ИНФОРМАЦИЯ
1. Энергия — это универсальная количественная характеристика физических систем. Не существует физических систем, которые бы не характеризовались энергией.	1. Информация — это универсальная количественная характеристика любых систем. Не существует систем, которые бы не характеризовались информацией.
2. Указанная величина характеризует <i>интенсивность</i> процессов, происходящих в физических системах.	2. Указанная величина характеризует <i>сложность</i> процессов, происходящих в системах, и сложность (неоднородность, структурированность) самих систем.
3. Величина энергии (как измеряемая на эксперименте, так вычисляемая теоретически) изображается числами.	3. Величина информации (как измеряемая, так вычисляемая теоретически) изображается числами.
4. Размерность энергии: [масса] [длина] ² [время] ⁻² .	4. Размерность информации: [бит] или [энергия].
5. В различных по своей природе системах энергия имеет разные проявления и характеристики. Они отражаются в способах экспериментального наблюдения и измерения энергетических эффектов, а также в способах (моделях) их теоретического описания.	5. В разных по своей природе системах информация имеет различные проявления и характер. Они отражаются в способах экспериментального наблюдения и измерения информационных эффектов, а также в способах (моделях) их теоретического описания.
6. Различные проявления энергии для разных систем и уровней рассмотрения принято называть формами энергии (механической, электрической, магнитной и т.д.). Фундаментальное свойство — закон сохранения: в замкнутой системе все процессы происходят так, что энергия, превращаясь из одной формы в другую, остается постоянной величиной.	6. Различные проявления информации для различных систем и уровней рассмотрения принято называть формами информации. Информация существует в форме данных, текстов, знаний, моделей и т.п. Информационные процессы сопровождаются преобразованиями информации из одной формы в другую.
7. Для различных уровней материи характерны свои, специфические способы энергетического описания, представляемые с помощью соответствующих теоретических моделей и математических формализмов.	7. Для разных типов систем характерны свои, специфические способы информационного описания, представляемые с помощью соответствующих теоретических моделей и математических формализмов (например, модели Хартли, Шеннона, Колмогорова, ...).

который представляет пример механизма, в некотором смысле преобразующего информацию в энергию, то есть использующего в качестве «горючего» информацию, что впервые было отмечено Лео Сциллардом¹¹. Пример такого процесса описан в книге¹². Данный эффект дает основания полагать, что аналогичные процессы происходят и в социотехнических системах; они подробно описаны в работах¹³.

Таким образом, мы констатируем, что информационные процессы не составляют замкнутой системы и очень быстро выводят в другую область описания свойств вещей — энергетическую. Это побуждает специалистов, занимающихся информационными свойствами языка, постепенно

осознавать недостаточность одного лишь информационного подхода даже в своей, по преимуществу информационной области. И — по мере всё более глубокого проникновения в свойства языка — представления, факты и методы других наук, по нашему убеждению, будут применяться в языковедении всё чаще. В этом проявляется смысл трансдисциплинарности — основной определяющей черты современного развития науки.

Отметим, что далеко не любая информация может быть преобразована в полезный ресурс. В частном случае этот процесс, образно выражаясь, способен реализовать даже один информационный демон Максвелла. В реальных же ситуациях для этой цели создаются целые комплексы организаций и институций (следуя Селфриджу¹⁴, употребим для них название «пандемониумы»), обеспечивающие производство, многократные преобразования информации и ее транзакции до тех пор, пока она не поступит в социотехническую систему в нужный момент, да еще и в нужной форме, адаптированной к восприятию упомянутой системой. Именно процессы продуцирования и

to Tait». Life and Scientific Work of Peter Guthrie Tait. Cambridge University Press. p. 215.)

¹¹ Szillard L. Über die Entropievermindung in einem Thermodynamischen System bei Eingriffen intelligenter Wesen. — Zs. f. Phys., 1929. — 53, 11–12, Heft. — P. 840–856.

¹² Стратонович Р.Л. Теория информации. — М.: Сов. радио, 1975. — 423 с., гл. 12.

¹³ Широков В.А. Інформаційно-енергетичні трансформації та інформаційне суспільство. //Українсько-польський науково-практичний журнал «Наука, інновація, інформація». — Київ, 1996. — № 1. — С. 48–66. В.А.Широков. Інформаційна теорія лексикографічних систем, — К.: Довіра, 1998, 331 с. В.А.Широков. Феноменологія лексикографічних систем. — К.: Наукова думка, 2004, 326 с.

¹⁴ Selfridge O.G. Pandemonium: a paradigm for learning. In: Mechanisation of thought processes. London. HMSO., 1959. — P. 511–531.

целенаправленного многократного преобразования и транспортирования информации и придают ей новое качество, позволяющее квалифицировать ее как *знание*. Следовательно, теперь мы можем дать такое рабочее определение понятия знания: *знание — это информация, форма которой является носителем трансформаций, адаптирующих ее рецепцию социальной системой*.

Подчеркнем, что это определение не претендует на абсолют и является рабочим, даже техническим. Заметим, что в нем совершенно отсутствуют конструкции типа «правильное отображение закономерностей мира в голове человека», характерные для философских определений знания. Данное нами определение, как мы полагаем, является максимально объективированным, поскольку оно содержит объективное понятие информации и достаточно ясное представление о роли ее трансформаций в социальной системе. Отсутствие референций к «правильности» и «закономерностям мира» мы считаем совершенно оправданным, поскольку, как известно, существуют ложные знания, а также знания, не отражающие закономерностей мира. Несмотря на такую лапидарность нашего определения, из него следуют вполне конкретные выводы.

Во-первых, если *информация*, как таковая, представляет определенные *объективные* свойства вещей, то *знание* несет в себе потенциал *субъективного*. Ведь перед тем как попасть в производственно-экономическую систему информация должна стать фактом сознания — сначала индивидуального, а затем и коллективного, по крайней мере, группового. При этом она, понятно, подвергается целому ряду трансформаций в соответствии со спецификой функционирования сознания. А конструкция человеческого интеллектуального аппарата такова, что в нем происходят многочисленные превращения и взаимодействия между ментальными и языковыми структурами.

Исследования великих лингвистов и психологов, среди которых отметим Вильгельма фон Гумбольдта, В. Бехтерева, Л. Выготского, Н. Хомского, Р. Шенка и многих других, убедительно подтвердили, что любой ментальный процесс у человека имеет свою рефлексию в языковой сфере. Так что корректно говорить о едином *мыслеречевом* процессе. Подчеркнем также и ту важную роль, которую играет специализация языковой подсистемы как основного коммуникативного инструмента в человеческом обществе.

Из сказанного следует еще одно, довольно простое определение знания как *информации, вербализованной согласно законам естественной языковой системы*. Ведь в языковой форме скомпрессованы такие важные аспекты человеческого бытия как

культурный код, научная и языковая картины мира, языковое сознание и подсознание (коллективное и индивидуальное) и т. д.

И вот, научная или контекстно-предметная картина мира, которая обычно является первичной в информации, превращаясь в общественно значимое знание, активно взаимодействует с языковой системой и «набирается» от нее свойств, присущих этой системе. Вообще, языковая и контекстно-предметная компоненты составляют два вполне различных аспекта функционирования знания. Причем, именно языковая форма превалирует на достаточно большом отрезке функционирования знания, вплоть до акта его непосредственного потребления социумом. Отметим, что языковая форма, приобретаемая знанием, сама по себе довольно удобна и гибка.

Во-первых, она универсальна: мы имеем убеждение, что практически любая информация может быть вербализована — это нашло отражение в формулировке первой из наших Презумпций.

Во-вторых, каждый человек вполне квалифицированно оперирует этой формой, усваивая ее с раннего детства.

В-третьих, языковая форма линейна, вследствие чего она легко поддается различным информационным операциям — кодированию, преобразованию, хранению, транспортировке и т.п.

Но обратной стороной этого удобства и гибкости оказывается то, что в естественной языковой форме информации более или менее эксплицитно представлены лишь элементы *формы* языковой системы, в то время как *онтологическая* или *семантическая* составляющие представлены только имплицитно. Из этого логически следует, что «экстракция» семантической информации из вербализованного знания, то есть естественной языкового текста неминуемо ведет к применению принципов языковой системы, так сказать, в «обратном» направлении. А именно, если символически изобразить процесс вербализации информации, как:

$$L: I \rightarrow T, \quad (2)$$

где I — информация, L — «преобразователь» информации в языковой объект, T — текст, то добывание семантической информации будет выглядеть таким образом:

$$L^{-1}: T \rightarrow I_{\text{сем}}. \quad (3)$$

Последний процесс мы и квалифицируем как экстракцию знаний из текста.

Понятно, что написать формулы (2) — (3) намного легче, чем организовать соответствующие процессы, представленные этими формулами. В последнее время наблюдается значительное возрастание потока научных публикаций на данную

тему. Немало из них посвящено методам автоматизированного построения моделей знаний по тексту на естественном языке. Однако особо выдающихся успехов в этих направлениях пока что нет.

Во-первых, следует отметить, что все методики на самом деле очень зависимы от конструкции языковой системы конкретного национального языка, на котором представлена входная информация. Поэтому для перехода к интеллектуальному анализу текста для каждого языка необходимо выполнить определенные этапы анализа языковых структур, неявно содержащихся в конструкциях операторов L и L^{-1} . Такими этапами, в частности, являются: анализ знаковых систем, грамматический анализ языковых структур, контекстный, синтаксический, семантический, прагматический, статистический анализ и т. д. Создание эффективных автоматических процедур, ориентированных на выполнение отмеченных разновидностей анализа является довольно сложной научно-технической задачей.

На данный момент в мире уже существуют немало программных средств, способных в определенной мере выполнять такие операции. Наиболее распространенными среди них являются поисковые системы, в частности, интернетовские, способные выполнять целый ряд функций естественного языка. Тем не менее, полного удовлетворения пользователям эти средства пока что не приносят. Наоборот, все хорошо знают, насколько «шумящими» являются поисковые средства Интернета. Наш опыт свидетельствует, что такие средства и, соответственно, технологии должны носить комплексный характер и объединять целый ряд концептуальных парадигм, ведь вряд ли стоит надеяться, что одна схема способна охватить все многообразие когнитивных ситуаций, возникающих при интеллектуальной обработке текстов на предмет экстракции из них знаний.

Информационный подход к языку требует привлечения максимально формализованных представлений о природе информации в её отношении к природе языка. В ряду таких представлений важное место занимает теории, позволяющие квалифицировать и оценивать информацию с количественной стороны. В настоящее время наиболее распространёнными являются меры количества информации, введение которых связывают с именами Хартли, Шеннона и Колмогорова. Обзор данных представлений и необходимые литературные ссылки приведены в наших книгах¹⁵.

¹⁵ В.А.Широков. Інформаційна теорія лексикографічних систем, — К.: Довіра, 1998, 331. . В.А.Широков. Феноменологія лексикографічних систем. — К.: Наукова думка, 2004, 326 с.

Широков В. А. та ін. Корпусна лінгвістика: Моногр. / Широков В. А., Бугаков О. В., Грязнухіна Т. О., Костишин О. М.,

2. Проявления системной триады «структура — субстанция — субъект» в научных описаниях естественного языка

2.1. Системность и словоизменение флективных языков

Для определения понятия системы¹⁶ в первой части статьи было предположено символическое равенство:

$$C = C + C + C, \quad (4)$$

где «С» левой части обозначает понятие «система», а правая часть демонстрирует наличие и взаимодействие основных образующих компонент этого понятия, а именно «структуру», «субстанцию» и «субъект»¹⁷.

Системная триада «структура — субстанция — субъект» ярко проявляется практически во всех областях теоретического описания языка. Рассмотрим, например, такой характерный лексический феномен флективных языков как словоизменение и проанализируем его системные свойства на примере украинского языка, учитывая установленную нами системную триаду (1).

1. Структура:

Каждое слово в украинском языке (да и в иных флективных языках) имеет структуру:

$$x = \rho(x) * \omega(x),$$

где символом $\rho(x)$ обозначена квазиоснова слова x , то есть часть лексемы, которая остается неизменной в процессе словоизменения x (эта часть одинакова для всех словоизменительных форм лексемы); $\omega(x)$ — квазифлексия, то есть часть лексемы x , которая подвергается изменениям в процессе построения парадигмы¹⁸. Символом «*» обозначена конкатенация. Введем дополнительные обозначения: $[x]$ — полная парадигма слова x ; $[x] = \rho(x) * [\omega(x)]$; $[\omega(x)]$ — набор квазифлексий, входящих в состав парадигмы $[x]$. Структура и субстанциальное наполнение $[\omega(x)]$ определяется словоизменительной классификацией. Приведем пример парадигмы для слова «інстинкт». Структура данной парадигмы приведена в таблице.

Кригін М. Ю.; НАН України, Укр. мов.-інформ. фонд. — К.: Довіра, 2005. — 472 с. Широков В. А. Елементи лексикографії: Моногр. / Укр. мов.-інформ. фонд НАН України. — К.: Довіра, 2005. — 304 с.

¹⁶ Автор вполне осознаёт, что точного определения понятия системы быть не может в силу фундаментальности, первичности данного понятия, поэтому здесь слово «определение» в применении к понятию системы употребляется в несколько «пиквикском» смысле.

¹⁷ Автор также осознаёт те трудности, которые встретятся на пути определения понятий *структуры, субстанции и субъекта*. Поэтому данным вопросом (которому, впрочем, посвящена колоссальная библиография) в настоящей работе мы заниматься не будем, апеллируя к интуиции читателя.

¹⁸ В этом разделе под словом парадигма мы подразумеваем словоизменительную парадигму лексемы.

Падеж	Число	[x]	$\rho(x)$	$[\omega(x)]$
Називний	Однина	Інстинкт	інстинкт	$\emptyset$
Родовий	Однина	інстинкту	інстинкт	y
Давальний	Однина	інстинкту інстинктові	інстинкт	y ові
Знахідний	Однина	Інстинкт	інстинкт	$\emptyset$
Орудний	Однина	інстинктом	інстинкт	ом
Місцевий	Однина	інстинкті	інстинкт	i
Кличний	Однина	інстинкте*	інстинкт	e
Називний	Множина	інстинкти	інстинкт	и
Родовий	Множина	інстинктів	інстинкт	ів
Давальний	Множина	інстинктам	інстинкт	ам
Знахідний	Множина	інстинкти	інстинкт	и
Орудний	Множина	інстинктами	інстинкт	ами
Місцевий	Множина	інстинктах	інстинкт	ах
Кличний	Множина	інстинкти*	інстинкт	и

2. Субстанция

Субстанциальное наполнение данной парадигмы задается набором квази-флексий $[\omega(x)] = \{\emptyset; y; (y, ovi); \emptyset; om; i; e; и; iv; am; и; ами; ax; и\}$. Он представляет *субстанцию* парадигмы [інстинкт], определяющую словоизменительный класс $K(x)$, которому принадлежит лексема «інстинкт» и которому приписываются все парадигматические атрибуции согласно правилам украинского словоизменения имен существительных.

3. Субъект

К данному члену системной триады относятся:

- алгоритмы грамматической (морфологической) идентификации;
- алгоритмы построения разложения

$$x = \rho(x) * \omega(x)$$

и построения парадигмы:

$$[x] = \rho(x) * [\omega(x)];$$

- алгоритмы словоизменительной классификации, то есть построения соответствия:

$$[\omega(x)] \Leftrightarrow K(x);$$

– алгоритмы лемматизации, то есть правила реконструкции грамматического значения по виду соответствующей текстовой формы. Эта задача может иметь неоднозначное решение даже в пределах фиксированной парадигмы вследствие явления грамматической омонимии. Например, форма «інстинкта» имеет грамматические значения: «Родовий; Однина» и «Давальний; Однина». Эта омонимия может быть снята только при наличии достаточно широкого контекста. Только тогда «субъект» языкового процесса может однозначно идентифицировать грамматическое состояние соответствующей формы.

Заметим, что элементы системной триады «структура – субстанция – субъект» в Украинском языково-информационном фонде реализованы в

Виртуальных грамматических лексикографических лабораториях (ВГЛЛ), а именно: для украинского языка (реестр более 258 тыс. единиц), русского языка (реестр более 180 тыс. единиц), немецкого языка (реестр более 60 тыс. единиц) и агглютинативного турецкого языка (существительное; реестр более 30 тыс. единиц). Созданы также ВГЛЛ для испанского, французского, польского языков. Указанные системы являются инструментальными и могут быть использованы для грамматических исследований в среде соответствующих языков.

2.2. О других системных отношениях

Автор убежден, что проявления системности в форме триады «структура – субстанция – субъект» имеют статус лингвистических универсалий и в том или ином виде характерны для грамматики любого языка. Однако эти проявления характерны и для разнообразных семантических свойств и отношений, развивающихся в языковой системе. Этому вопросу мы предполагаем посвятить отдельную заметку. В ней будут изложены системные элементы семантики, абстрагируемые из структур больших толковых словарей, рассматриваемых как лексикографические системы. Всего будет рассмотрено четыре системных отношения, а именно, следующие: 1. Замкнутость и полнота лексикографических систем. Автоморфизмы, гиперцепи и гиперциклы на лексикографических системах. 2. Строение эквисемантических рядов лексикографических систем. 3. Исследование формул квазисемантики и ядра лексической системы. 4. Применение теории лингвистических состояний к единицам лексической системы и коллокациям, что позволяет построить унифицированную грамматическую теорию для этих, весьма различающихся между собой классов языковых единиц.

Выводы

В заключение приведем без комментариев десять достаточно общих утверждений, носящих общесистемный характер, и поэтому имеющих весьма неопределенную область применения. Вследствие такой неконтролируемой общности к данным утверждениям неприменим модус рассуждения, обычно трактуемый как «доказательство». Поэтому мы, квалифицируя эти утверждения как «теоремы», прибавляем к ним эпитет «системные», полагая, что он подчеркивает отмеченный модус. Можно считать, что приведенные утверждения отражают жизненный и научный опыт автора, не претендующего, впрочем, на личное авторство относительно каждого из данных утверждений. Итак:

1-я системная теорема.

В каждой системе развиваются противоположные тенденции.

2-я системная теорема

Сложное содержит в себе еще сложнее.

3-я системная теорема.

Глобальное возмущение сложной системы ведет к последствиям, противоположным тем, которые ожидалось.

4-я системная теорема.

Отрицанием простой истины является ложь, отрицанием сложной истины является другая сложная истина.

5-я системная теорема.

Локальная оптимизация в системе может привести к стратегическим просчетам.

6-я системная теорема.

В каждой системе переход от одного состояния порядка к другому состоянию порядка возможен только через некое состояние хаоса.

7-я системная теорема.

Переход из состояния хаоса в состояние порядка происходит через процесс, который напоминает

лексикографический эффект в информационных системах¹⁹.

8-я системная теорема.

Система только тогда имеет шанс достичь совершенства, когда осознает необходимость составить полный словарь о самой себе. (Критерий совершенства)

9-я системная теорема.

Совершенство недостижимо

10-я системная теорема.

В интеллектуальной семантической сети при превышении среднего уровня интеллекта ее узлов некоторой критической величины уровень интеллекта целой сети начинает падать.

(Принцип А.С. Грибоедова: «Горе от ума»).

¹⁹ Понятие лексикографического эффекта в информационных системах впервые было введено и обосновано нами в книге В.А.Широков. Інформаційна теорія лексикографічних систем, – К.: Довіра, 1998, 331.

УДК 81'322.2'33

**О.В. Лазаренко¹, Д.И. Панченко², Е.Ю. Айвас³**¹ХГУ «НУА», г. Харьков, Украина, lazolvlad@gmail.com²ХГУ «НУА», г. Харьков, Украина, panchenko.di2013@gmail.com³ХГУ «НУА», г. Харьков, Украина, b.u.elena@mail.ru

МОДЕЛИРОВАНИЕ БАЗОВЫХ СОСТАВЛЯЮЩИХ ПРОЦЕССА ПОНИМАНИЯ ТЕКСТА В СИСТЕМЕ АВТОМАТИЧЕСКОГО РЕФЕРИРОВАНИЯ

В статье рассматривается процедура смыслового анализа текста с использованием концептуальных инвариантов текста, текстовых баз, описывающих главные смысловые аспекты текста, и прообразов рефератов для синтеза автоматических рефератов на уровне глубинной семантики. Предложенная процедура позволяет обеспечить универсализацию алгоритма смыслового анализа текстов различной тематики за счет создания ситуационных моделей при разработке системы автоматического реферирования.

АВТОМАТИЧЕСКОЕ РЕФЕРИРОВАНИЕ, СИТУАЦИОННАЯ МОДЕЛЬ, ТЕКСТОВАЯ БАЗА, КОНЦЕПТУАЛЬНЫЙ ИНВАРИАНТ ТЕКСТА, ПРООБРАЗ РЕФЕРАТА

Введение

Изучение и моделирование процесса понимания человеком текста, проводимые в различных областях современных научных исследований таких как исследование механизмов работы мозга (Дж. Хокинс [1]), разработка стратегий понимания дискурса (А. ван Дейк [2]), моделирование процесса реферирования (О.В. Лазаренко [3]) и др. прямо и косвенно подтверждают тот факт, что человек в процессе распознавания объектов и ситуаций использует наиболее важные их характеристики, хранящиеся в его памяти в виде инвариантных форм.

В своих исследованиях при разработке системы автоматического реферирования мы вышли на понимание этих механизмов через изучение особенностей смысловой структуры реферата в сравнении с первичным текстом и его заголовком [3,4,5,6 и др.]. В связи с чрезвычайной сложностью процесса реферирования, мы на начальном этапе наших исследований сознательно допустили теоретическую и эмпирическую неполноту, ограничившись изучением индикативных рефератов научных текстов. Оттолкнувшись, таким образом, от понимания особенностей и закономерностей наиболее структурированного объекта (индикативного реферата) и последовательно расширяя и углубляя область исследования, мы пришли к пониманию определенных механизмов сжатия смысла в цепочке «текст-реферат-заголовок». И оказалось, что в основе этих механизмов лежит использование инвариантных форм представления информации.

В ходе разработки процедуры семантического анализа текста с целью сжатия его смысла была построена семантико-контекстную модель реферирования, включающая модель заголовка и текстовую базу. Заголовок рассматривается нами как смысловой инвариант текста, а текстовая база как «информационное ядро» текста, содержащее

информацию о ситуации, описанной в тексте. В текстовую базу входят предложения, отражающие основные смысловые аспекты исходного текста.

В своих дальнейших исследованиях процесса понимания текста [3] мы пришли к выводу о том, что процедура смыслового анализа текста с построением текстовых баз при выборе предложений, описывающих главные смысловые аспекты текста, позволяет обеспечить универсализацию алгоритма смыслового анализа текстов различной тематики и различных предметных областей. Инструментом такой универсализации стала ситуационная модель. В разрабатываемой нами системе ситуационная модель формируется в виде накопителя текстовых баз определенной тематики, автоматически извлекаемых из текста в процессе его смыслового анализа в соответствии с разработанным алгоритмом извлечения основных смысловых аспектов текста.

Предложенный подход к смысловому анализу текста позволяет обеспечить более качественный результат автоматического реферирования за счет:

1) выделения макроструктуры текста в виде главных смысловых аспектов и построения из них текстовой базы, являющейся «смысловым ядром» текста;

2) формирования в автоматическом режиме ситуационных моделей в виде накопителей текстовых баз с целью выявления инвариантных репрезентаций ситуаций;

3) использование инвариантных репрезентаций ситуаций для извлечения из текста информации, необходимой при построении реферата любого вида.

Основной целью наших исследований на данном этапе является разработка процедуры построения ситуационных моделей и инвариантной репрезентации ситуации, обеспечивающих анализ глубинной семантики текста.

1. Формирование прообраза реферата с использованием текстовой базы

На очередном этапе исследования процесса понимания текста с целью выделения необходимых смысловых аспектов для построения реферата мы пришли к выводу о том, что использование текстовых баз для выбора предложений, описывающих главные смысловые аспекты текста, позволяет более точно определить предложения-претенденты для формирования прообраза реферата. При разработке методики построения текстовой базы, представляющей «информационное ядро» текста, мы опирались на использование заголовка как смыслового инварианта текста и слов-указателей на необходимые для реферата смысловые аспекты: объект, результат, метод, область исследования, цель исследования [5]. В соответствии с этой методикой для каждого текста создавались две текстовые базы.

В первом случае в текстовую базу были выделены все предложения в тексте, содержащие слова из заголовка.

Во втором в текстовую базу выбирались предложения, содержащие слова из заголовка и слов-указатели на необходимые для реферата смысловые аспекты.

Анализ двух типов текстовых баз показал, с одной стороны, их схожесть в значительной степени, а с другой, более точный выбор предложений при использовании дополнительно слов-указателей.

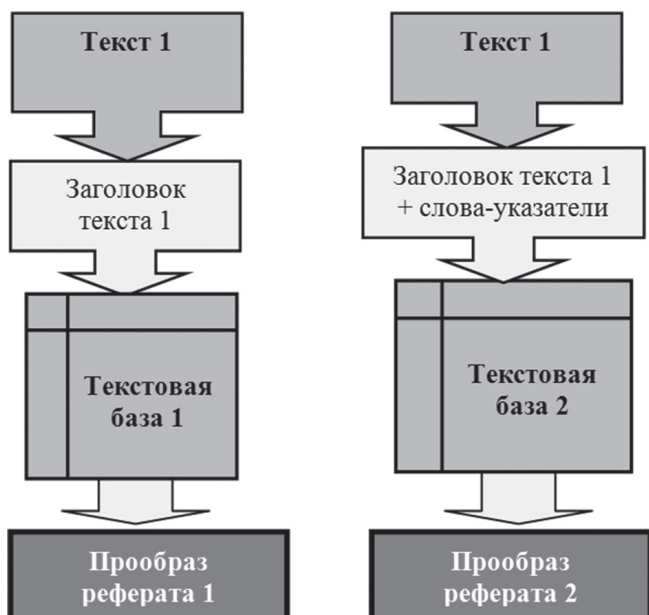


Рис. 1. Процедуры отбора предложений для создания прообраза реферата

По результатам анализа полученных текстовых баз для построения прообраза реферата были выделены первые три и последние два предложения из текстовой базы.

Пример 1.

Заголовок: Интенсификаторы в современном английском языке

Предложения из ТБ:

1. В данной статье в фокусе внимания находятся семантические и структурные свойства интенсификаторов современного английского языка, используемых в современных художественных, публицистических и газетных текстах.

2. Под интенсификаторами мы понимаем разноразличные единицы языка, функционирующие как усилители признака в широком смысле.

3. Иначе говоря, интенсификатор — это слово или фраза, которая добавляет силу или эмфазу высказыванию.

...

33. В современной литературе, отражающей спонтанную разговорную речь, ощутимо присутствуют экспрессивные интенсификаторы-дисфемизмы, среди которых преобладают так называемые “four-letter-words”.

34. Избыточность представления интенсификации признака является стилистическим маркером экспрессивности, показателем особенностей индивидуального стиля автора или портретной / психологической характеристики героя.

Анализ выделенных предложений:

Все предложения достаточно информативны и описывают общий смысл статьи.

Пример 2.

Заголовок: Роль концептуальной метафоры “death” в экспликации концепта “vampire”

Предложения из ТБ

1. Согласно мнению американских ученых Дж. Лакоффа и М. Джонсона, метафоры являются концептуальными, поскольку одновременно существуют в двух концептах, сферах и тем самым показывают взаимосвязь между ними [Lakoff, Johnson, 1980, p. 37].

2. Понятие концепта является, на сегодняшний день, неоднозначным, однако, мы придерживаемся мнения Ю.С. Степанова, который предложил разделение концепта на слои: актуальный, этимологический и пассивный [Степанов, 2001, с.45].

3. Пассивный слой — особенности, типичные для определенных носителей данного концепта.

...

15. Таким образом, можно сделать вывод, что использование концептуальной метафоры “DEATH” для экспликации концепта “VAMPIRE” является очень распространенным средством, особенно в традиционном готическом романе.

16. Использование данной концептуальной метафоры увеличивает степень влияния на читателя и делает атмосферу романа еще более устрашающей.

Анализ выделенных предложений:

Выделенные предложения передают общий смысл статьи, лишним является третье предложение, т.к. оно содержит описание одного конкретного слоя концепта, что относится к деталям, не отражающимся в индикативных рефератах.

Пример 3.

Заголовок: Фразеологизмы с кулинарным компонентом в контексте гастрономического кода немецкой национальной культуры

Предложения из ТБ:

1. На протяжении столетий язык передаёт векования, заблуждения, национально-культурные установки, фиксируя таким образом национально-специфическое видение окружающего мира.

2. Собранные в этих словарях фразеологические единицы (далее ФЕ) включают в себя устойчивые словесные комплексы эпохи раннего средневековья, они снабжены примерами, авторскими комментариями, пометами об их распространённости в конкретной местности Германии, нередко приводятся эквиваленты из других языков.

3. Сравнительный анализ словарей XIX века и современных лексикографических источников позволяет сделать вывод о том, актуальна ли та или иная идиома в современном немецком языке.

...

36. При этом некоторое число аналогичных устойчивых словесных комплексов в других европейских языках свидетельствует о сходном отношении к пище, сложившемся, в частности, под влиянием христианской культуры.

37. Путём привлечения экстралингвистической информации — сведений о повседневной жизни, социально-культурных установках, религиозных представлениях народа — и соотнесения их с лингвистической формой выражения представляется возможным описать пока ещё недостаточно изученный сегмент гастрономического кода национальной культуры.

Анализ выделенных предложений:

Первые три предложения содержат общую информацию, в которой отсутствует указание на то, чему посвящена статья. Наиболее информативным является последнее предложение.

Общие выводы:

1. Некоторые прообразы реферата получились достаточно информативными (пример 1).

2. Первые три предложения текстовой базы могут содержать общую вводную информацию, которая не является полезной для реферата (пример 2).

3. Некоторые выбранные предложения содержат описание деталей, для понимания которых необходима дополнительная информация (пример 3).

Таким образом, прообразы рефератов не всегда получаются из выбранных первых трех и последних двух предложений из текстовой базы. Это подтвердило нашу гипотезу о том, что при выборе предложений для прообраза реферата, нужно ориентироваться на предложения в ТБ, выбранные из первых четырех абзацев текста, которые практически со 100% вероятностью позволяют найти в них указание на объект, область и цель исследования. Вместе с тем, последние два предложения из ТБ всегда представляют описание результата.

2. Оптимизация процедуры выбора предложений из текстовой базы для прообразов реферата

Очень часто тексты научных статей начинаются с общего описания проблемы, изучения и полученных к данному моменту результатов предшествовавших исследований. Иными словами, с описания истории вопроса. Естественно, что постановка задачи, цель и метод исследования при этом приводятся после вводной части. Вместе с тем во многих статьях такая информация либо полностью отсутствует, либо сведена к нескольким предложениям. Чтобы найти интересующие нас смысловые аспекты при любом варианте написания статьи, в нашем алгоритме поиска предложений в текстовой базе для прообраза реферата мы на начальном этапе отбираем предложения, выделенные в текстовую базу из первых четырех абзацев исходного текста. Экспериментальная проверка подтвердила достаточность этих предложений для выбора объекта (всегда), а также цели, метода и области исследования, если в статье есть на них указание.

Ниже приведен сравнительный анализ двух подходов для выбора предложений из текстовых баз для составления прообраза реферата:

$$R_{pr} \in \text{TextBase}_1 (P_1, P_2, P_3, P_{n-1}, P_n)$$

и

$$R_{pr} \in \text{TextBase}_2 (P_k, P_{k+1}, \dots, P_{n-1}, P_n),$$

где R_{pr} — прообраз реферата, TextBase — текстовая база, P_1, P_2, P_3 — первые три предложения из текстовой базы первого типа, P_k — первое из предложений, выбранных из первых четырех абзацев текста для текстовой базы второго типа, P_n — последнее предложение в текстовой базе любого типа.

Статья: Роль концептуальной метафоры “death” в экспликации концепта “vampire”

$$R_{pr} \in \text{TextBase}_1 (P_1, P_2, P_3, P_{n-1}, P_n)$$

$$P_1 - \{$$

$$P_2 - \{$$

$$P_3 - \{$$

$P_{n-1} - \{$ Таким образом, можно сделать вывод, что использование концептуальной метафоры “DEATH” для экспликации концепта “VAMPIRE”

является очень распространенным средством, особенно в традиционном готическом романе}

P_n - {Использование данной концептуальной метафоры увеличивает степень влияния на читателя и делает атмосферу романа еще более устрашающей}

$Rpr \in \text{TextBase2}(P_k, P_{k+1}, \dots, P_{n-1}, P_n)$

P_k - {В нашей статье внимание направлено на актуальный слой концепта "VAMPIRE", а именно на использование концептуальной метафоры "DEATH" как репрезентанта данного концепта}

P_{k+1} - {Языковое воплощение лингвокультурологического концепта "VAMPIRE" имеет четкий ареал в литературе — это готические романы ужасов, романы о вампирах}

P_{n-1} - {Таким образом, можно сделать вывод, что использование концептуальной метафоры "DEATH" для экспликации концепта "VAMPIRE" является очень распространенным средством, особенно в традиционном готическом романе}

P_n - {Использование данной концептуальной метафоры увеличивает степень влияния на читателя и делает атмосферу романа еще более устрашающей}

Приведенное сравнение демонстрирует тот факт, что использование второго типа текстовых баз позволяет создать более полный и точный прообраз реферата.

Как отмечалось в работе [3], в ходе наших исследований мы пришли к выводу о целесообразности поэтапного приближения к выбору необходимых предложений из текста. Поскольку предложений, указывающих на определенный смысловой аспект, может быть несколько, следует выделить их из текста в текстовую базу, которая представляет собой расширенное информационное ядро текста. А затем на основе анализа заголовка, слов-указателей на смысловые аспекты и предложений из текстовой базы выбрать необходимую информацию для прообраза реферата.

Проведенные исследования подтвердили эффективность такого подхода и соответствие полученных результатов концепции понимания, предложенной в работах голландского лингвиста ван Дейка [2], согласно которой для описания глобального содержания текста необходимо построение схемы, обеспечивающей «быстрый анализ поверхностных структур и выстраивание относительно простой и жесткой семантической конфигурации». В нашем случае это текстовые базы и прообразы рефератов. Такие структуры текста представляют собой обобщенное описание основного содержания дискурса, которое читатель строит в процессе понимания, и являются фактически рефератом

или резюме. А это то, что и является конечной целью наших исследований.

В результате последних исследований мы вплотную подошли к моделированию процесса реферирования на уровне глубинной семантики текста. Следующим шагом на этом пути будет автоматическое построение ситуационной модели для создания инвариантных репрезентаций ситуаций, описываемых в текстах, и представляющих собой набор наиболее важных признаков, выделенных на основе относительных характеристик ситуации, в которых возможны существенные упущения в сравнении с конкретной ситуацией, описываемой в конкретном тексте.

Выводы

В статье рассмотрена процедура построения прообразов рефератов на основе использования текстовых баз, позволяющая обеспечить более качественный результат автоматического реферирования.

Подтвержден тот факт, что для создания связанного прообраза реферата необходимо использовать предложения из текстовой базы, выделенные из первых четырех абзацев текста, и последние два предложения текстовой базы.

Построенные таким образом прообразы рефератов являются хорошим базисом для создания инвариантных репрезентаций ситуаций в виде набора наиболее важных признаков обобщающего порядка в сравнении с конкретной ситуацией, описываемой в тексте.

Список литературы:

1. Хокинс Дж., Блейкли С. Об интеллекте / Дж., Хокинс, Блейкли С. — М.: Издательский дом "Вильямс", 2007. — 240 с.
2. Дейк ван Т. А. Стратегии понимания связанного текста / Т. А. ван Дейк, В. Кинч // Новое в зарубежной лингвистике. — Вып. 23: Когнитивные аспекты языка. — М., 1988. — С. 153–211.
3. Лазаренко О. В. Моделирование процесса понимания текста с использованием инвариантной репрезентации ситуаций в системе авто-реферирования / О. В. Лазаренко // Бiонiка iнтелекту: навч.-техн. журнал. 2014. — Вып. 2(83). С. 15–19.
4. Лазаренко О. В. Моделювання процесу узагальнення в системі автоматичного реферування / О. В. Лазаренко, А. А. Яковенко. — Х.: Изд-во НУА, 2007. — 136 с.
5. Лазаренко О. В. Моделювання семантичних зв'язків «Текст-Реферат» в системах автоматичного реферування / О. В. Лазаренко, Д. І. Панченко. — Х.: Изд-во НУА, 2014. — 176 с.
6. Буряк Е. Ю., Лазаренко О. В., Панченко Д. И. Разработка алгоритма смыслового анализа текста для синтеза реферата в системе автоматического реферирования / Е. Ю. Буряк, О. В. Лазаренко, Д. И. Панченко / Бионика интеллекта: науч.-техн. журнал — Харьков: ХНУРЭ, 2015. — Вып. 2 (85) — С. 127 — 130.

Поступила в редколлегию 9.11.2016

УДК 004.853

**Л.Э. Чалая¹, С.Г. Удовенко², Е.В. Кушвид¹**¹ ХНУРЭ, г. Харьков, Украина, larysa.chala@nure.ua¹ ХНУРЭ, г. Харьков, Украина, eugenii.kushvid@nure.ua² ХНЭУ, г. Харьков, Украина, serhii.udovenko@nure.ua

МЕТОД ДВУХЭТАПНОЙ КЛАССИФИКАЦИИ ЭЛЕКТРОННЫХ ТЕКСТОВ

Предложен метод классификации электронных текстов, основанный на комбинированном применении полиномиальной модели классификатора Байеса, лингвистических дескрипторов и модифицированного метода Гинзбурга. Метод содержит два этапа: этап предварительной фильтрации псевдоспама и этап классификации документов основного массива. Результаты тестирования подтверждают эффективность применения метода для обработки и классификации документов в больших политематических текстовых массивах.

КЛАССИФИКАЦИЯ ТЕКСТОВ, КЛАССИФИКАТОР БАЙЕСА, ЛИНГВИСТИЧЕСКИЙ ДЕСКРИПТОР, ПРОГРАММНЫЙ МОДУЛЬ

Введение

В области автоматической обработки электронных текстов сложился ряд относительно самостоятельных направлений: извлечение объектов и признаков, реферирование, классификация, кластеризация, интеллектуальный поиск, семантический анализ и т.п. [1, 2]. В частности, задача организации эффективного доступа к неструктурированной тематической информации непосредственно связана с задачей классификации электронных текстов, извлекаемых из ресурсов сети Интернет или электронных библиотек. Для решения последней разработано множество эффективных методов, некоторые из которых характеризуются качеством классификации, сравнимым с результатами классификации, выполняемой квалифицированными экспертами [3]. К наиболее используемым методам относятся, в частности, метод опорных векторов, метод логистической линейной регрессии, метод k-ближайших соседей, нейронные классифицирующие сети и т.п. В то же время следует отметить, что данные методы характеризуются высокой сложностью разработки и реализации классифицирующих алгоритмов, высокой вычислительной сложностью, сложностью коррекции функций обучения.

В последнее время получили распространение алгоритмы классификации объектов, лишенные подобных недостатков, например, наивный классификатор Байеса (НКБ) [4, 5]. Этот классификатор обладает простой и реализуемой в реальном времени обучающей процедурой, которую легко модифицировать под особенности конкретной решаемой задачи. Однако базовый классификатор Байеса основан на предположении об условной независимости переменных, что во многих случаях может оказаться неприемлемым, так как гипотеза полной независимости результата классификации от сочетания признаков оказывает существенное влияние на качество классифицирующей

процедуры во многих реальных задачах обработки электронных текстов (в частности, текстов научно-технического содержания). Когда же допущение о независимости выполняется, НКБ, как правило, превосходит другие алгоритмы классификации (в том числе, и многоклассовой) и при этом использует меньший объем обучающих данных.

Одним из подходов, позволяющим учесть совокупность классификационных признаков в анализируемых текстах, является нечеткая классификация [6]. При такой классификации фрагменты данных могут принадлежать нескольким классам, связанным с каждым элементом набором степеней принадлежности. Они указывают силу ассоциации между элементом данных и определенным классом. Нечеткая классификация — процесс определения этих степеней принадлежности и использования их для присвоения элементов данных двум или более классам. В реальных случаях не может быть никаких резких границ между классами, и тогда применение методов нечеткой логики может оказаться перспективным направлением повышения эффективности процедур автоматической классификации текстов. Однако практическая реализация нечетких классификаторов предполагает необходимость использования ряда эвристических назначений (в частности, задания функций принадлежности и интервалов дискретизации значений лингвистических переменных), что зачастую не дает возможности гарантировать точность.

Другим направлением повышения эффективности решения задач, связанных с автоматической классификацией политематических текстовых ресурсов, является гибридное применение байесовской классификации и методов, связанных с возможностью учета в схеме классификатора операций выделения из текстов дескрипторов и ключевых слов. В общем случае, при необходимости классификации большого объема разнородных классов с целью получения конечного множества

массивов, содержащих документы заданной тематики, целесообразно разработать простую в вычислительном отношении процедуру построения соответствующего классификатора.

Целью настоящей работы является разработка и тестирование метода двухэтапной классификации политематических электронных текстов, основанной на комбинированном применении НКБ и модифицированной процедуры байесовской классификации с предварительной фильтрацией исходного массива текстовых документов и выделением лингвистических дескрипторов.

1. Структура и описание предлагаемого метода классификации электронных текстов

Рассмотрим задачу классификации по тематическим рубрикам большого массива документов, имеющего политематический и составной характер. Таким массивом может быть, например, электронный сборник материалов конференции, содержащих аннотации и основной текст; массив документов большой электронной библиотеки, посвященных некоторому общему научному направлению, и т.п. Проблемами, возникающими при реализации такой классификации, являются: наличие служебных элементов и посторонних блоков текста, не относящихся к основной тематике документа; наличие в обучающем массиве аномальных документов (пустых, в неизвестных кодировках и т.п.); сложность автоматического формирования решающих правил для рубрик из-за негативного влияния посторонней информации; снижение качества классификации из-за наложения нескольких рубрик друг на друга; сложность интерпретации результатов классификации из-за неопределенности расположения в тексте информации, релевантной рубрике, и т.п.

Предлагаемая схема решения этой задачи приведена на рис. 1.

Схема предполагает последовательную реализацию процедур двухэтапной обработки исходного текстового массива документов $M1$.

Предполагается, что потенциальный пользователь определил рубрики, в соответствии с которыми должен быть сформирован классифицированный массив документов, содержащий N классов.

Первый этап метода содержит процедуры фильтрации исходного массива $M1$ с целью удаления документов, относящихся к *псевдоспаму*. Под *псевдоспамом* будем понимать все документы массива, не представляющие тематический интерес для потенциального пользователя, вероятность отнесения которых к одному из N классов результирующего массива очень низка.

Выделение псевдоспама будем осуществлять с использованием минимального словаря

дескрипторов, формируемых по тематической направленности рубрик, интересующих потенциального пользователя. Дескриптор – лексическая единица (слово, словосочетание, аббревиатура), служащая для описания основного смыслового содержания документа или формулировки запроса при поиске документа в информационно-поисковой или классифицирующей системе [7]. Дескриптор однозначно ставится в соответствие группе ключевых слов естественного языка, отобранных из текста, относящегося к определенной области знаний. Это позволяет создать уникальные словари дескрипторов для учета их в классификаторе с целью повышения точности классификации. Процедура выделения псевдоспама может быть основана на применении классического подхода, использующего машинное обучение с учителем [5]. В этом подходе требуется наличие обучающей коллекции содержащих дескрипторные термины текстов, на базе которой строится статистический или вероятностный классификатор.

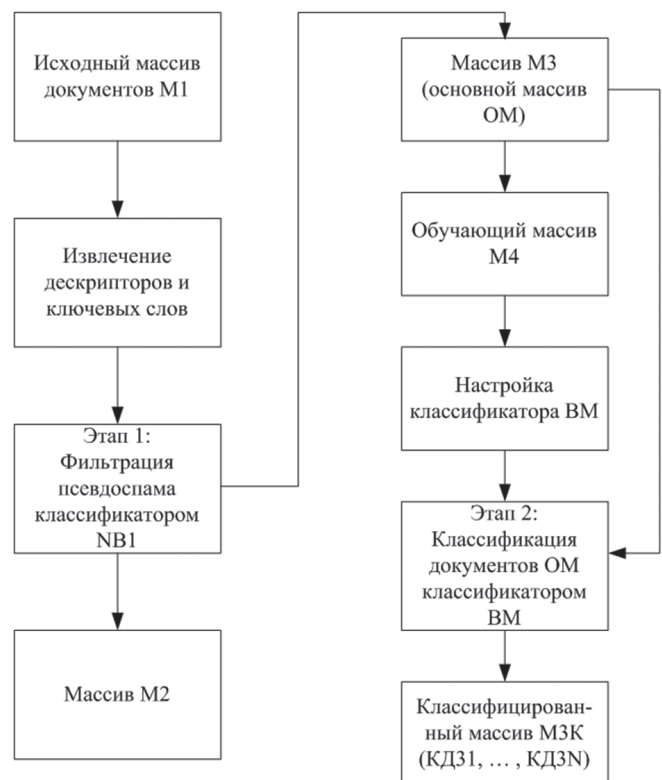


Рис. 1. Схема двухэтапного метода классификации

На первом этапе предлагаемого метода для обработки исходного массива документов с целью выделения псевдоспама используется полиномиальная модель метода Naive Bayes (NB), на основании которой реализуется наивный бинарный классификатор Байеса NB1.

При этом для характеристики анализируемых текстов используется векторная модель представления документов d исходного массива $M1$:

$$d = (t_1, t_2, \dots, t_n), \quad (1)$$

где $t_1, t_2, \dots, t_n$ — дескрипторные термины (признаки) текстов; n — общее количество учитываемых терминов.

Компонентам вектора (1) могут быть поставлены в соответствие бинарная или частотная функции взвешивания. Бинарные векторы представляются как последовательность нулей и единиц: если конкретный термин из словаря выборки встречается в тексте — вес термина будет равен 1, в противном случае — 0. Частотные векторы формируются на основе количества вхождений определенного термина в классе документов.

Для классификатора псевдоспама NB1 выберем бинарную функцию взвешивания, полагая, что при решении задачи фильтрации псевдоспама наличие термина в документе важнее, чем его частота.

Полиномиальная модель классификатора NB1, как и базовая модель байесовской классификации, базируется на понятии условной вероятности принадлежности документа d классу c . В соответствии с теоремой Байеса, такая вероятность определяется следующим образом:

$$P(c|d) = \frac{P(c) \cdot P(d|c)}{P(d)}. \quad (2)$$

Наиболее вероятным классом (по оценке апостериорного максимума), к которому принадлежит документ d , в классификации по методу NB1 является тот, при котором условная вероятность принадлежности документа d классу c максимальна:

$$c_{map} = \arg \max_c P(t_1, t_2, \dots, t_n | c) * P(c). \quad (3)$$

Отметим, что в (3) не учитывается знаменатель из отношения (2), так как для одного и того же документа d вероятность $P(d)$ будет одинаковой.

Предположим, что позиция термина в предложении документа не важна. Тогда условную вероятность для признаков $t_1, t_2, \dots, t_n$ можно представить следующим образом:

$$P(t_1 | c) * P(t_2 | c) * \dots * P(t_n | c) = \prod_i P(t_i | c). \quad (4)$$

Для нахождения наиболее вероятного класса для документа с помощью классификатора NB1 необходимо определить условные вероятности принадлежности документа d для каждого из представленных классов (в задаче классификации псевдоспама рассматриваются два возможных класса: массив M3 (ham) документов, отбираемых для дальнейшего анализа, и массив M2 (spam) документов псевдоспама) отдельно и выбрать класс, имеющий максимальную вероятность:

$$c_{map} = \arg \max_c \left[P(c) * \prod_i P(t_i | c) \right]. \quad (5)$$

В уравнении (5) умножается множество условных вероятностей, что может привести к потере значимости при выполнении операций с плавающей точкой для относительно больших текстов. Поэтому целесообразно применять логарифмическое представление уравнения (5):

$$c_{map} = \arg \max_c \left[\log P(c) + \sum_i \log P(t_i | c) \right]. \quad (6)$$

Вероятности классов $P(c)$ в (5) и (6) оцениваются как отношение количества документов класса в обучающей выборке к общему количеству документов в выборке:

$$P(c) = \frac{N_c}{N}. \quad (7)$$

Условные вероятности для признаков (терминов) t_i оцениваются как отношение количества терминов t_i в классе к общему количеству терминов в этом классе:

$$P(t_i | c) = \frac{\text{count}(t_i | c)}{\sum_t \text{count}(t, c)}, \quad (8)$$

где $t \in V$; V — словарь обучающей выборки.

Здесь принимается предположение о том, что позиция слова в тексте не влияет на вероятность его появления. В этом случае условные вероятности слова, занимающего две разные позиции $k1$ и $k2$ в разных местах документа, будет одинаковой:

$$P(t_{k1} | c) = P(t_{k2} | c). \quad (9)$$

При вычислении вероятностей возможна ситуация, когда какое-либо слово из текста для классификации ни разу не присутствовало в обучающей выборке какого-либо класса, тогда вероятность данного слова в классе и полная вероятность соответствия документа данному классу будут равны нулю. Для устранения таких выбросов используем аддитивное сглаживание Лапласа (Laplace smoothing). Принцип такого сглаживания состоит в том, что к частотам появления всех терминов из словаря искусственно добавляется единица. При этом термины, которые не присутствовали в документах обучающей выборки, получают незначительную, но отличную от нуля вероятность появления и дают возможность отнести документ к одному из формируемых классов. При этом формула (8) трансформируется следующим образом:

$$P(t_i | c) = \frac{\text{count}(t_i | c) + 1}{\sum_t (\text{count}(t, c) + 1)} = \frac{\text{count}(t_i | c) + 1}{(\sum_t \text{count}(t, c)) + |V|}, \quad (9)$$

где в знаменатель правой части добавляется количество слов в словаре обучающей выборки V .

Альтернативой полиномиальной модели классификатора NB1 может быть многофакторная модель или модель Бернулли метода NB. Модель

Бернулли оценивает условные вероятности для признаков (терминов) t_i как часть документов класса c , которые содержат термин t_i по отношению ко всем документам класса c . При классификации тестового документа модель Бернулли игнорирует количество вхождений слова в документ, в то время как полиномиальная модель отслеживает все вхождения термина. В результате этого модель Бернулли обычно делает много ошибок при классификации длинных документов. Условная вероятность для признаков (терминов) t_i в методе Бернулли рассчитывается следующим образом:

$$P(t_i|c) = \frac{N_{ct_i} + 1}{N_c + 2}. \quad (10)$$

Для оценки эффективности классификатора NB1 используем простую метрику эффективности. Пусть в результате классификации документов тестовой выборки, к классу M3 правильно отнесены T3 документов, неправильно — F3, а к классу M2 правильно были отнесены T2 документов, неправильно — F2.

Тогда точность классификации на первом этапе с помощью NB1 определится следующим образом:

$$\text{Prec NB1} = (T3 + T2) / (T3 + T2 + F3 + F2) * 100\%. \quad (11)$$

Второй этап метода содержит процедуры обработки основного массива M3, сформированного после фильтрации псевдоспама, с целью формирования результирующего классифицированного массива документов, содержащего N классов (по количеству тематических рубрик, интересующих пользователя). Отметим, что удаление псевдоспама из исходного массива позволяет существенно сократить количество анализируемых в дальнейшем документов, что дает возможность применения на втором этапе более сложных и точных процедур классификации.

Для решения задачи классификации на втором этапе предлагается использовать модифицированный классификатор Байеса (BM)). Этот классификатор, в отличие от классификатора NB1, имеет следующие особенности: применение локальных словарей дескрипторов для каждого из формируемых классов; возможность учета связей между дескрипторными терминами при реализации процедуры классификации.

При обучении классификатора BM для каждого встреченного в текстах слова рассчитывается и сохраняется его вес — оценка вероятности того, что текст с этим словом принадлежит к одному из возможных классов (одна из возможных аппроксимаций такой оценки состоит в замене бинарной функции взвешивания терминов t_i в векторной модели (1) частотной функцией, ставящей в соответствие этим терминам частоту β_i их появления в

анализируемом документе). Обучающая выборка, используемая при настройке BM, содержит документы, гарантированно принадлежащие хотя бы к одному из заданных тематических классов. Классификация таких документов осуществляется с помощью локальных словарей дескрипторов для каждого из формируемых классов.

На основе данных о классификации лингвистических дескрипторов при составлении локальных словарей производится расчет целесообразности выбора тех или иных его вариантов с учетом требований к системе, по каждому выбранному классу дескрипторов в различных классификациях. При этом учитывается общее соотношение слов, распределенных для каждой из категорий выбранных подклассов дескрипторов. Для формирования локальных словарей дескрипторов для каждого из формируемых на втором этапе классов рассматриваются такие виды классификации, как «Части речи» и «Частотность». Анализ научно-технических текстов позволяет оценить общее процентное разделение по категориям дескрипторов каждого вида классификации: классификация «Части речи»: существительные (60%), глаголы (13,75%), наречия (6,25%), прилагательные (16,25%), причастия (3,75%); классификация «Частотность»: высокочастотные дескрипторы (60%); среднечастотные дескрипторы (22,5%); низкочастотные дескрипторы (17,5%).

На основе данной классификации, соотношения числа дескрипторов и возможностей добычи их из неструктурированной текстовой информации можно утверждать, что для повышения точности классификации: рациональней всего использовать существительные; по частотной встречаемости видов дескрипторов: целесообразно использовать высокочастотные и низкочастотные дескрипторы. Наиболее соответствующей этим требованиям структурой являются аббревиатуры. Аббревиатуры — это существительные, состоящие из усеченных слов, входящих в исходное словосочетание, или из усеченных частей исходного сложного слова, а также из названий начальных букв этих слов (или их частей). Аббревиатуры являются низкочастотными дескрипторами, которые легко получить из текстов (особенно из их аннотаций). В случае, если в текстах отсутствуют аббревиатуры, целесообразно использовать высокочастотные дескрипторы для максимального охвата всего массива текстовой информации. Для качественной классификации текстов на основе выделенных дескрипторов необходимо прибегнуть к предварительной его очистке от шумов, под которыми понимается категория лексем, мешающих адекватной классификации данных (служебных слов, стоп-слов и т.п.). На

втором этапе рассматриваемого метода предусмотрена предварительная обработка текста документов из массива МЗ: приведение слов в начальную форму, удаление служебных слов, вычисление веса для целых фраз, транслитерация.. К стоп-словам будем относить: союзы и союзные слова; местоимения; предлоги; частицы; междометия; указательные слова; цифры; знаки препинания; отдельно стоящие буквы алфавита; вводные слова. При классификации текстов стоит обращать внимание на наличие стоп-слов из вышеперечисленных категорий и их соотношение с общей массой слов и дескрипторов. Общие шумовые слова часто не учитываются классификатором ВМ, однако, они заменяются специальным маркером. Данное обстоятельство имеет практическое значение при составлении классификатора и оценки плотности ключевых слов разного рода, так как игнорирование стоп-слов влияет на некоторые показатели, которые в свою очередь влияют на точность классификации текстовой информации.

В классификаторе ВМ предусмотрена также возможность учета связей между дескрипторными терминами при реализации процедуры классификации на втором этапе. Для этого реализуется специальная процедура поиска слов, вероятность использования которых в качестве связей между выделенными дескрипторными терминами для каждого из классов высока, с последующим выбором дескрипторов, которые может соединить в будущем локальном словаре дескрипторов данная связь. Это позволяет дополнить начальную совокупность дескрипторов, используемых классификатором ВМ, связными тройками дескрипторов вида «дескриптор1-связка12-дескриптор2», наиболее характерными для документов рассматриваемого класса.

Учет в классификаторе таких связных дескрипторов позволяет снизить влияние на качество классифицирующей процедуры электронных текстов гипотезы полной независимости результата классификации от сочетания признаков, которая является одним из наиболее существенных недостатков применения метода Naive Bayes. Формирование совокупности связных дескрипторов осуществляется в предлагаемом методе с применением модифицированного метода Гинзбурга, этапы которого описаны в работе [8]. Процедуры (6)-(11), используемые на первом этапе для определения параметров вероятностного классификатора и оценки качества классификации, дополняются на втором этапе процедурами, реализующими описанные дополнительные функции классификатора ВМ.

Итоговым результатом классификации является формирование классифицированного массива

МЗК, содержащего совокупность документов КД31, КД32,..., КД3N, распределенных по соответствующим классам.

2. Программная реализация и тестирование предлагаемого метода классификации

Для программной реализации классификаторов NB1 и ВМ, а также алгоритмов выделения наиболее значимых атрибутов были использованы объектно-ориентированный язык программирования Java и среды программной разработки (IDE), в частности, Eclipse и NetBeans.

Разработанный программный модуль Booker V позволяет осуществлять классификацию больших объемов текстовой информации по пользовательским категориям, учитывая возможность распределенного хранения массивов информации.

Общая модель архитектуры модуля BookerV приведена на рис. 2.

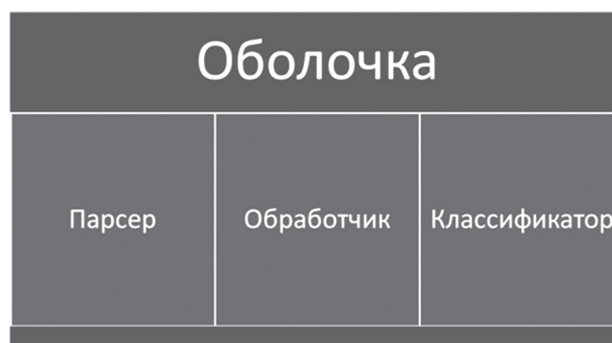


Рис. 2. Общая структура модуля BookerV

Разделение модуля на четыре основных блока позволит вносить изменения в систему в любом из блоков, улучшая его работу, при сохраняющемся интерфейсе взаимодействия с пользователем. В блоке «Оболочка» происходит взаимодействие с пользователем с целью получения от него ссылки на библиотеку классифицируемых данных и выдачи пользователю данных о принадлежности к определенному классу некоторой текстовой информации. Схема передачи данных в модуле BookerV приведена на рис. 3.

Для запуска программы необходимо инициировать исполняемый файл BookerV.jar.

Интерфейс модуля построен с помощью визуальных средств проектирования в среде разработки NetBeans с помощью технологии jxml.

Визуальная оболочка программы состоит из четырех основных структурных блоков: блок загрузки библиотеки; блок системных настроек; блок ввода информации на классификацию; блок вывода информации.

Загрузка входных данных из библиотеки производится через один шлюз, с возможностью дополнительных настроек.

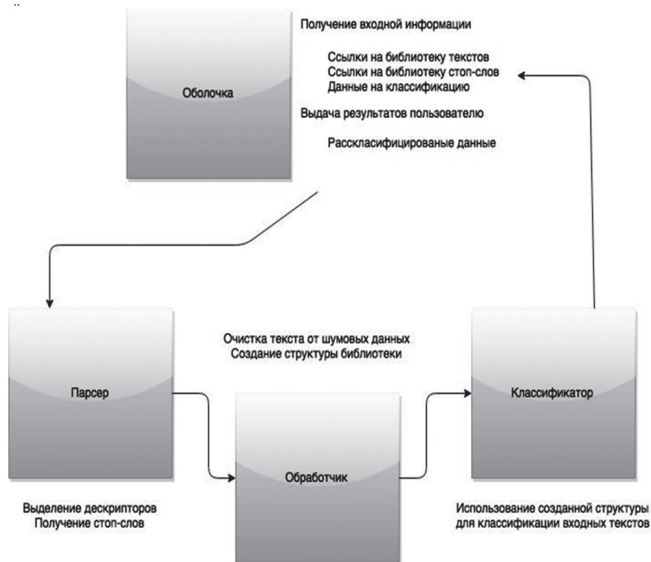


Рис. 3. Передача данных в модуле

В качестве библиотеки может выступать файл или директория операционной системы. При этом есть возможность подгрузки библиотеки из десериализованного файла библиотеки. При выборе файла (LIB_NAME.data / struct.csv) или директории необходимо сделать правильные преднастройки для анализа библиотеки. Пример настроек приведен на рис. 4.

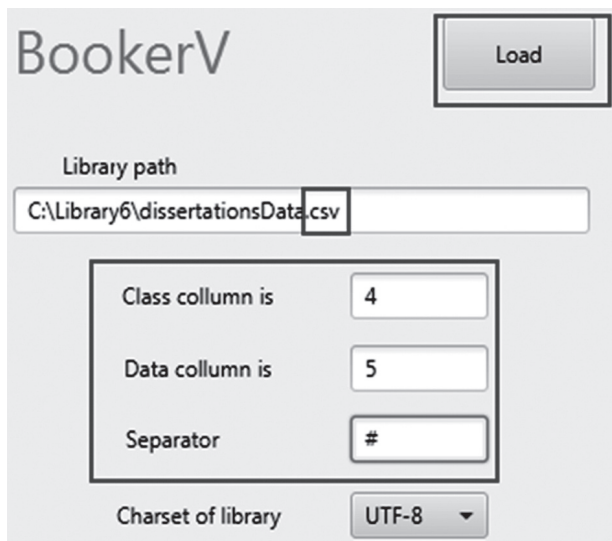


Рис. 4. Настройки модуля (для csv подобных структур)

Для упрощения взаимодействия с системой был реализован алгоритм определения типа входящей информации перед ее классификацией, что позволяет вводить информацию для классификации следующим образом: ссылки на файлы в расширениях .txt и .csv; списки ссылок, воспринимаемые построчно; текст в одну строку; текст многострочный; текстовые csv-подобные структуры.

В блоке вывода информации отображаются результаты классификации документов массива M3. Структура библиотеки сохраняется

самостоятельно по окончании сканирования указанного источника. Для выбора пользовательского имени десериализованного файла библиотеки необходимо ввести его заранее, используя полный путь к желаемому создаваемому файлу. При отсутствии пользовательского файла файл создается в корне хранимого источника библиотеки.

Разработанное приложение создано с целью обработки большого количества текстовой информации, его можно использовать на серверных платформах, в электронных библиотеках, а также в организациях с большим объемом документооборота.

Структура части программного модуля Booker V, реализующей функции классификатора NB1, состоит из двух основных частей (рис.5 и рис.6): первая содержит исходный код модуля, интерфейсы взаимодействия с данными и реализацию классификатора, вторая содержит визуальную оболочку для используемого модуля. Каждый файл с расширением .java реализует различные функции программы: обучение на различных видах выборки; классификацию текстовых документов на элементы M3 (ham), отбираемые для дальнейшего анализа, и документы массива M2 (pseudospam); анализ производительности реализованных методов на тестовых выборках. Визуальный интерфейс (вторая часть модуля) реализован с помощью технологии jxml для качественного отображения информации и простоты использования на различных платформах, включая OS-X и Android.

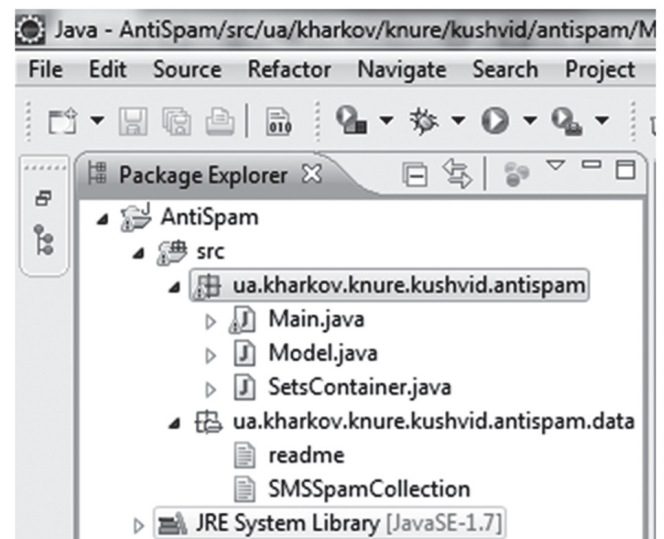


Рис. 5. Структура программного модуля Booker V (1)

Фрагмент программы, реализующей основные функции классификатора NB1, приведен на рис. 7.

Структура части программного модуля BookerV, реализующей функции классификатора BM для документов из массива M3, фрагмент которой приведен на рис. 8, содержит файлы, иницирующие

реализацию функций предварительной очистки документов от шумов; применения локальных словарей дескрипторов для каждого из формируемых классов; выделения и учета связей между дескрипторными терминами при реализации процедуры классификации.

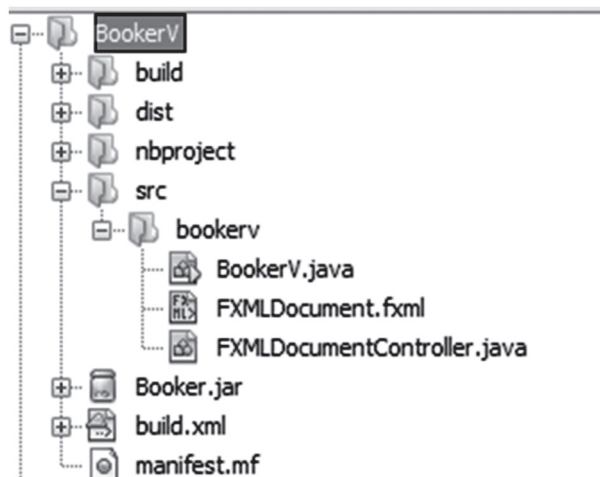


Рис. 6. Структура программного модуля Booker V (2)

```
public static String classification(String string, Model model,
    boolean preprocess) {
    boolean isSpam = true;
    TreeMap<String, Integer> hamMap = model.hamMap;
    TreeMap<String, Integer> spamMap = model.spamMap;
    double laplaceFactor = model.laplaceFactor;

    HashSet<String> uniqueWordSet = model.uniqueWordSet;

    double ham, spam;

    ham = Math.log((double) hamMap.size()
        / (hamMap.size() + spamMap.size()));
    spam = Math.log((double) spamMap.size()
        / (hamMap.size() + spamMap.size()));

    if (preprocess) {
        string = reformation(string);
    }

    String[] words = string.split(" ");

    for (String word : words) {
        if (hamMap.containsKey(word))
            ham += Math
                .log((hamMap.get(word) + laplaceFactor)
                    / (hamMap.size() + laplaceFactor
                        * uniqueWordSet.size()));
        else
            ham += Math
                .log((laplaceFactor)
                    / (hamMap.size() + laplaceFactor
                        * uniqueWordSet.size()));

        if (spamMap.containsKey(word))
            spam += Math.log((spamMap.get(word) + laplaceFactor)
                / (spamMap.size() + laplaceFactor
                    * uniqueWordSet.size()));
        else
            ham += Math.log((laplaceFactor)
                / (spamMap.size() + laplaceFactor
                    * uniqueWordSet.size()));
    }

    if (ham >= spam)
        return "ham";
    else
        return "spam";
}
```

Рис. 7. Программная реализация основных функций NB1

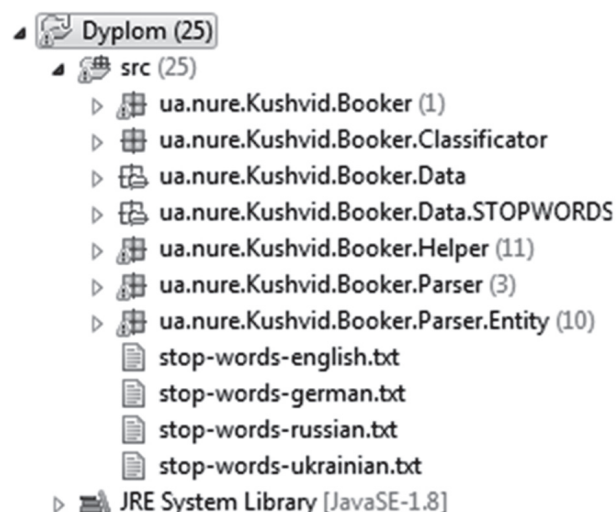


Рис. 8. Структура программного модуля Booker V (3)

Система тестировалась на основе бесплатных электронных библиотек, тестовой выборки из массива «Reuters-21578», а также свободно распространяемых списков стоп-слов от компании Google.

При тестировании полиномиальной модели классификатора NB1 использована часть архива выборки электронных документов, приведенного в ресурсе <http://archive.ics.uci.edu/ml/dataset/SMS+Spam+Collectio>.

По результатам тестирования (с применением сглаживания по Лапласу) точность классификации (Prec NB1) достигала для большинства текстовых выборок 97%.

При тестировании классификатора ВМ использована часть электронной библиотека «ITAR» объемом 21 ГБ и коллекция аннотаций авторефератов диссертационных работ объемом 2 ГБ.

Поскольку в классификаторе ВМ реализована возможность пользовательского создания или редактирования личных списков стоп-слов путем указания пути хранения текстовых файлов с их наборами, каждый отдельный файл считался отдельной группой и проход по группам пользовательских стоп-слов происходил в порядке их хранения в директории.

На разных выборках библиотек формировались персональные наборы дескрипторов трех видов: аббревиатуры (например, HDL, IP, JPEG, MIMO, БА, БИС, БПИС, БТС, ВА, ВЗ, ВКАС и т.п.), связанные дескрипторы (например, «использование в конкретной функциональной задаче - позволяет расширить - функции администрирования», «модель актуализации оперативных данных - позволяет автоматизировать - функция обработки документов» и т.п.) и общие ключевые слова (например, центральный процессор, оценка быстродействия, программное обеспечение и т.п.).

При исследовании возможностей классификатора ВМ были использованы два метода тестирования выборок, программно реализованные в модуле BookerV: `testModel` и `crossValidation`. Данные функции вызываются после обучения классификатора на обучающей выборке во время его проверки на тестовых данных. Это позволяет: сократить вычислительное время благодаря сокращению числа слов, учитываемых в каждом классе; уменьшить возможность переобучения; снизить влияние выбросов и шумов на результат классификации.

В процессе тестирования осуществлялся подбор наилучшего значения коэффициента размытия по Лапласу (KPM). Исследовано влияние этого коэффициента на точность классификации. Анализ был проведен на примере построения динамических выборок с разными коэффициентами KPM. Кроме того, исследовался алгоритм выбора вида дескрипторов в зависимости от характера выборки для дополнительного повышения качества классификации. Точность двухэтапной классификации текстовых документов с помощью модуля BookerV, для разных серий тестовых экспериментов находилась в диапазоне от 84% до 99,26%.

Выводы

Применение рассмотренного метода классификации является перспективным для классификации по тематическим рубрикам больших массивов документов, имеющего политематический и составной характер. Такими массивами могут быть, например, электронные сборники материалов конференции, содержащих аннотации и основной текст; массив документов большой электронной библиотеки, посвященных некоторому общему научному направлению, и т.п. Первый этап предложенного метода, реализованного в виде программного модуля BookerV, осуществляет фильтрацию исходного массива с целью удаления документов, не представляющие тематический интерес для потенциального пользователя, вероятность отнесения которых к одному из классов результирующего массива очень низка. На втором этапе осуществляется процедура основной классификации, использующая модифицированный полиномиальный байесовский классификатор. Модификация

состоит в применении связанных дескрипторов, что позволяет снизить влияние на качество классифицирующей процедуры электронных текстов гипотезы полной независимости результата классификации от сочетания признаков, которая является одним из наиболее существенных недостатков применения метода Naive Bayes.

Результаты тестирования предложенного метода подтверждают целесообразность его использования при решении широкого класса задач классификации политематических электронных текстов. Перспективным развитием метода является проведение экспериментов по усовершенствованию классификаторов NB1 и ВМ для работы с тематическими документами, характеризующимися наличием неравномоощных подтем, а также выбор наиболее эффективных комбинированных критериев оценки качества классификации.

Список литературы:

1. *Sebastiani F.* Machine learning in automated text categorization/ *Sebastiani F.* // ACM Computing Surveys, 34(1), 2002. — pp. 1–47.
2. Автоматическая обработка текстов на естественном языке и компьютерная лингвистика : учеб. пособие / Большакова Е.И., Клышинский Э.С., Ландэ Д.В., Носков А.А., Пескова О.В., Ягунова Е.В. / — М.: МИЭМ, 2011. — 272 с.
3. *Епрев А.С.* Автоматическая классификация текстовых документов / А.С. Епрев // Математические структуры и моделирование. — 2010 — Вып. 21. — С. 65–81.
4. *Domingos P.* On the optimality of the simple Bayesian lassifier under zero-one loss / P. Domingos, M. Pazzani // Machine Learning. — 1997. — № 29. — pp. 103–137.
5. *Петровский М.И.* Алгоритмы машинного обучения для задачи анализа и рубрикации электронных документов / М.И. Петровский, В.В. Глазкова // Вычислительные методы и программирование. — 2007. — Т. 8. — № 2. — С. 57–69.
6. *Рыжов А.П.* О качестве классификации объектов на основе нечетких правил / А.П. Рыжов // Интеллектуальные системы — 2005 — Т. 9. — С. 253–264.
7. *Чалая, Л. Э.* Оценивание пертинентности лингвистических дескрипторов в системах информационного поиска документов [Текст] / Л.Э. Чалая, Ю.Ю. Харитонова // Восточно-европейский журнал передовых технологий. — 2015. — № 1/9(73). — С. 46–53.
8. *Чалая, Л. Э.* Метод поиска пертинентных связей между концептами проектируемых онтологий [Текст] / Л.Э. Чалая, А.В. Чижевский, Е.Б. Волощук // Комп'ютерно-інтегровані технології: освіта, наука, виробництво. — 2016. — №22. — С. 50–56.

Поступила в редколлегию 16.11.2016

УДК 519.7



А.С. Пузик
ХНУРЭ, Украина, as308@mail.ru

ЛЕКСИКОГРАФИЧЕСКАЯ СИСТЕМА ЭЛЕКТРОННОГО ТЕРМИНОЛОГИЧЕСКОГО СЛОВАРЯ

Статья посвящена описанию программной системы русско-украинско-английского терминологического словаря. В статье описаны подходы к первичной обработке исходных данных. Описана организация модели данных для трехязычного словаря и перспективы развития модели при переходе к многоязычным словарям. Приведена база данных и архитектура программной системы. Показаны базовые принципы работы со словарем с точки зрения пользователя.

ЭЛЕКТРОННЫЙ СЛОВАРЬ, АЛГЕБРА КОНЕЧНЫХ ПРЕДИКАТОВ, OFFICE AUTOMATION,
БАЗА ДАННЫХ, ПАРСИНГ, ШАБЛОН MVVM

Введение

Компьютерные системы давно и прочно вошли в наш мир. На них основаны как простые веб-сайты, так и сложные информационные порталы. С конца прошлого века начала свое развитие компьютерная лексикография. Ручной труд лексикографа переводится в электронную форму, на смену бумажным словарям приходят электронные, которые доступны в том числе и через интернет. Электронные словари могут использоваться как база для дальнейшего развития систем интеллектуальной обработки естественной речи, так и в качестве средства получения актуальной переводной информации. Кроме того электронные словари, в отличие от бумажных, позволяют вносить исправления по мере необходимости, без больших затрат, тем самым позволяя оперативно реагировать на возможные недоработки, изменения в языке, тем самым поддерживая актуальность словаря.

В данной статье описывается подход к построению программной системы электронного трехязычного терминологического словаря по информатике и радиоэлектронике. Описаны проблемы, связанные с представлением, обработкой и хранением данных, приведена лексикографическая база данных и описана внутренняя архитектура системы. Кроме того приведены примеры работы с программной системой.

1. Постановка проблемы и цели разработки

Построение электронного словаря является трудоёмкой процедурой состоящей из нескольких шагов. При построении программной системы трехязычного словаря необходимо учитывать особенности описания типов и структур данных в базе данных, способы изменения состояний данных и способы извлечения данных, а также аспект их целостности. Часть подходов к решению этих проблем решаются при помощи средств лексикографических систем [1, 2] и алгебры конечных предикатов [3, 4].

При построении программной системы электронного словаря в качестве исходных данных использовались предварительно отсканированные и распознанные страницы из двуязычного русско-украинского словаря по информатике и радиоэлектронике. Одной из проблем преобразования, является корректура отсканированных текстов. В данной работе требуется показать подходы к обработке отсканированных текстов словарных статей, построить адекватную модель данных для трёхязычного словаря и реализовать программную систему для трехязычного терминологического русско-украинско-английского электронного словаря по информатике и радиоэлектронике.

2. Обработка отсканированных текстов

Создание электронных словарей состоит из нескольких этапов. Сначала сканируются и распознаются словарные статьи, производится коррекция артефактов распознавания. Следующим шагом полученный текст разбивается на массив отдельных словарных статей, а потом производится их декомпозиция по формальным признакам [2].

Документы формата «doc» были исходными для данной работы. Формат «doc» - это бинарный формат файлов, который используется в программе MSWord.

Исходный двуязычный словарь построен на основе алфавитно-гнездового принципа [5]. Русское слово-термин является заголовочным словом. Гнездо включает терминологические словосочетания, элементом которых является заголовочный термин. Заголовочное слово заменяется в терминологических словосочетаниях на тильду, а сами терминологические словосочетания строятся таким образом, чтобы тильда была на первом месте. На основании этой информации можно извлечь нужные данные из отсканированных файлов.

Формат «doc» имеет ряд особенностей, из-за которых напрямую использовать данные документов

довольно затруднительно. Разные секции документов содержат распознанный текст, который надо сгруппировать по определенным признакам. Также в тексте содержится довольно большое количество артефактов распознавания, которые надо исправить для дальнейшей обработки информации (рис. 1).

Одним из способов получения необходимой информации является доступ к содержимому документов при помощи технологии Office COM Automation посредством VBA[6]. Этот способ является наиболее простым подходом, поскольку MSWord предоставляет программные интерфейсы для парсинга документов.

При подходе к созданию многоязычных словарей необходимо учитывать такой аспект представления данных, как кодировка. Существует формат кодирования юникод (Unicode), который позволяет представлять символы любого языка в едином формате. Таким образом после обработки документов MSWord, данные были извлечены и сохранены в текстовый юникодный формат. В полученных файлах хранилась только информация о терминах без учета форматирования, что облегчило их дальнейшую обработку.

Неправильно распознанные символы, знаки переносов, пустые строки, латинские буквы вместо кириллицы, неправильные скобки являлись основной проблемой корректуры. Однако в этих

ошибках были определенные закономерности, что позволило организовать их исправление в автоматизированном режиме при помощи регулярных выражений.

Пример из разбитого на переводы строк, но не обработанного файла:

...
1.(моно, не)хроматическая (
2.моно, не)- хроматична аберація
...
1.(-виток
2.ампер-виток
...

В первом случае скобка из верхней строки должна принадлежать нижней, во втором скобка — артефакт распознавания.

Для заполнения внутренних структур словаря при дальнейшей обработке были выбраны такие характеристики терминов, как отрасль знаний, семантика, изменяемая часть слов и прочее.

3. Организация модели данных трехязычного словаря

Обработка, хранение и представление пользователю являются основными проблемами при построении программной системы. Сложная структура лингвистического материала является одной из причин, которые возникают при создании электронного словаря. Частично проблемы и подходы

А

аббревіатора абрезіатора	нечётно-
абераці́йний абераци́йний	симметри́чная ~ непáрно-
абера́ція аберация	симметри́чна аберация
~ антáнны аберация антáни	оптёческая ~ оптёчна аберация
~ восстані́вленного фр́нта волні	поперáчная ~ поперáчна
абера́ція відні́вленого фр́нту	абера́ція
хвѐлі	проді́льная ~ позді́жня
~ восстані́вленной волні	абера́ція
абера́ція відні́вленої хвѐлі	произвільная ~ дов'ільна
~ ви́сшего пор́ядка аберация	абера́ція
вѐшого пор́ядку	сагитта́льная ~ сагитáльна
~ гологра́ммы аберация	абера́ція
гологра́ми	стигматёческая ~ стигматёчна
~ зáркала аберация дзáркала	абера́ція
~ изобра́жения аберация	сферёческая ~ сферёчна
зобра́ження	абера́ція
~ луча́ аберация пріменя	термооптёческая ~ термооптёчна
~ па́рвого пор́ядка аберация	абера́ція
па́ршого пор́ядку	угло́вая ~ кутова́ аберация
~ полож́ения аберация полі́ження	чётно-симметри́чная ~ па́рно-си-
~ при сканё́ровании аберация при	метри́чна аберация
сканува́нні	електрі́нно-оптёческая ~ элект-
~ свáта аберация св'тла	рі́нно-оптёчна аберация

Рис. 1. Пример входных данных

к их решению описаны для двуязычных словарей и лексикографических систем в целом [1, 2].

Двуязычный словарь можно представить в виде набора переводных эквивалентов для каждого термина. При этом некоторые термины могут иметь многозначную семантику, соответственно при построении связи это надо учитывать. Когда один из языков является основным, то переводные эквиваленты приводятся относительно этого языка. В случае двустороннего перевода появляется второй аналогичный список для другого языка. При этом семантика самого термина становится размытой между переводными эквивалентами языков.

В базе данных электронного словаря переводные эквиваленты и связи между ними будут храниться в соответствующих таблицах. При этом для омонимичных терминов количество связей будет увеличиваться, а выделение семантики термина будет являться дополнительной задачей, которая потребует дополнительных усилий для решения. При переходе к многоязычным словарям, особенно, когда планируется, свободное переключение между языками, появляется проблема увеличения количества связей пропорционально количеству языков. Это случай связи многие ко многим.

Для решения данной проблемы предлагается введение дополнительного уровня косвенности. Им будет являться абстракция, обозначающая семантику термина. Это позволит перейти от отношения многие ко многим к отношению один ко многим (рис. 2). Также при таком подходе появляется привязка термина к семантике, что открывает возможность использовать контекст для перевода текстов. При дальнейшем развитии словаря терминам можно будет добавлять толкования и семантически однозначные примеры использования.

Таким образом база данных электронного словаря будет содержать таблицу терминов и таблицу переводных эквивалентов. Благодаря такому подходу останется возможность получения всех переводов термина, включая семантически разные, путем простой выборки из таблицы переводных эквивалентов.

Данный подход к построению электронного словаря позволяет перейти от двуязычного словаря к многоязычному словарю.

4. Описание лексикографической базы данных и архитектуры программной системы

На основе вышеизложенного материала была построена лексикографическая база данных трехязычного словаря (рис. 3).

Таблица Conception:

- а) Id – идентификатор термина;
- б) ParentId – идентификатор родительского термина, используется для терминологических словосочетаний;
- в) Semantic – идентификатор семантики термина;
- г) Topic – идентификатор отрасли знаний термина.

Таблица Description:

- а) Id – идентификатор переводного эквивалента;
- б) Description – написание переводного термина в именительном падеже, единственного числа на любом из языков;
- в) ConceptionId – идентификатор термина, к которому относится переводной эквивалент;
- г) Language – идентификатор языка, к которому относится переводной эквивалент;
- д) PartOfSpeech – идентификатор части речи, для терминов, отличных от существительных;

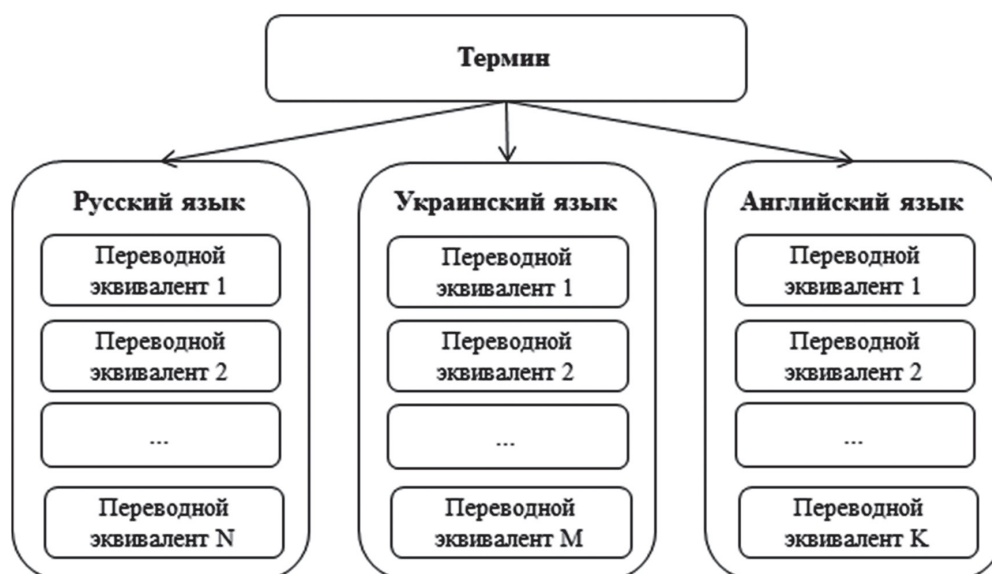


Рис. 2. Предлагаемая модель хранения данных

е) **ChangableType** – идентификатор изменяемой части слова (например род. – родительный падеж, мн. – множественное число);

ж) **ChangablePart** – написание изменяемой части речи (например окончание в родительном падеже «-та»);

Таблица **Semantic**:

а) **Id** – содержит только идентификатор семантики, в дальнейшем будет расширена.

Таблица **SemanticTranslation**:

а) **Id** – идентификатор перевода семантики;

б) **SemId** – идентификатор семантики, для которой этот перевод;

в) **LangForId** – идентификатор языка, для которого этот перевод;

д) **Translation** – непосредственно сам перевод на языке **LangForId**.

Таблица **Topic**:

а) **Id** – содержит только идентификатор отрасли знаний, в дальнейшем будет расширена.

б) **TopicTranslation**

а) **Id** – идентификатор перевода отрасли знаний;

б) **TopicId** – идентификатор отрасли знаний, для которой этот перевод;

в) **LangForId** – идентификатор языка, для которого этот перевод;

д) **Translation** – непосредственно сам перевод на языке **LangForId**.

Таблица **PartOfSpeech**:

а) **Id** – содержит только идентификатор части речи, в дальнейшем будет расширена.

Таблица **PartOfSpeechTranslation**:

а) **Id** – идентификатор перевода части речи;

б) **PartOfSpeechId** – идентификатор части речи, для которой этот перевод;

в) **LangForId** – идентификатор языка, для которого этот перевод;

д) **Translation** – непосредственно сам перевод на языке **LangForId**.

Таблица **ChangablePartType**:

а) **Id** – содержит только идентификатор изменяемой части слова, в дальнейшем будет расширена.

Таблица **ChangablePartTypeTranslation**:

а) **Id** – идентификатор перевода изменяемой части слова;

б) **PartOfSpeechId** – идентификатор изменяемой части слова, для которой этот перевод;

в) **LangForId** – идентификатор языка, для которого этот перевод;

д) **Translation** – непосредственно сам перевод на языке **LangForId**.

Таблица **Languages**:

а) **Id** – идентификатор языка, используемого в словаре;

б) **Name** – название языка по умолчанию.

Таблица **LangTranslation**:

а) **Id** – идентификатор языка, используемого в словаре;

б) **LangIdNeeded** – идентификатор языка, который переводится;

в) **LangForId** – идентификатор языка, для которого этот перевод;

д) **Translation** – непосредственно сам перевод на языке **LangForId**.

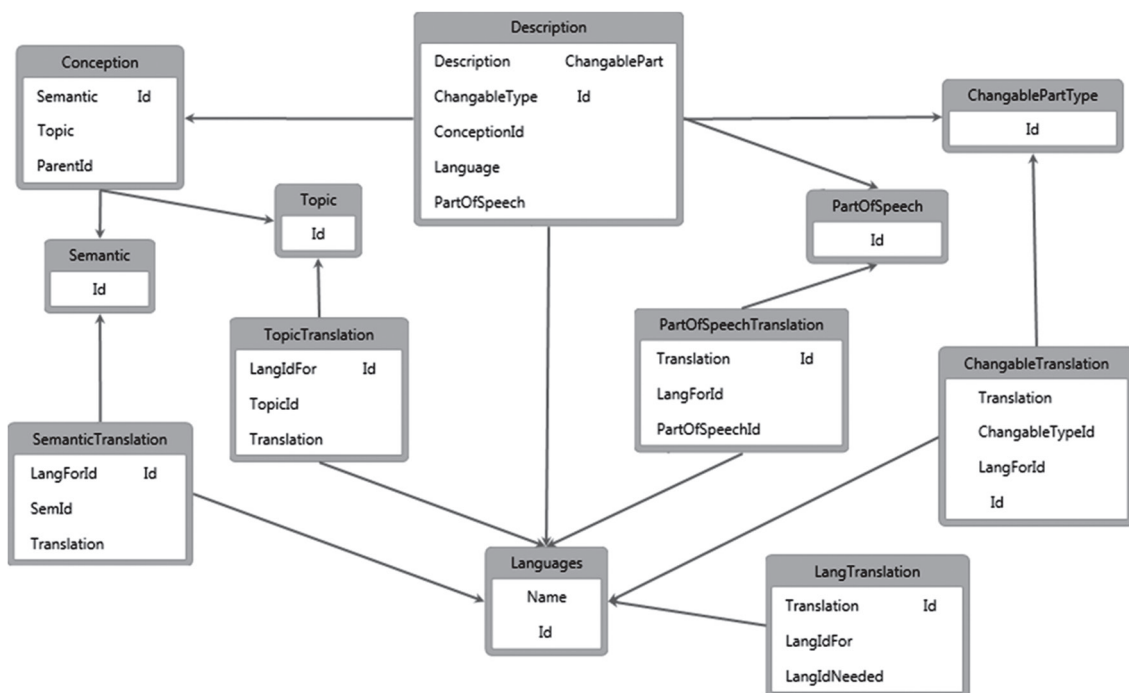


Рис. 3. Лексикографическая база данных трехязычного словаря

Для построения программной системы электронного словаря был использован шаблон проектирования MVVM (Model-View-ViewModel, Модель—Представление—Модель представления) (рис. 4).

Модель — это термин обозначающий классы реализации внутренних данных и логику взаимодействия между ними в программной системе. Представление — это термин, который обозначает классы, обеспечивающие взаимодействие пользователя с графическим интерфейсом. Модель представления — это посредник между моделью и представлением. Модель представления с одной стороны обрабатывает логику представления данных, а с другой стороны передает информацию о действиях пользователя в модель.

Благодаря такому подходу интерфейс пользователя отделён от основной логики программы, что позволяет вносить изменения в различные компоненты программной системы независимо друг от друга.

5. Примеры использования программной системы трёхязычного словаря

Программная система трёхязычного словаря предназначена для создания, редактирования, просмотра терминов и их переводных эквивалентов на русском, украинском и английском языках (рис. 5).

Основной язык выбирается в выпадающем списке сверху по центру главного окна. В левой панели находятся все переводные термины сгруппированные по алфавиту основного языка. Словарь поддерживает функцию поиска, для этого в окне «Поиск» надо ввести интересующую часть слова. Поиск может осуществляться как с учетом ударения, когда ударная буква помечается символом «#», так и без него. После нажатия кнопки «Найти» в левой панели выбирается термин отвечающий первому найденному совпадению.

Центральная панель отображает переводные эквиваленты для всех языков для выбранного термина.

Правая панель предназначена для редактирования терминов. Для добавления переводного эквивалента в поле «Введите текст» надо ввести текст переводного эквивалента термина для выбранного языка. После этого надо нажать кнопку «Добавить» в секции «Переводной эквивалент». Для изменения существующего переводного эквивалента, его надо выбрать в центральной панели, в поле «Введите текст» обновить значение и нажать кнопку «Изменить». Для удаления переводного эквивалента надо его выбрать в центральной панели и нажать кнопку «Удалить».

Кроме того для переводных эквивалентов может быть указана часть речи в выпадающем меню «Часть речи». Если необходимо показать изменение фонем или ударения, то в поле «Изменяемая часть» вводится текст изменяемой части. Это может быть либо слово целиком (например, для слова «вісь» в этом поле будет «осі»), либо другая изменяемая часть (например, для слова «дейтрон» надо ввести «-на», а для слова «нелінійність» — «-ності»). В поле «Тип изменяемой части» вводится тип изменяемой части (например, родительный падеж «род.»)

Для термина в целом можно ввести дополнительный семантические пометы. Условное обозначение отрасли в которой применяется термин (например, мат., физ. и т.д.) выбирается в выпадающем меню «Тема». В меню «Семантика» выбирается краткое толкование семантики термина (например, слово «акт» может иметь два значения «действие» и «документ»). После нажатия кнопки применить к термину данные из этих полей вносятся в базу данных. При нажатии кнопки «Добавить термин» пользователю будет представлено окно,

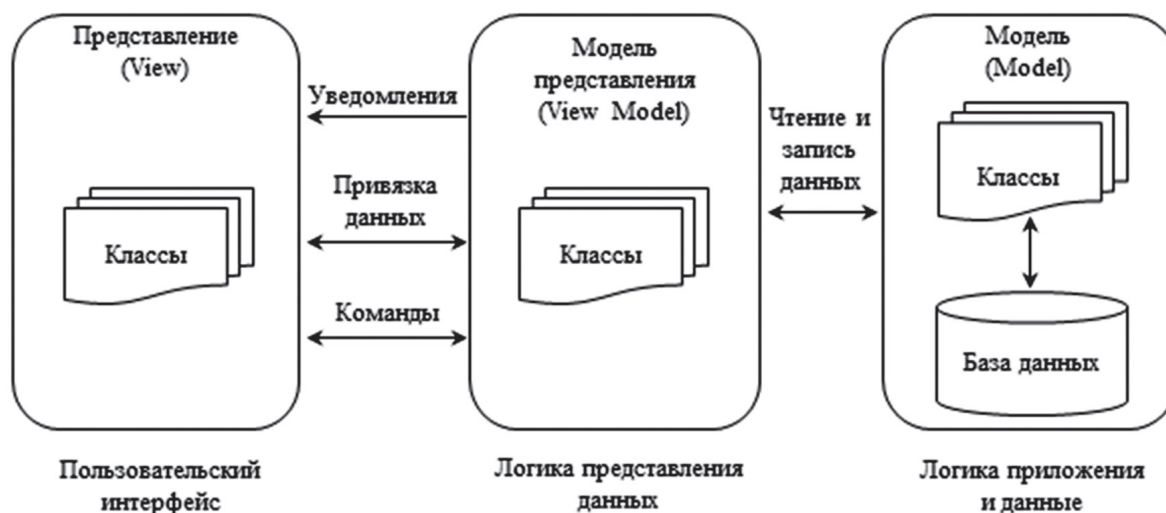


Рис. 4. Архитектура программной системы

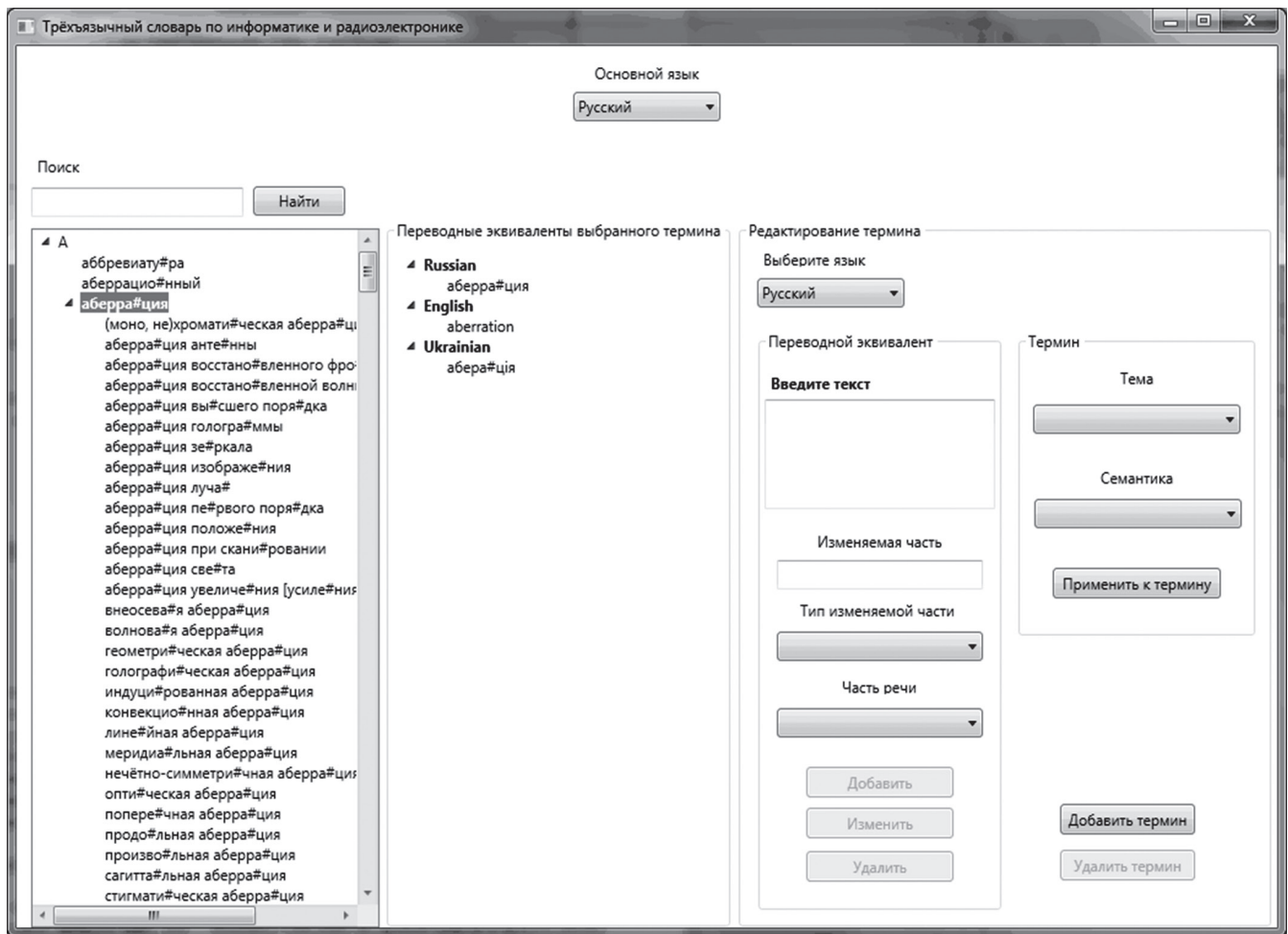


Рис. 5. Главное окно

в котором можно добавить термин, заполнив необходимые поля. При нажатии кнопки «Удалить термин», выбранный в левой панели термин будет удален, со всеми переводными эквивалентами.

Выводы

Таким образом, была создан электронный трехязычный терминологический русско-украинско-английский словарь по информатике и радиоэлектронике. При его создании были использованы подходы, позволяющие избежать ранее описанных проблем, возникающих при построении электронных словарей [1, 5]. В частности при помощи регулярных выражений удалось откорректировать входные данные. При помощи математического аппарата теории лексикографических систем и алгебры конечных предикатов [7] удалось избежать избыточности в решениях. Предложенный подход к построению модели данных словаря можно также применить и для многоязычных словарей.

В статье приведено детальное описание лексикографической базы данных словаря и архитектуры программной системы. Приведены примеры использования словаря пользователем.

Одним из достоинств системы является изначальное равноправие языков, что позволяет назначать основной язык в зависимости от текущих потребностей.

Список литературы

1. Широков В.А. Комп'ютерна лексикографія. — Київ: науково виробниче підприємство «Видавництво «Наукова думка» НАН України», 2011. — 352 с.
2. Рабулець О.Г., Широков В.А., Якименко К.М. Дієслово в лексикографічній системі — К.: Довіра, 2004. — 259 с.
3. Бондаренко М.Ф., Шабанов-Кушнарченко Ю.П. Теория интеллекта: учеб. — Харьков: Изд-во СМИТ, 2006. — 571 с.
4. Бондаренко М.Ф., Шабанов-Кушнарченко Ю.П. Мозгоподобные структуры: справочное пособие. Том первый — К.: Наукова думка, 2011. — 460 с.
5. Остапова И.В. Лексикографическая структура этимологических словарей и их представление в цифровой среде // Прикладная лингвистика и лингвистические технологии: сборник научных трудов. — 2007. — С. 236-245.
6. Word VBA reference // [Електр. Ресурс]. — Режим доступа: <https://msdn.microsoft.com/en-us/library/office/ee861527.aspx>
7. Вечирская И.Д. Разработка трехязычного терминологического словаря на основе алгебры конечных предикатов // Бионика интеллекта: науч.-техн. журнал. — 2011. — № 2(76). — С. 109-113.

Поступила в редколлегию 17.11.2016.



А.Л. Ерохин¹, А.С. Нечипоренко², А.С. Бабий³, А.П. Турута⁴

¹ ХНУРЭ, г. Харьков, Украина, andriy.yerokhin@nure.ua

² ХНУРЭ, г. Харьков, Украина, alinanechiporenko@gmail.com

³ ХНУРЭ, г. Харьков, Украина, apratster@gmail.com

⁴ ХНУРЭ, г. Харьков, Украина, alexey.turuta@gmail.com

ПРИМЕНЕНИЕ ГЛУБОКИХ СВЕРТОЧНЫХ НЕЙРОННЫХ СЕТЕЙ ДЛЯ КЛАССИФИКАЦИИ РИНОМАНОМЕТРИЧЕСКИХ ДАННЫХ

Статья посвящена разработке метода предварительной обработки риноманометрических данных и программной системы для оценки функции носового дыхания, что позволяет осуществлять автоматизированное обнаружение и удаление некорректных измерений в режиме реального времени. Предлагается решение задачи классификации риноманометрических данных на основе использования глубоких сверточных нейронных сетей.

ГЛУБОКИЕ СВЕРТОЧНЫЕ НЕЙРОННЫЕ СЕТИ, КЛАССИФИКАЦИЯ, АЭРОДИНАМИЧЕСКИЕ ХАРАКТЕРИСТИКИ, РИНОМАНОМЕТРИЯ

Введение

Около 20% населения земного шара страдает от заложенности носа. Заложенность носа является распространенным клиническим симптомом у пациентов с заболеваниями носа и околоносовых пазух. Медицинский диагноз в ринологии — довольно сложный процесс, который зависит от результатов субъективных и объективных методов оценки функции носового дыхания и опыта врача-отоларинголога. Использование комплекса методов субъективной и объективной диагностики позволяет повысить эффективность диагностики функции носового дыхания.

Согласно рекомендациям комитета по стандартизации и объективной оценке носового дыхания, «золотым стандартом» объективной диагностики является метод активной передней риноманометрии (ПАРМ) [1]. Результатом измерений по методу ПАРМ являются два сигнала: расход воздушного потока через носовую полость и интраназальное дифференциальное давление воздушного потока. На основании этих измерений вычисляется основной диагностический параметр ПАРМ — коэффициент носового сопротивления [2, 3]. Таким образом, по риноманометрическим данным определяют соотношение между давлением воздуха и интраназальным воздушным потоком. Наиболее полный обзор и анализ критериев оценки функции носового дыхания содержится в работе [4].

Проблема регистрации биомедицинских сигналов занимает центральное место при проектировании медицинских диагностических систем. Для регистрации физиологических сигналов необходимо использование высокочувствительных датчиков. Использование таких датчиков всегда связано с регистрацией шумов. Поэтому необходима предварительная обработка сигналов. Важной задачей предварительной обработки является удаление искажения сигнала. Этап предварительной

обработки представляет собой фильтрацию шумов и сглаживание сигналов. Основные методы фильтрации, применяющиеся на данном этапе, приведены в работах [5, 6].

С другой стороны, мы должны учитывать такие технические параметры систем как разрешение, дрейф напряжения смещения, дрейф коэффициента усиления, частоту среза и др. Все они могут привести к искажению сигналов. Во избежание ложных данных измерений в работах [7, 8] была предложена процедура калибровки. Во время процедуры измерения могут наблюдаться смещение маски, нарушение соединений между маской и датчиком расхода воздуха или датчиком дифференциального давления, попадание инородных масс в соединительные трубки частей системы [7]. В результате, мы имеем дополнительные искажения сигнала.

Все описанные выше факторы формируют некорректные измерения. Они влияют на обработку риноманометрических данных и последующую диагностику. Кроме того, только высококвалифицированный персонал может контролировать точность измерений по методу ПАРМ. Таким образом, все эти факты снижают диагностическую ценность метода активной передней риноманометрии.

Целью данной работы является разработка компонента программно-аппаратной системы для риноманометрических измерений. Данный компонент должен реализовывать две функции: автоматизированное обнаружение и удаление неправильных измерений в режиме реального времени и классификацию риноманометрических данных.

1. Анализ литературы и постановка задачи

Известно множество методов оценки функции носового дыхания. В большинстве случаев в клинической практике вычисление коэффициента носового сопротивления рассчитывается по формуле:

$$R = \frac{\Delta p}{Q}. \quad (1)$$

Данный метод основан на предположении о линейной зависимости между объемной скоростью воздушного потока и перепадом давления. Это означает, что режим потока в полости носа считается полностью ламинарным [8, 9]. В работах [10, 11] анализируется коэффициент сопротивления, который определяется в соответствии с выражением:

$$O = \frac{\Delta p}{Q^2}. \quad (2)$$

Выражение (2) описывает режим турбулентного течения в носовой полости. Наиболее адекватным с физической точки зрения методом описания аэродинамических характеристик процесса дыхания является метод Рехрера [12]. В соответствии с данным методом, дифференциальное давление вычисляется по формуле:

$$\Delta p = k_1 Q + k_2 Q^2, \quad (3)$$

где k_1 — коэффициент ламинарного потока, k_2 — коэффициент турбулентного потока. Расчёт коэффициентов k_1 и k_2 в соответствии с уравнением Рехрера приведен в работах [13, 14, 15]. Однако, следует отметить, что все коэффициенты, приведенные выше, имеют размерность. Данный факт существенно снижает их диагностическую ценность.

В работе [16] авторами была предложена методика расчета гидродинамического коэффициента сопротивления носовой полости, который является безразмерным, учитывает режимы потока через носовые дыхательные пути, а также не зависит от анатомо-физиологической конфигурации носовой полости.

Данные для анализа представлены в виде зависимости $\Delta p = f(Q)$ для повторяющихся циклов дыхания. С помощью метода наименьших квадратов определяются коэффициенты k_1 и k_2 . Для получения зависимости $\zeta = f(Re)$ использована указанная выше методика. Полученная графическая зависимость $\zeta = f(Re)$ приведена на рис. 1.

Подробное описание программно-аппаратной системы приведено в работах [17, 18, 19]. Система состоит из двух подсистем: аппаратной и программной. Аппаратная часть основана на измерительном модуле, в состав которого входят два датчика: датчик малых дифференциальных давлений и двунаправленный датчик расхода воздушного потока. Датчики данного типа обеспечивают высокую стабильность и производительность, а также имеют компактную конструкцию.

Отображаемый диапазон измерений дифференциального давления ± 1200 Па. Предел приведенной погрешности измерения дифференциального давления $\gamma_p = \pm 0,25\%$. Диапазон измерения

объемной скорости воздушного потока составляет ± 1200 см³/с. Предел относительной погрешности измерения скорости воздушного потока $\delta_p = \pm 3\%$. Особенностью предлагаемой технической реализации является устранение дополнительных потерь измерения дифференциального давления [20].

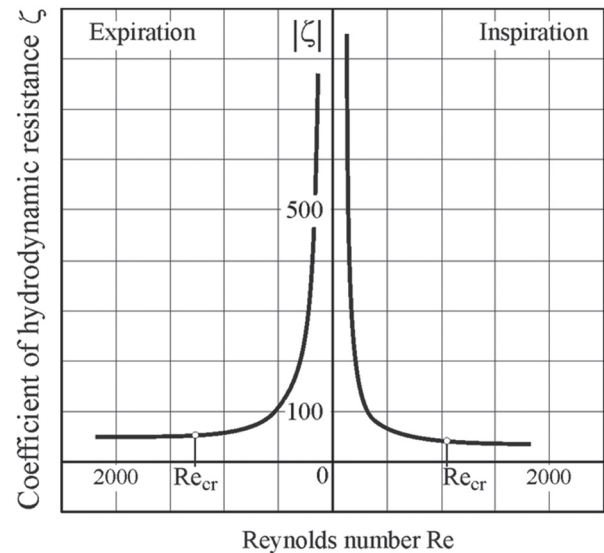


Рис. 1. Графическая зависимость коэффициента гидродинамического сопротивления от числа Рейнольдса

Риноманометрические данные регистрируются и анализируются с помощью программного обеспечения подсистемы (C #, SQLite, платформа «NET») в режиме реального времени. Частота опроса измерительных каналов 100 Гц. Подсистема программного обеспечения состоит из модуля записи входных данных, модуля анализа данных, модуля хранения данных. Система прошла сертификацию и получила свидетельство о государственной регистрации № 14777/2015 от 06.12. 2015. Подробную информацию о системе можно найти в работе [17].

2. Метод предварительной обработки риноманометрических данных

В статье разрабатывается метод предварительной обработки риноманометрических данных, который реализуется в программной подсистеме. С целью определения ошибочных измерений, предлагается использовать методику искусственного интеллекта, основанную на нейронных сетях.

На первом этапе метода выполняется подготовка исходных данных, которая включает фильтрацию риноманометрических данных, расчёт элементов матриц, нормализацию.

Исходный набор клинических данных состоит из множества риноманометрических сигналов, представленных на рис. 2. Рассмотрим графические зависимости на рис. 2, а и 2, б. Визуально сложно выявить существенные отличия, поэтому

предлагается использовать логарифмическое масштабирование с целью повышения качества визуализации данных.

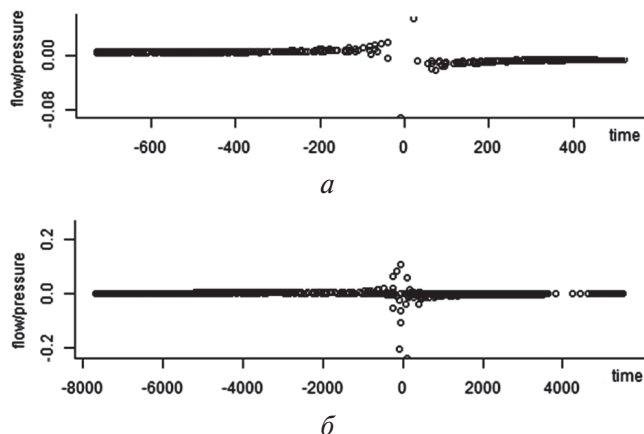


Рис. 2. Риноманометрические сигналы без обработки:
а – правильно выполненное измерение;
б – некорректное измерение

Подготовим данные для классификации. Пусть матрица m размером $N \times M$ содержит риноманометрические данные, каждый элемент имеет значение в интервале $0..K$. Начальные значения матрицы m состоят из риноманометрических сигналов Q и $\dot{p}$. На первом шаге необходимо выявить неверные значения, например «N/A» или «Inf». На втором шаге, необходимо вычислить каждый элемент матрицы по формуле:

$$m_{i,j}^* = \text{count}(\Delta p, Q) \cdot \begin{matrix} i=[\Delta p / \max(\Delta p) * N] \\ j=[Q / \max(Q) * M] \end{matrix} \quad (4)$$

На третьем шаге, необходимо определить максимальное значение K среди элементов матрицы m , нормализовать все элементы матрицы m относительно найденного максимального K .

$$m = \left[\ln \left| \frac{m^* \cdot K}{\max(m)} \right| \right] \quad (5)$$

Размерность матрицы m – N, M ограничена диапазоном измерений сигнала. Максимальная размерность матрицы 1200×1200 определяется точностью измерения расхода воздушного потока и дифференциального давления. Используя метод классификации Random Forest можем определить значения K при которых значения ошибок минимальны. В табл. 1 приведены данные зависимости ошибок от K , графически данные представлены на рис. 3.

Таблица 1

Зависимость ошибки от параметра K
для метода Random Forest

Данные \ K	32	64	128	256
ООВ, %	23	18,42	10,53	11,84
Test, %	10	5	10	25

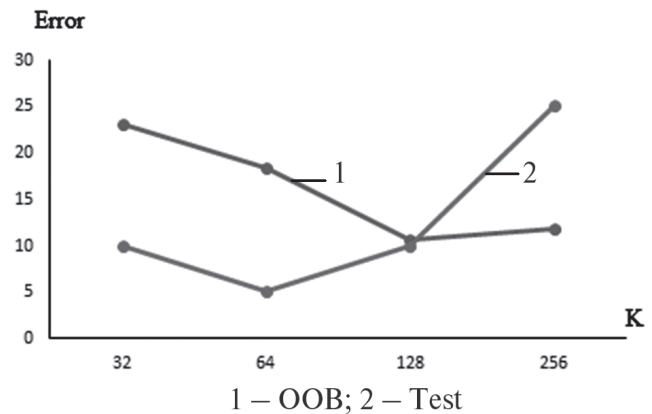


Рис. 3. Зависимость ошибки от K

На втором этапе предлагаемого метода выполняется визуализация зависимости $\zeta = f(Re)$. Коэффициент гидродинамического сопротивления рассчитывается по формуле:

$$\zeta = \ln \left| \frac{A}{Re} + B \right|, \quad (6)$$

где Re – число Рейнольдса, A и B – безразмерные константы, которые зависят от формы местного сопротивления. [16]. Таким образом, безразмерные константы A и B учитывают различные анатомические конфигурации носовых полостей..

Зависимость коэффициента гидродинамического сопротивления от числа Рейнольдса имеет существенное значение для диагностики носового дыхания функции патологии.

На третьем этапе, анализируются матрицы изображений зависимости $\zeta = f(Re)$, полученные на втором этапе, с целью удаления некорректных измерений.

Для классификации изображений предлагается использовать глубокие сверточные сети типа DCN (Deep convolutional networks). Эти модели основываются на работах [23, 24]. Данные искусственные нейронные сети состоят из слоев свертки, макс-пул слоев и нормализации, которые соединены одним полносвязным слоем. Мы предлагаем использовать структуру нейронной сети изображенную на рис. 4. Система классифицирует подготовленные наборы данных на два класса – «правильный» и «ошибочный». Данные классы относятся к двум ситуациям: когда протокол измерения риноманометрических данных был соблюден и ситуации когда протокол измерений был нарушен. Для обучения используется обучение с учителем. Примеры исходных данных двух классов без предварительной обработки изображены на рис. 4.

При использовании данной нейронной сети, достаточно важно сформировать вывод информации в наглядном виде. Для визуализации ключевых признаков выделенных глубокой искусственной нейронной сетью на разных уровнях

масштабирования графической информации, предлагается использовать метод, который основывается на [21]. Этот подход заключается в использовании многослойной сети деконволюции [22] для визуализации признаков в процессе обучения.

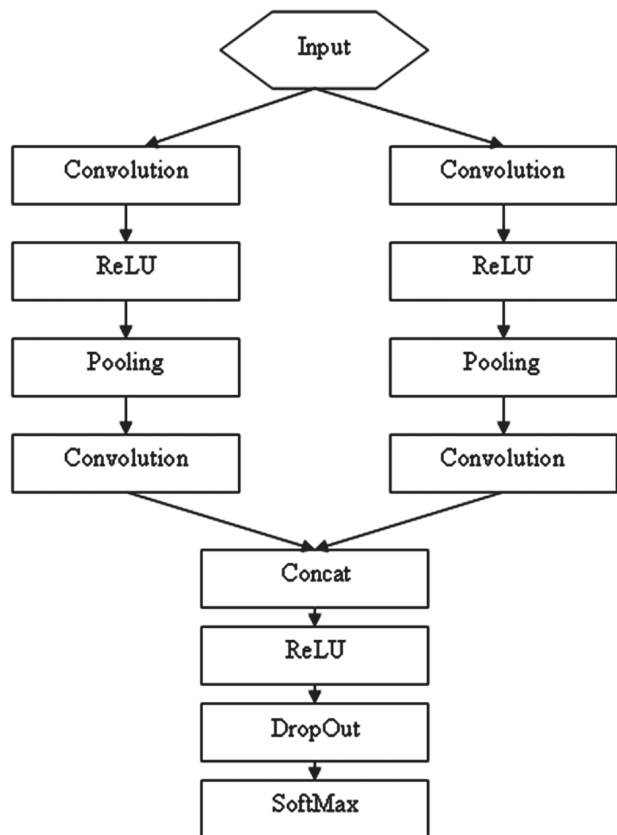


Рис. 4. Архитектура сверточной искусственной нейронной сети для классификации данных риноманометрии

Информация, полученная в результате деконволюции, может использоваться для создания карт признаков, которые могут быть интерпретированы специалистами в ринологии и проанализированы в будущих исследованиях.

3. Экспериментальные результаты

Результаты масштабирования после применения формулы (6) представлены на рис. 5. Откуда можно сделать вывод о визуально различимой разнице между правильными измерениями (5, а) и измерениями, полученными в результате нарушения протокола измерения (5, б).

На данном изображении четко выделяются две горизонтальные линии на фигуре 5, б. Они указывают на нарушение процедуры измерения. На рис. 5, а данные линии отсутствуют. Результаты предварительной обработки изображены на рис. 6.

Для проведения классификации были использованы три метода: Random Forest, Support Vector Machine (SVM) и глубокие конволюционные сверточные сети (DCNN). В результате проведенного

эксперимента, были получены результаты в виде соотношения правильно классифицированных измерений данных риноманометрии для пациентов с ринитом.

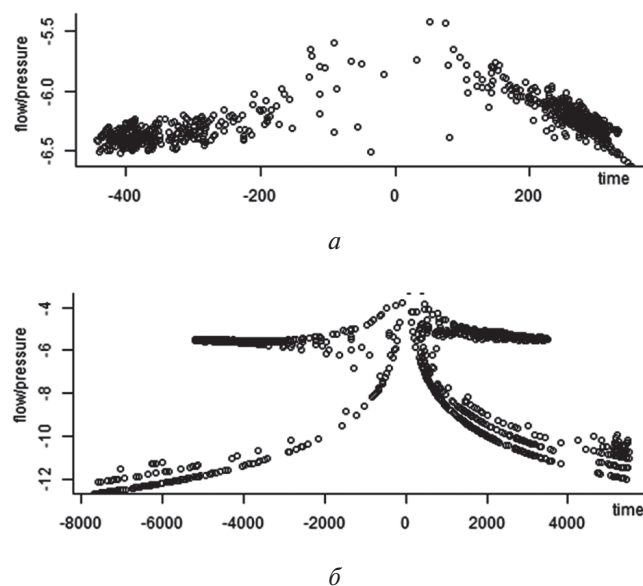


Рис. 5. Препроцессинг данных риноманометрии для выделения значимых признаков:
а — «правильно» проведенное измерение;
б — «ошибочно» проведенное измерение

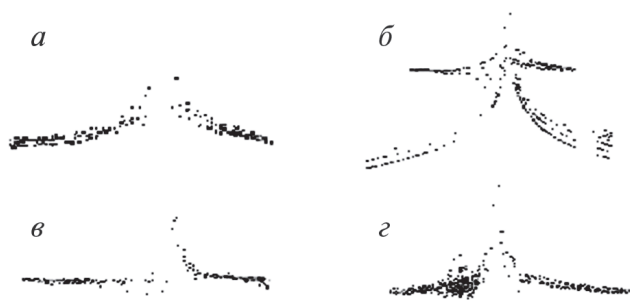


Рис. 6. Препроцессинг сигналов риноманометрии:
а — «правильное» измерение, ринит;
б — «ошибочное» измерение, ринит;
в — норма; г — патология

Результаты 89,4 % для Random Forest, 88,2 % для SVM и 90,1 % для предложенной структуры DCNN.

Выводы

В работе проведен анализ основных методов объективной оценки функции носового дыхания. Разработан метод предварительной обработки риноманометрических данных, который позволяет пользователю визуально оценить правильность проведенного измерения. Предложенный метод позволяет уменьшить количество хранимой информации в программно-аппаратной системе. Использование многослойной сети деконволюции делает возможным создание визуальных карт признаков. В проведенном исследовании был

продемонстрирован потенциал использования глубоких сверточных сетей, которые показали наибольшую точность (0,91%) среди методов, применяемых во время исследования. Также в результате обучения изменяются параметры предварительной обработки данных. Таким образом, данный компонент системы риноманометрических измерений позволяет уменьшить влияния человеческого фактора на результаты измерений.

Список литературы:

1. *Clement P. A.* Standardisation Committee on Objective Assessment of the Nasal Airway. Consensus report on acoustic rhinometry and rhinomanometry / P.A. Clement, F. Gordts // *Rhinology*. — 2005. — 43. — P. 169–179.
2. *Thulesius H. L.*, Rhinomanometry in clinical use. A tool in the septoplasty decision making process : doctoral dissertation, clinical sciences / H.L. Thulesius — 2012. — 67 p. H. L.
3. *Malm L.* Guidelines for nasal provocations with aspects on nasal patency, airflow, and airflow resistance / L. Malm, R. G. v. Wijk, C. Bachert // International Committee on Objective Assessment of the Nasal Airways, International Rhinologic Society, " *Rhinology*, vol. 38(1), pp. 1–6, March.
4. *Гарюк О.Г.* Риноманометрия. Сообщение 2: Современное состояние и перспективы / О.Г. Гарюк / *Ринология*, № 3, с. 32–45, 2013.
5. *Goras L.* Preprocessing method for improving ECG signal classification and compression validation / L. Goras, M. Fira // *PHYSICON 2009*, Catania, Italy, September, 1–September, 4 2009.
6. *Файнзильберг Л.С.* Адаптивное сглаживание шумов в информационных технологиях обработки физиологических сигналов / Л.С. Файнзильберг // *Математические машины и системы*, № 3, с. 96–104, 2002.
7. *Vogt K. A.* 4-Phase-Rhinomanometry (4PR) — basics and practice 2010 / K. Vogt, A. Jalowsky, W. Althaus, C. Cao, D. Han, W. Hasse, H. Hoffrichter, R. Mosges, J. Pallanch, K. Shah-Hosseini, K. Peksis, K. D. Wernecke, L. Zhang and P. Zaporoshenko // *Rhinology Suppl.* 21, pp. 1–50, 2010.
8. *Kern E. B.* Committee report on standardization of rhinomanometry // *Rhinology*, vol. 19(4), pp. 231–236, 1981.
9. *Clement P. A.* Committee report on standardization of rhinomanometry // *Rhinology*, vol. 22(3), pp. 151–155, 1984.
10. *Eichler J.* Power-Curves in Rhinomanometry Leistungskurven in der Rhinomanometrie // *Biomedizinische Technik*, Band 34, pp. 42–45, 1989.
11. *Chometon F.* Aerodynamics of nasal airways with application to obstruction / F. Chometon, P. Gillieron, J. Laurent et al. // *Proceedings of the 6th Triennial International Symposium on Fluid Control, Measurement and Visualization*, pp. 65–71, 2000.
12. *Rohrer F.* The flow resistance in the human respiratory tract // *European Journal of Physiology*, no. 162, pp. 225–295, 1915.
13. *Naito K.* Comparison of calculated nasal resistance from Rohrer's equation with measured resistance at delta P 150Pa / K. Naito, T. Mamiya, Y. Mishima, Y. Kondo, S. Miyata and S. Iwata // *Rhinology*, vol. 36(1), pp. 28–31, 1998.
14. *Solow B.* Nasal airflow characteristics in a normal sample / B. Solow, A. Sandham // *Eur J Orthod*, no. 13(1), pp. 1–6, Feb., 1991.
15. *Mlynski G.* Diagnosis of respiratory function of the nose / G. Mlynski, A. Beule // *Springer medizin verlag*, no. 56(1), pp. 81–99, 2008.
16. *Чмовж В.В.* Аэродинамика носовой полости человека / В.В. Чмовж, О.Г. Гарюк, А.С. Нечипоренко // *Матеріали ХХ Міжнародної науково-технічної конференції "Гідроаеромеханіка в інженерній практиці"*, Київ, 26–29 травня, — 2015, с. 70–72.
17. *Нечипоренко А.С.* Технические аспекты риноманометрии / Нечипоренко А.С. Восточно-европейский журнал передовых технологий, — 2013, № 4/9 (64), — с. 11–14.
18. *Ерохин А.Л.* Особенности измерения дифференциального давления при передней активной риноманометрии / А.Л. Ерохин, В.В. Чмовж, А.С. Нечипоренко, О.Г. Гарюк // *Вісник національного технічного університету "ХПІ", серія "Інформатика і моделювання"*, — 2014, № 62(1104), — с. 49–57.
19. *Yerokhin A. A.* Usage of F-transform to Finding Informative Parameters of Rhinomanometric Signals / A. Yerokhin, A. Nechiporenko, A. Babii, O. Turuta // *Proc. of the X International Scientific and Technical Conference "Computer Science and Information Technologies CSIT 2015"*, Lviv, 14–17 September, — 2015, с. 129–132.
20. *Нечипоренко А.С.* спосіб вимірювання диференційного тиску для оцінки носового дихання / А.С. Нечипоренко, О.Г. Гарюк, В.В. Чмовж, О. Б. Касьяненко // Патент України на винахід № 10785 від 25.02.2015.
21. *Zeiler M. D.* Visualizing and understanding convolutional networks / M. D. Zeiler, R. Fergus. // *CoRR*, abs/1311.2901, 2013.
22. *Zeiler M.D.* Adaptive deconvolutional networks for mid and high level feature learning / M. D. Zeiler, G. Taylor and R. Fergus // In: *ICCV*, 2011.
23. *LeCun Y.* Learning methods for generic object recognition with invariance to pose and lighting / Y. LeCun, F.J. Huang, and L. Bottou. // *Computer Vision and Pattern Recognition*, 2004. *CVPR 2004. Proceedings of the 2004 IEEE Computer Society Conference on*, vol. 2, pp. 11–97. IEEE, 2004.
24. *LeCun Y.* Convolutional networks and applications in vision. / Y. LeCun, K. Kavukcuoglu, and C. Faretat // *Circuits and Systems (ISCAS)*, *Proceedings of 2010 IEEE International Symposium on*, pp. 253–256. IEEE, 2010.

Поступила в редколлегию 22.11.2016

УДК 681.513



А.С. Свиридов, Ю.Ю. Завизиступ, О.Ф. Михаль
ХНУРЭ, г. Харьков, Украина, svyrydovartem@gmail.com

МЕТОДОЛОГИЧЕСКИЕ АСПЕКТЫ СЕГМЕНТНОЙ ОБРАБОТКИ КИРЛИАН-ОБЪЕКТОВ

Применительно к алгоритмизации обработки *Кирлиан-объектов* для медицинской диагностики, показана допустимость и методологическая целесообразность подхода в целом, определена общая структурная схема, сформулирована этапность комплексного решения. Как вариант реализации, предложено фрагментирование *Кирлиан-объекта* и использование пофрагментного фрактального описания. Ключевые элементы рассмотренных алгоритмических решений отмакетированы с демонстрацией их работоспособности.

РАСПОЗНАВАНИЕ ОБРАЗОВ, СЛАБАЯ СТРУКТУРИРОВАННОСТЬ, КИРЛИАН-ЭФФЕКТ, ФРАКТАЛЬНОСТЬ

Введение

Распознавание объектов окружающего мира есть базовое проявление *человеческого интеллекта* (ЧИ). Эволюционирование ЧИ включает использование *инструментальных средств*, которые на современном этапе сами становятся *интеллектуальными* (ИИС), поскольку не только «расширяют границы наблюдаемости» окружающего мира, но и реализуют предварительную обработку информации. Особый интерес в связи с тематикой применения ИИС представляют *слабо структурированные образы* (ССО) (объекты, изображения) [1], обработка (интерпретация) которых обычно (традиционно) трактуется как прерогатива традиционной (не компьютеризированной) части ЧИ. Применительно к этому классу задач, содержащими примерами являются *Кирлиан-объекты* [2], относящиеся к изучению живой природы (биология, медицина). Достаточно широкий класс *Кирлиан-объектов* может быть охарактеризован как кольцевые структуры, допускающие фрагментирование и пофрагментную фрактальную интерпретацию.

Цель настоящей работы – разработка элементов алгоритмической части для специализированных ИИС, применительно к распознаванию характерных фрагментов кольцевых ССО, имеющих место при интерпретации определённого класса *Кирлиан-объектов*.

1. Слабо структурированные объекты

Распознавание объектов окружающего мира (существенная часть которых может быть в той или иной степени отнесена к ССО), как базовый элемент деятельности ЧИ, предполагает выделение из общего информационного потока наиболее информационно ёмких элементов, с целью получения в дальнейшем, при последующей их обработке, определённых преимуществ от их использования. Существенная часть информации поступает к человеку на обработку в виде двумерных образов.

По своей сути, распознавание есть *предварительная обработка* информации. Полагается, что далее для распознанных образов будет найдено целевое применение. Распознавание включает выделение, согласно контексту обработки, отдельных элементов входного информационного потока, в соответствии с определёнными распознавательными признаками.

Эволюционирование ЧИ включает использование *инструментальных средств*. Первоначально это – раздельно средства наблюдения за объектом и средства манипулирования объектом [3]. Постепенно, по мере эволюционирования, их функции трансформируются и взаимно сближаются: реализуется манипулирование для лучшего наблюдения за объектом и наблюдение для лучшего манипулирования объектом. Слияние (объединение) этих функций, установление прямых и обратных связей между ними – знаменует собой образование ИИС. Интеллектуальность ИИС проявляется в том, что они не только меняют масштабы (пространственные и временные) и расширяют диапазоны (температурные, частотные и др.), но и включают интеллектуальные (алгоритмизированные компьютеризированные) элементы предварительной обработки информации. Этим ЧИ разгружается от рутинных процедур и, как следствие, может использовать высвободившиеся интеллектуальные ресурсы на решение задач *интерпретационного* уровня. В частности, ИИС могут «подхватывать» и эффективно реализовывать процедуры фильтрации, выделения из фона, сопоставления с набором распознавательных признаков, реализации распознавания, идентификации и др. Тогда ЧИ (собственно человек-специалист) может уже использовать результаты работы ИИС в сопоставлении с конкретным контекстом, в рамках определённой целевой ориентации.

Распознавание образов (изображений) актуально для множества сфер деятельности человека.

Различные формы (конфигурации объектов) задают различные требования к тому, каким образом и сколь точную информацию нужно получить из изображения. Одним из таких частных случаев является обработка ССО.

Слабое структурированные может быть охарактеризовано как форма группировки данных, не соответствующая строгой (детерминированной) структуре таблиц и отношений в рамках моделей типа реляционных баз данных. Впрочем, подобная характеристика является самоочевидной. В ходе познавательного процесса, объекты окружающего мира сначала вычленяются (идентифицируются, как отличающиеся от других), затем группируются (сопоставляются между собой), и только после этого строится некая структура (типа реляционной базы данных), являющаяся текущей версией модели окружающего мира. Далее, по ходу последующего познавательного процесса (накопления и упорядочения информации), версии сменяются, структура неоднократно переделывается и уточняется. Попутно, соответствующим образом модифицируется структура ИИС ЧИ [3]. Но всякий раз структура строится на основе познавательного процесса, а не познавательный процесс укладывается в готовую структуру, поскольку модель — вторична по отношению к реальному окружающему миру. «Инвертирование» процесса познания оказывается контрпродуктивным, чему имеется множество примеров в истории науки. Поэтому представляется более правильным характеризовать *слабое структурирование* как некоторый промежуточный (продвинутый, но ещё не завершённый) этап формирования научного знания.

Существует большое количество разных типов ССО. В отношении изображений, их можно объединить в классы изображений по конфигурационному принципу. Как отмечалось, «мощным источником» ССО являются объекты живой природы. Их изображения могут быть классифицированы (рассортированы) по видовым, морфологическим, функциональным и др. признакам; однако даже внутри каждого из таких классов сохраняется большое разнообразие. Наличие этого разнообразия является фактором продолжения структуризации модели. В качестве примеров ССО, только из области изучения устройства и принципов жизнедеятельности организма человека (биология, медицина), могут быть указаны, в частности, изображения сетчатки глаза, радужной оболочки глаза (ириодиагностика), отпечатки пальцев (дактилоскопия), а так же *Кирлиан-объекты*. Типовый пример — след контакта подушечки пальца — представлен на рис. 1.

2. Кирлиан-эффект

А. Эйнштейн высказал [4] следующую мысль о динамике развития науки. Новые парадигмы приходят на смену прежним не потому, что сторонники прежних теорий переучиваются или признают свою неправоту, а потому что они (прежние классики) просто постепенно вымирают. Так, Ньютоновская механика пришла на смену Аристотелевским представлениям, а затем Эйнштейновская теория относительности — на смену Ньютоновской механике. Легко видеть, что интервалы смены этих базовых концепций — достаточно велики. За это время успевают «сойти со сцены» не только «пророки» и «отцы-основатели», но и последующие сторонники примкнувших научных школ и направлений. Почему так происходит? Потому что *прижизненно*, в науке, как и во многих других областях человеческой деятельности, — «противоборствуют и бушуют» человеческие страсти и честолюбия. Должны проявиться противоречия, демонстрирующие несостоятельность прежней концепции; должны уйти прежние авторитеты, которые, вопреки противоречиям, поддерживали прежнюю концепцию; должна накопиться «критическая масса» фактов, подтверждающих новую концепцию. Как результат, — должна так или иначе реализоваться *правомочность* новых представлений.

В указанном контексте, показательна динамика развития представлений, связанных с *Кирлиан-эффектом* [5]. Разумеется, сам эффект несколько «мельче рангом», чем Эйнштейновская теория относительности. Тем не менее, по-видимому, он доставлял определённый дискомфорт «признанным авторитетам» в соответствующих областях. На текущий момент, мы находимся где-то в середине процесса становления: основная волна «борьбы со лженаучностью» концепции уже схлынула, но «по инерции» продолжают вяло «расклеиваться ярлыки». Показательна заметка в Большой Российской Энциклопедии (БРЭ), где в статье «Аура» [6] сказано:

«В парапсихологии — гипотетич. излучение «тонкой материи» вокруг всех одушевлённых, живых и неживых предметов, представляемое в виде некоего нефизич. силового поля (биополе, биорадиация, энергоинформационное поле и т. п.), свойства которого можно измерить электромагнитными приборами или сфотографировать (эффект С. Кирлиана и В. Кирлиан)».

В данной заметке тенденциозно выглядит само упоминание *Кирлиан-эффекта* — явления давно известного, подтверждённого, воспроизводимого, защищённого авторским свидетельством [5] — в связи с парапсихологией. Кроме того, автор явно

«не дружит» с логикой и философией. Ведь согласно даже «Ленинскому определению материи», если объект «...копируется, фотографируется...», то он *материален*, т.е. принадлежать к реальному *физическому* миру. Между тем, согласно заметке, свойства будто бы *нефизического* поля — фотографируются и измеряются электромагнитными приборами. Итак, в целом, речь идёт о довольно неуклюжей подмене понятий.

В связи с этим, — не в защиту парапсихологии, а исключительно для демаркации рассматриваемой предметной области, — в таблице 1 сформулированы *Положения*, инвариантные относительно инсинуаций вокруг *Кирлиан-эффекта*.

Таблица 1

Поз.	Формулировка
1	Имеется наблюдаемый эффект. Наличие и воспроизводимость эффекта — подтверждены [3]
2	Проявление эффекта (Поз. 1) объективно (аппаратурно) регистрируется на материальных носителях
3	После того, как результаты зафиксированы на материальном носителе (Поз. 2), они могут быть интерпретированы. Интерпретация результатов не является частью эффекта (Поз. 1)
4	Формализованные процедуры интерпретации (Поз. 3) в принципе (при желании) могут быть алгоритмизированы. Алгоритмы не являются частью эффекта (Поз. 1)
5	Принятие решений относительно применения интерпретационных результатов (Поз. 3), полученных с применением алгоритмов формализованных процедур (Поз. 4), не является частью эффекта (Поз. 1), интерпретации (Поз. 2) или алгоритмов формализованных процедур (Поз. 4)

Легко видеть, что основное содержание настоящей работы всецело касается Поз. 4, а область применения — медицинская диагностика — Поз. 5.

3. Применение в медицинской диагностике

Как отмечалось выше, *Кирлиан-объекты* интересны в плане исследования, в особенности в связи с изучением объектов живой природы. Живые сущности при применении *Кирлиан-эффекта*, ведут себя несколько иначе, чем объекты «мёртвой природы». *Кирлиан-объекты* живых сущностей более изменчивы и динамичны; что естественно, поскольку в них активно проходят сложные комплексы химических реакций. Как следствие, у живых сущностей постоянно в динамике находятся различные внешние проявления: температура, диэлектрическая проницаемость поверхности, собственные электромагнитные излучения, качество поверхности контакта биологического объекта с подложкой при фиксации изображения, а так же другие известные и не известные факторы. Эти внешние проявления неизвестным (не достаточно

исследованным) образом сказываются на картине *Кирлиан-эффекта*, что *метафорически* (в литературно-художественном смысле) может быть названо биополем, информационным полем, или как угодно ещё. Таким образом, картина столь многообразна, что нельзя отрицать, что в целом *Кирлиан-объекты* на настоящий момент «не достаточно изучены» в рамках стандартных представлений о классической науке.

(Попутно в скобках стметим, что «*недостаточная изученность*» — *естественное и неотъемлемое состояние истинной науки*. Так, следует напомнить: в продолжение смены крупномасштабных парадигм, в настоящее время глобальная космология, астрофизика и просто физика («святая-святых» человеческого научного понимания общемирового устройства) находятся в очередном глубоком кризисе, т.к. оказалось, что 80-85% массы Вселенной приходится на «чёрную материю», никак *в настоящее время* не наблюдаемую и аппаратно не регистрируемую. Таким образом, по степени обособанности «чёрная материя» мало чем отличается от «божьего промысла».)

В связи с отмеченной «недостаточной изученностью», *Кирлиан-объекты* порой «обвешиваются ярлыками» типа экстрасенсорности или антинаучности. По данному вопросу выше мы «провели демаркационные линии» (Таб. 1). Легко видеть, что любое «обклеивание» *Кирлиан-эффекта* «ярлыками» может касаться только *интерпретации* — Поз. 5, Таб. 1. А так же само «обклеивание» является по существу *интерпретацией*. А *интерпретация* — дело субъективное и глубоко личное. Что же касается содержания настоящей нашей работы, то оно всецело относится к Поз. 4, Таб. 1.

Отдельного рассмотрения требует вопрос о *принятии решений* относительно применения интерпретационных результатов (Поз. 5, Таб. 1). Более конкретно, речь может идти о правомочности использования *Кирлиан-объектов* в медицине, в частности, в медицинской диагностике.

Медицина, как область человеческого знания, имеет «две ипостаси»: *научную* и *прикладную*. Различие состоит в следующем.

Научная медицина выглядит как стандартная классическая наука. Её цель — накопление знаний и построение теорий, т.е. *познание* одного из объектов окружающего мира — устройства, принципов действия и закономерностей эффективного функционирования *Человека*. Накопление и обработка информации в *научной* медицине идёт согласно канонам классической науки. Научное знание должно быть объективным. Для этого применяется стандартный познавательный процесс: наблюдение, гипотеза, эксперимент, обработка

результатов, обобщение, построение теории, верификация, корректировка теории, и т.д.

Прикладная медицина опирается на результаты *научной* медицины, но цель её другая: максимальная помощь конкретному больному. Этот конкретный больной из *прикладной* медицины есть единичная реализация *Человека* из *научной* медицины. В стандартной науке субъективность недопустима. При этом объективность достигается, в частности, максимальным накоплением данных (статистики) и статистической обработкой результатов. В *прикладной* медицине статистическая достоверность выборки просто не достижима, поскольку конкретный больной существует в единственном экземпляре. Поэтому, с точки зрения *прикладной* медицины, функция набора статистики (для нужд *научной* медицины) — является уже вторичной; а первично — принятие текущих (по возможности, максимально эффективных) решений в интересах конкретного больного.

Следует отметить, что *прикладная* медицина, в основном, никого насильственно не лечит. Больной обращается и желает лечиться, а врач (дагност) консультирует и рекомендует. Таким образом (и это есть тонкие моменты, которые необходимо принимать во внимание), в *прикладной* медицине есть место для субъективного врачебного опыта и интуиции. Кстати, интуиция, по существу, есть принятие во внимание особенностей, обстоятельств и закономерностей, которые входят в человеческий опыт (индивидуальный опыт специалиста), но, возможно, не формулируются в явном виде. Таким образом, в *прикладной* медицине (в частности, в диагностике) применимо практически всё, что по мнению специалиста может пойти на пользу больного, а не исключительно только то, что имеет своё научно-теоретическое обоснование.

В связи с этим, использование *Кирлиан-эффекта* в медицинской диагностике допустимо, оправдано и не противоречит принципам *прикладной* медицины, если имеется специалист дагност, который имеет опыт интерпретации *Кирлиан-объектов*.

4. Анализ Кирлиан-объектов

Метод получения *Кирлиан-объектов* (изображений проявления *Кирлиан-эффекта*) [5], — газоразрядная визуализация (ГРВ). В литературе имеется так же термин ГРВ-грамма. Помимо ГРВ, существуют некоторые другие варианты этого метода. В основе его лежат общие принципы электрофотографии. Технические подробности целесообразно опустить, так как они не относятся к рассматриваемым в настоящей работе вопросам обработки информации.

При обработке и изучении изображений биологических объектов, полученных с использованием *Кирлиан-эффекта*, возникает необходимость их последующей классификации. На рис. 1 представлен типовой вид *Кирлиан-объекта* отпечатка (следа контакта) подушечки пальца. *Кирлиан-объекты* такой (подобной) конфигурации (обладающие кольцевой структурой) рассматриваются нами на данном этапе. Тип объекта рассмотрения определяется (изначально задаётся) простым прикладным применением (областью перспективного применения) — медицинской диагностикой отдельных видов заболеваний (не уточняется, так как выходит за пределы предмета разработки методов формализации изображений и алгоритмизации процедуры обработки) по *Кирлиан-объектам* отпечатков (следов контакта) пальцев пациента.

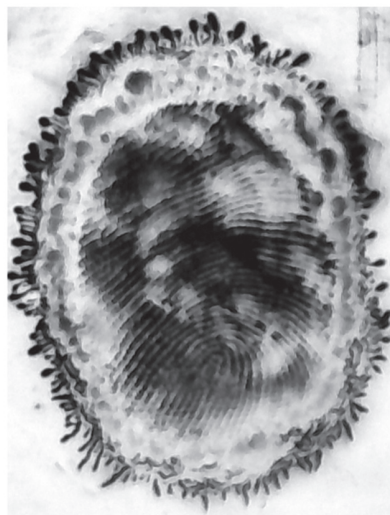


Рис. 1 Кирлиан-объект — след подушечки пальца

На фотографии (рис. 1) видны (могут быть идентифицированы) по крайней мере три области. *Центральная область* — дактилоскопический узор — отпечаток пальца. Это область непосредственного контакта (касания) с поверхностью. *Промежуточная область* — окаймляет центральную, является более светлой и содержит, судя по виду, некоторые неоднородности. Характер этих неоднородностей, а так же ширина промежуточной области, по-видимому, специфичны в зависимости от конкретного применяемого метода получения изображения (*Кирлиан-объекта*). *Периферийная область* — всплески и выбросы, образующие некоторый «узор», существенно более тёмного цвета по сравнению с промежуточной областью. Периферийная область и есть часть *Кирлиан-объекта*, несущая полезную информацию.

Полезное изображение (периферийная область) имеет форму, близкую к кольцевой. Информационную ценность имеет конфигурация

отдельных всплесков. Интенсивности и распределения всплесков по фрагментам (секторам) кольцевой структуры могут описывать отдельные классы состояний, и могут интерпретироваться как соотносящиеся с определёнными свойствами исследуемого объекта. Исходя из этого, целесообразна следующая структура обработки *Кирлиан-объектов*.

Шаг 1. Демаркация границ раздела центральной, промежуточной и периферийной областей *Кирлиан-объекта*.

Шаг 2. Аппроксимация демаркационных контуров стандартными геометрическими кривыми. Определение параметров этих кривых и характера их изменения для приведения их к окружностям.

Шаг 3. Удаление (очистка за ненадобностью) центральной и промежуточной областей.

Шаг 4. Деформация, масштабирование и приведение периферийной области к стандартной круговой конфигурации.

Шаг 5. Разбиение стандартной (приведённой) круговой конфигурации периферийной области на требуемое число секторов (равных или различных), согласно ориентации объекта, выявленной (уточнённой) при проведении процедуры приведения, и с учётом требований интерпретации, применительно к данной предметной области.

Шаг 6. Посегментная формализация — характеристика «узора» периферийной области, с использованием детерминированного алгоритма обработки.

Шаг 7. Посегментная фиксация параметров формализации

Шаг 8. Представление посегментных наборов параметров в виде, удобном для интерпретации.

Данная 8-Шаговая структура обработки допускает различные варианты подходов при реализации отдельных Шагов. В связи с этим, буквально каждый из Шагов требует отдельного детального рассмотрения и сопоставления вариантов. Рассмотрим далее более детально (но, возможно, не полно) некоторые (не все) из перечисленных Шагов.

5. Фрактальная размерность изображения

Может быть предложен подход, основанный на *фрактальном* представлении секторов периферийной области *Кирлиан-объекта*. Он пригоден на этапе подготовки данных для интерпретации (Шаг 6). Отдельные элементы его алгоритмической реализации применимы так же при «черновой» (предварительной) подготовке информации (Шаги 1 - 3); и частично, для скорректированной визуализации результатов для рассмотрения их специалистом-интерпретатором (Шаг 8).

В рамках рассматриваемого *фрактального* представления — принадлежность объектов изображения к тому или иному классу определяется путём посегментного сравнения эталонных и текущих *фрактальных* размерностей полученного изображения.

Так как наше изображение имеет не четкие границы, для повышения качества распознавания целесообразно применять размерность Минковского [7]. Последняя — является одним из способов задания *фрактальной* размерности ограниченного множества в метрическом пространстве, и определяется как:

$$D = \lim_{\epsilon \rightarrow 0} \frac{\log N(\epsilon)}{\log \frac{1}{\epsilon}} = \lim_{\epsilon \rightarrow 0} \frac{\log N(\epsilon)}{-\log \epsilon}, \quad (1)$$

где: $N(\epsilon)$ — минимальное число множеств диаметра ϵ , которыми можно покрыть исходное множество.

Если предел не существует, то рассматриваются верхний и нижний пределы и говорится, соответственно, о верхней и нижней размерностях Минковского, которые тесно связаны с размерностью Хаусдорфа [8]. Рассматриваемая задача по распознаванию в размерности Минковского выглядит как двумерный случай; хотя аналогичное определение распространяется и на n -мерный случай.

Рассмотрим некоторое ограниченное множество в метрическом пространстве. Исходное изображение — это черно-белая картинка. Добавим к ней равномерную сетку с шагом ϵ , и закрасим те ячейки сетки, которые содержат хотя бы один элемент рассматриваемого множества.

Следующим шагом, начнём уменьшать размер ячеек (значение ϵ), а размерность Минковского будет вычисляться согласно (1), исследуя скорость изменения отношения логарифмов. Алгоритм выводится следующим образом. Обозначим приближенное значение размерности Минковского как D_{bc} . Обозначим определение этой размерности, убрав предел. Его мы будем имитировать в итерациях, в которых будет изменяться размер ячеек.

$$D_{bc} = \frac{\log N(\epsilon)}{\log \frac{1}{\epsilon}} \Leftrightarrow D_{bc} \log \frac{1}{\epsilon} =$$

$$\log N(\epsilon) \Leftrightarrow D_{bc} \log \frac{1}{\epsilon} - \log N(\epsilon) = 0.$$

При условии, что мы зафиксировали размеры ячеек ϵ и рассматриваем D_{bc} как неизвестное, получаем, что приведенное выражение является формулой линии.

Далее, в продолжение рассматриваемому, подлежат уточнению границы нашего изображения при условии отслеживания контуров изображения.

6. Отслеживание контуров изображения

Отслеживание контуров даёт набор пригодных для обработки данных. На базе математической морфологии существуют алгоритмы оконтуривания, но обычно используются гибридные алгоритмы или алгоритмы в связке.

Полученные границы достаточно просто преобразуются в контуры. Например, алгоритм Кэнни [9] оконтуривает автоматически, для остальных алгоритмов требуется дополнительная бинаризация. Получить контур для бинарного алгоритма можно, например, алгоритмом «жука» [10] (рис. 2).

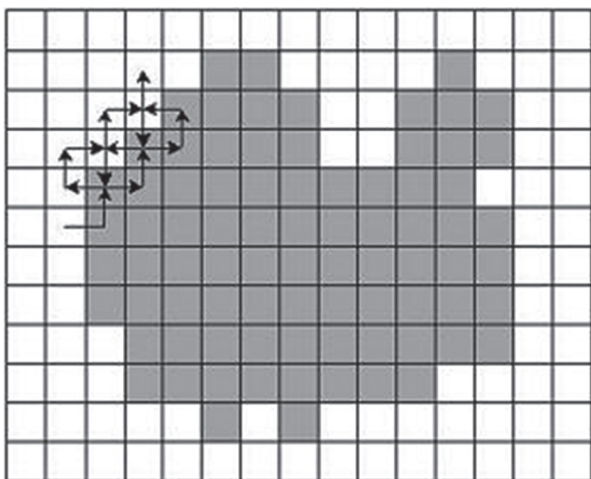


Рис. 2. Схема работы отслеживающего алгоритма «жука» [10]

Алгоритм работает следующим образом. Элементы («клеточки») поля последовательно рассматриваются (например по строке или по столбцу), пока «жук» не наткнется на объект. Далее «жук» захватывает объект и продвигается по его контуру. При этом, если элемент (текущее рассматриваемое поле) белый, «жук» поворачивается направо, иначе — налево. Поле, в котором «жук» побывал, помечается как *пройденное*. Процедура повторяется до тех пор, пока «жук» не вернется в исходную точку контура — наткнется на *пройденное* поле. Координаты точек перехода с чёрного на белое и с белого на чёрное — описывают границу объекта.

Результатом работы процедур выделения контуров является *контурный препарат*. Он представляет собой множество не связанных друг с другом краевых точек. Для дальнейшей работы с контурами необходимо из множества выделенных краевых точек сформировать кривые (ломанные).

Задача прослеживания контуров (edge following) — одна из ключевых, применительно к работе с определением границ, поэтому она так же является объектом рассмотрения. Метод базируется на достаточно продуктивном предположении, что точки, принадлежащие одному контуру, должны

иметь близкие значения модуля и направления вектора градиента. Иначе говоря, исключается возможность разрывности. В методе рассматривается окрестность точки (i, j) размером $M \times M$ (обычно используют окрестность 3×3), и в каждой точке (k, ℓ) окрестности проверяются условия:

$$|G_{i,j} - G_{k,\ell}| \leq \Delta G, \quad |\alpha_{i,j} - \alpha_{k,\ell}| \leq \Delta \alpha, \quad (2)$$

где (i, j) — центральная точка окрестности; G — модуль градиента; α — направление градиента в точке; ΔG — предельное значение расхождения модулей градиента в точках (i, j) и (k, ℓ) ; $\Delta \alpha$ — предельное значение расхождения направлений векторов градиента в точках (i, j) и (k, ℓ) . Если в точке (k, ℓ) выполняются условия (2), считается, что пара точек принадлежит контуру.

Для упрощенного вычисления направления края, весь диапазон возможных значений $0, \dots, 360^\circ$ разбивается на 8 направлений (секторов). Каждое направление отличается от соседнего на 45° . При этом поиск точек, принадлежащих одному контуру, следует проводить среди точек соседних секторов, имеющих расхождения градиентов меньше заданного порога.

Результатом выполнения *процедуры прослеживания* является дискретное представление контуров, при котором каждый контур определяется множеством точек, из которых он состоит.

Полученный *контурный препарат* в дискретном представлении далее подвергается анализу на предмет выделения на нем *точек ветвления* (точек соединения кривых). Наличие таких точек свидетельствует о сложной геометрической структуре объекта, что в свою очередь существенно затрудняет формальное описание и сам процесс распознавания объектов. Выделение *точек ветвления* упрощает структуру объекта путем разбиения единого контура на множество кривых.

Выявление четкого контура изображения позволяет перейти к внутренней составляющей изображения и применить *эталонные матрицы*.

7. Применение эталонных матриц

Для распознавания, собственно, изображения необходимо иметь различные эталонные части. Таким образом, надлежит реализовать набор эталонных образов в виде некоторых матриц, стандартной конфигурации, например, 3×3 . Такой набор описывает совокупность возможных состояний. Как вариант, набор может выглядеть согласно рис. 3. Чёрным и белым обозначены поля с детерминированной раскраской, серым — поля, раскраска которых безразлична. То есть, серым обозначены поля (пиксели), цвет которых не имеет значения: может быть либо чёрным, либо белым.

Восемь первых шаблонов (первая и вторая строки рис. 3), являются основной частью. Четыре остальных шаблона (третья строка рис. 3) – предназначены для устранения «шума». При этом, эти шаблоны могут быть так же повернуты на 90° , 180° и 270° градусов. То есть, в действительности «антишумовых» шаблонов имеется 13 штук. В варианте алгоритмической реализации с применением поворотов матриц «шумовых» эталонных образов, соответствующие совпадения ищутся повторными обходами *контурного препарата*.

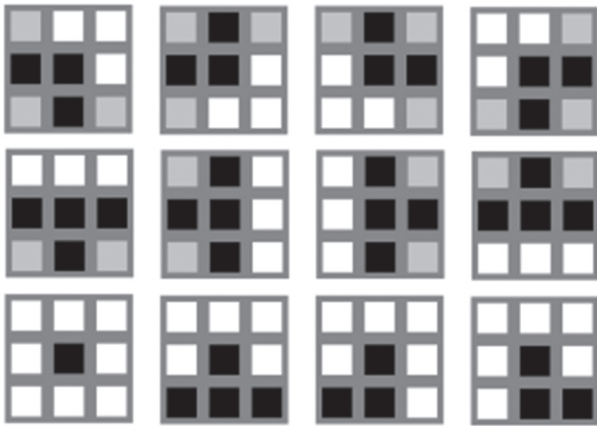


Рис. 3. Вариант реализации набора матриц эталонных образов

Также следует иметь ввиду что данный набор матриц не может покрыть все необходимые варианты сравнения. Тем самым принимается, что данный набор (рис. 3) является минимальным эталоном. Остальные сегменты должны быть реализованы (охвачены процессом распознавания фрагментов) в виде сочетаний этих эталонных матриц. С учётом этого, может быть реализовано покрытие всех существующих вариантов. В дальнейшем, может быть получена некоторая база данных образцов для сравнения, пригодная для использования при распознавании фрагментов контуров, выделяемых (отслеживаемых) согласно описанному выше, в п. 6.

Выводы

Рассмотрен структурно-методологический план алгоритмизации обработки *Кирлиан-объектов* для частного класса задач медицинской диагностики.

Показана допустимость и целесообразность самого подхода на основе *Кирлиан-эффекта*; очерчена граница реализуемости его возможной формализации. Сформулированы пункты (8 Шагов), определяющие этапность комплексного решения задачи, являющиеся одновременно крупномасштабной структурой системы (синтезируемого *интеллектуального инструментального средства*) в целом. Рассмотрены варианты реализации на алгоритмическом уровне, в числе которых предложено фрагментирование *Кирлиан-объекта* и использование пофрагментного фрактального описания. Ключевые элементы рассмотренных алгоритмических решений отмакетированы с демонстрацией их работоспособности. Намечена методологическая перспектива дальнейших исследований алгоритмов прикладной обработки *Кирлиан-объекта*, включающая, в частности, сопоставление фрагментов контура с базой эталонных (априорно заданных или ранее выделенных) фрагментов.

Список литературы:

1. Свиридов А.С. Метод обработки слабоструктурированных изображений кольцевой формы. // Проблемы информатизации. Тезисы докладов 4-й Международной н.-т. конф. – 2016 – с. 36.
2. Завизиступ Ю.Ю., Михаль О.Ф., Свиридов А.С. Метод обработки и классификации ГРВ подобных изображений. // Проблемы информатизации. Тезисы докладов 4-й Международной н.-т. конф. – 2016. – с. 24.
3. Михаль О.Ф. Эволюционирование мультиагентной системы как воспроизведение процесса формирования индивидуального человеческого интеллекта. // Тезисы докладов 1-й Международной научно-технической конференции «Полиграфические, мультимедийные и web-технологии» (PMW-2016). Т. 1. – Харьков: ХНУРЭ, 2016. – с. 73-74.
4. Эйнштейн А., Инфельд Л. Эволюция физики. / В кн. Эйнштейн А. Собрание научных трудов. Т. 4. М.: Изд. «Наука» – 1967, с. 357-543.
5. Кирлиан С.Д. Авт. свид. №106401, кл. G03B 41/00, 1949.
6. БРЭ. Аура – <http://bigenc.ru/medicine/text/1840161>.
7. Кроновер Р.М. Фракталы и хаос в динамических системах. Основы теории. М., 2000. 352 с.
8. Мандельброт Б. Фрактальная геометрия природы. – Москва: Инст. комп. исслед., 2002, 656 с.
9. Canny J. IEEE transactions on pattern analysis and machine intelligence, vol. pami-8, no. 6, november 1986.
10. Прэтт У. Цифровая обработка изображений. – М.: Мир, 1982. – Кн. 2 – 480 с.

Поступила в редакцию 24.11.2016

УДК 681.513



О.Ф. Михаль

ХНУРЭ, г. Харьков, Украина, oleg.mikhal@gmail.com

ЭВОЛЮЦИОНИРОВАНИЕ МУЛЬТИАГЕНТНОЙ СИСТЕМЫ КАК АНАЛОГ ФОРМИРОВАНИЯ ИНДИВИДУАЛЬНОГО ЧЕЛОВЕЧЕСКОГО ИНТЕЛЛЕКТА

Эволюционирование мультиагентной системы, как инструмента усиления человеческого интеллекта, сопоставлено с этапностью развития индивидуального человеческого интеллекта. Факт установления соответствия интересен в плане прогнозирования развития конкретных мультиагентных приложений.

МУЛЬТИАГЕНТНАЯ СИСТЕМА, ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ

Введение

Концепция *мультиагентной системы* (МАС) доминирует в проблематике *искусственного интеллекта* (ИИ) в силу её *антропоморфности*: определённые программные и аппаратные сущности (*агенты*) наделяются свойствами (атрибутами поведения), подобными тем, которые присущи *человеческому интеллекту* (ЧИ) [1]. «Очеловечение» характера технических сущностей придаёт им обозримость, удобство обслуживания и модернизации, а также удобство в непосредственной работе с ними (штатной эксплуатации) на уровне ЧИ. Кроме того, в субъективном плане — появляется наглядность и предсказуемость. По существу, ЧИ проецируется на технические сущности и воспринимает своё собственное отражение в них, как ИИ. Сложность (обозримость) *агента*, как интеллектуальной сущности, которая приобретает очертания *усилителя человеческого интеллекта* (УЧИ) [2], — становится соизмеримой (сопоставимой) с соответствующими человеческими показателями. С другой стороны, индивидуальный ЧИ формируется согласно достаточно простой многоуровневой схеме, частично рассмотренной в [3, 4]. Представляет интерес изучение соответствия между эволюционированием систем ИИ (включая становление, развитие и «размножение» (разрастание структуры) МАС) и динамикой формирования (развития) *индивидуального ЧИ*, то есть интеллекта, присущего конкретному человеку.

Цель настоящей работы — сопоставление (демонстрация преемственности) эволюционирования МАС, как УЧИ, в сравнении с этапностью становления (развития) индивидуального (личностного) ЧИ. Факт установления указанного соответствия, помимо общеметодологического значения, интересен в плане прогнозирования и прототипирования дальнейших шагов (выбора направлений) в развитии МАС в конкретных приложениях.

1. Об интеллекте

Везде, во все времена, практически в любой области человеческих знаний, разделение объектов и

явлений на классы, проведение «демаркационных линий» между отдельными классами, в том числе установление терминологии, происходит «начиная с середины», с последующим движением «внутрь» («в глубь») и «во вне» («в ширь»). История науки показывает, что обычно в любой области знаний рано или поздно оказывается, что принятая «классификация», «исходная система аксиом» или «терминологическая база» *безосновательна* при движении «в глубь» («к корням», «к истокам») или *не состоятельна* при развитии «в ширь» («к обобщениям», «к объединению со смежными областями», «к развитию на стыках наук»). Подобное имеет место и в связи с науками об интеллекте и человеческом сознании.

Термин *интеллект* первоначально появился для обозначения «разумности» человека. Позднее оказалось, что поведение животных часто оказывается так же «не лишено разумности». Так возник термин *интеллект животных*, в отличие от ЧИ. Появление средств *вычислительной техники* (ВТ) первоначально не было воспринято как формирование УЧИ [2]. То есть, ЧИ первоначально «не разглядел» в средствах ВТ своего собственного отражения, а увидел только *интеллектуальную сущность* и назвал её ИИ. Такова тенденция формирования понятия (этимологический план). В содержательном же отношении, то что обозначено словом *интеллект* (разумность), есть просто способность приспосабливаться к условиям внешнего мира — атрибут *живого* — определённая фаза развития жизни. Поэтому чёткой грани наличия-отсутствия *интеллекта* попросту не существует. Таким образом, термину *интеллект* присуща «определённая неопределённость», что может быть допустимо для абстрактных понятий или групповых явлений, но в целом — мало продуктивна. Приводимые далее концептуальные соображения относительно развития МАС и формирования уровней ЧИ предположительно могут служить методологическим каркасом для придания понятию *интеллекта* большей определённости (для формализации).

2. Эволюционирование мультиагентной системы

Ключевая особенность МАС – антропоморфность. Причина этого кроется в изначальном назначении *агента* – быть УЧИ [2]. Целесообразно, чтобы размер, сложность (обозримость) и структурная организация *агента*, как интеллектуального образования (сущности), были соизмеримы со средними человеческими показателями. Человек имеет определённые интеллектуальные и физические ограничения, сложившиеся в ходе биологической эволюции. Так, существенными интеллектуальными ограничениями человека являются принципиальная однозадачность (отсутствие долговременного полноценного бесстрессового распараллеливания внимания) и «магическое число» 7 ± 2 (число ячеек «оперативной памяти») [5]. Решение более сложных задач, выходящих за границы этих ограничений, реализуется задействованием социального аспекта (объединением интеллектуальных усилий нескольких человек), либо подключением УЧИ (в частности, средств ВТ). Антропоморфная система (объект, сущность) соответствует ограничениям, присущим человеку, не антропоморфная – превышает средние человеческие показатели.

Не антропоморфный *агент*, разумеется, может быть разработан, но для его создания, сопровождения и эксплуатации должны затрачиваться «сверхчеловеческие» усилия (ресурсы, коллективы разработчиков, группы сопровождения и др.) Антропоморфный *агент* (аппаратный или программный) – не предполагает «сверхчеловеческих» затрат: удобен в разработке и эксплуатации одним человеком. В особенности, принцип антропоморфности должен соблюдаться в «человеко-агентном» интерфейсе. Не антропоморфный интерфейс не может быть «дружественным».

Соблюдение принципа антропоморфности интерфейса приводит к иерархичности системы агентов. Рассмотрим динамику эволюционирования (поэтапного становления и развития) МАС, иллюстрируемую серией структурных схем рис. 1 (а–с).

Начальный этап (рис. 1, а) соответствует примитивному уровню взаимодействия человека (оператора) Ч с объектом O_1 (рис. 1, а). Индекс здесь и дальше – обозначает номер уровня в многоуровневой (иерархической) МАС. Стрелки обозначают направление потока информации ($Ч \leftarrow O_1$) или управляющее (физическое) воздействие ($Ч \rightarrow O_1$). По мере изучения объекта O_1 человеком, характер взаимодействия усложняется. Объект становится «более понятен человеку»: человек изобретает средства для лучшего наблюдения объекта, а также орудия труда для более удачного (точного) воздействия на объект.

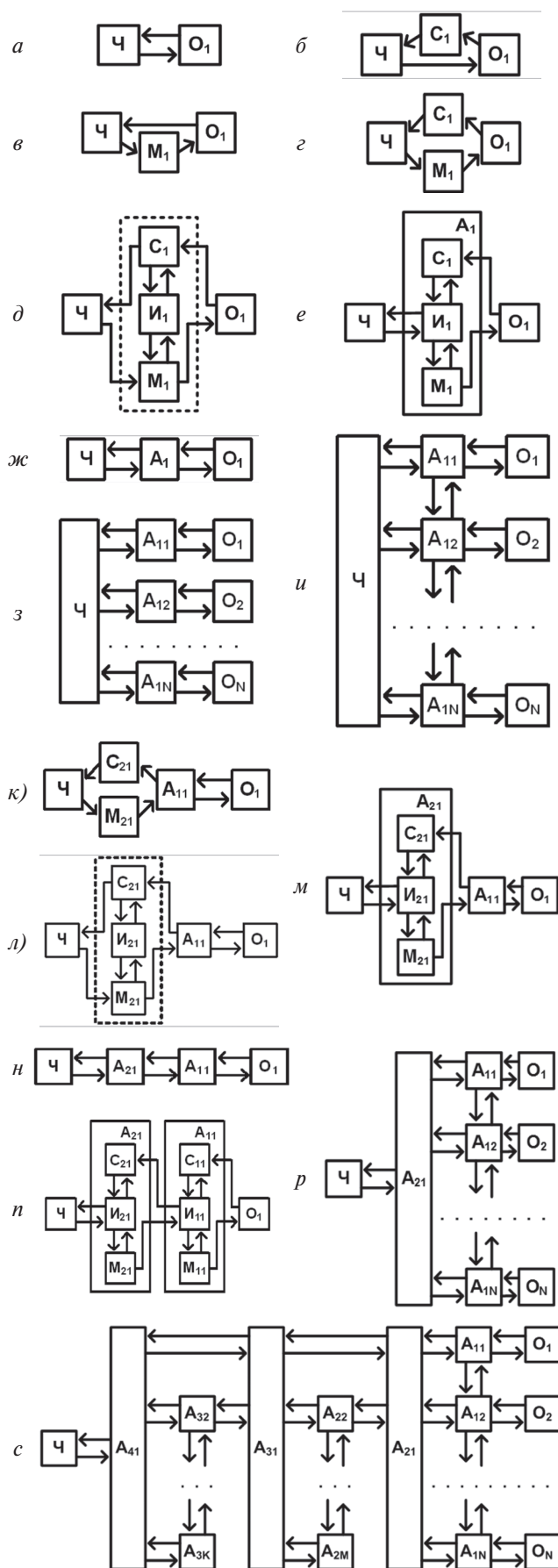


Рис. 1 (а–с). Динамика становления и развития агента, как функциональной и интеллектуальной сущности

Таким образом, Ч воспринимает информацию (рис. 1, б) или осуществляет воздействие (рис. 1 в) или и то и другое (рис. 1, з) — через промежуточное оборудование: сенсор C_1 и манипулятор M_1 . Сенсор, возможно, фильтрует, селективирует или преобразовывает информацию; манипулятор, возможно, усиливает, дозирует величину или более точно позиционирует воздействие на объект.

C_1 и M_1 развиваются и совершенствуются. Качественное изменение проявляется на последующем этапе, когда C_1 и M_1 достигают такого уровня сложности, что между ними устанавливается связь ($C_1 \rightleftharpoons M_1$): появление информационных потоков (показаны на рис. 1, д стрелками) и блока обработки информации — интерпретатора I_1 . В простом варианте функция I_1 — согласование форматов данных; в более развитом — преобразование информации, вплоть до выработки решений. Структура рис. 1, д является «*преагентной*»: объект O_1 наблюдается и обрабатывается человеком Ч, связка $C_1 \rightleftharpoons I_1 \rightleftharpoons M_1$ является вспомогательной конструкцией — «аватаром», управляемым Ч. «Агентность» (относительная автономность и самостоятельность в принятии решений) проявляется только при подключении Ч к I_1 (рис. 1, е). Взаимодействие $Ч \rightleftharpoons I_1$ сводится в конце концов к регулированию параметров автономно работающей интеллектуальной сущности (*агента*) A_1 .

Качественный характер перехода состояний «преагентность» (рис. 1, д) → «агентность» (рис. 1, е) состоит в том, что для Ч появляется возможно, работать не с единственным интеллектуальным агентом (рис. 1, ж), а параллельно с N агентами: $Ч \rightleftharpoons (A_{11}, A_{12}, \dots, A_{1N})$ (рис. 1, з). Здесь первый индекс обозначает уровень *агента*, второй — его номер. Распараллеливание означает снижение рутинной управленческой нагрузки на ЧИ, преобладающее переключение ЧИ на функции регулирования, контроля и принятия «стратегических решений». В целом это соответствует идеологии УЧИ. Очередной шаг в этом направлении — интеграция *межагентских* связей $A_{11} \rightleftharpoons A_{12} \rightleftharpoons \dots \rightleftharpoons A_{1N}$ — иллюстрируется рис. 1, и. Координация действий однотипных *агентов*, выполняющих сходные функции, позволяет дополнительно разгрузить ЧИ, что равносильно наращиванию числа N активных *агентов*.

Дальнейшее наращивание МАС в направлении *многоуровневости* происходит в связи со следующими представлениями. В ситуации $Ч \rightleftharpoons A_1 \rightleftharpoons O_1$ (рис. 1, ж) человек с O_1 непосредственно не взаимодействует, хотя то, с чем Ч взаимодействует, для него оказывается также объектом. Поэтому в методологическом плане для Ч ситуация ($Ч \rightleftharpoons A_1$) не отличается от ($Ч \rightleftharpoons O_1$). Тогда, применяя

«эволюционный процесс», аналогичный описанному (рис. 1, а–е), получаем формирование блоков второго уровня: сенсора C_2 и манипулятора M_2 (рис. 1, к); затем включение между ними интерпретатора второго уровня I_2 (рис. 1, л); и наконец полномасштабное формирование агента второго уровня A_2 (рис. 1, м). Результат — получение иерархической системы $Ч \rightleftharpoons A_2 \rightleftharpoons A_1 \rightleftharpoons O$ с двухуровневой упорядоченностью агентов (рис. 1, н, н). При этом Ч ещё на один слой управления отдалается от реального физического объекта O_1 .

Очевидны последующие «ветви» эволюции: Одна из них — множественное «агент-агентное» управление. На рис. 1, р *агент* второго уровня управляет N агентами первого уровня: $Ч \rightleftharpoons A_2 \rightleftharpoons (A_{11}, A_{12}, \dots, A_{1N})$. Другая «ветвь» эволюции — наращивании иерархии: $Ч \rightleftharpoons A_3 \rightleftharpoons A_2 \rightleftharpoons A_1 \rightleftharpoons O$ и т. д. При этом, «ветви» не мешают друг другу: на каждом из вновь возникающих иерархических уровней допустимо распараллеливание (множественное «агент-агентное» управление): $Ч \rightleftharpoons (A_{21}, A_{22}, \dots, A_{2M})$, $Ч \rightleftharpoons (A_{31}, A_{32}, \dots, A_{3K})$, и т. д. Здесь M и K — число *агентов* второго и третьего уровней, соответственно.

В качестве иллюстрации, на рис. 1 с компактно представлено четыре уровня *агентов*. Первый, второй и третий уровни *многоагентные*; на четвёртом уровне имеется единственный *агент* A_{41} . К первому агенту 2-го уровня A_{21} подключено N агентов 1-го уровня; к первому агенту 3-го уровня A_{31} подключено M агентов 2-го уровня; к первому (единственному) агенту 4-го уровня A_{41} подключено K агентов 3-го уровня. Остальное «ветвление агентур» — условно не показано.

Представленная структура может быть дополнена внутриуровневыми и межуровневыми межагентскими связями (рис. 1 и, р, с). Последнее может рассматриваться как образование внутриуровневых и межуровневых «интеллектуальных агентских сообществ», о существовании и активности которых Ч (при надлежащем автономном саморазвитии интеллектуальных агентов и их сообществ) может даже и не знать. Внутриуровневые межагентские связи — в целом укладываются в общую иерархическую структуру системы. Межуровневые межагентские связи следует рассматривать как расширение (дополнение) иерархической структуры. Использование таких связей в системе, по-видимому, продуктивно, т. к. существенно наращивает многовариантность; но требует отдельного исследования, так как в силу той же многовариантности, они могут таить в себе опасности для системы.

В качестве действующего прототипа, следует указать человеческий мозг. Иерархические

представления об организации структуры его нейронных связей могут быть приписаны ему, по-видимому, лишь в некотором укрупнённом масштабе. Кроме того, сама идея иерархичности есть «человеческое изобретение», т.е. модель, предположительно описывающая реальность.

Характерно, что даже при столь простом (приведённом выше, рис. 1) структурировании процесса эволюционирования МАС, могут быть указаны перспективные направления («ветви») развития, а также направления, «выходящие из-под контроля» человека (межагентские неиерархические внутриуровневые и межуровневые связи) и потому содержащие потенциальные угрозы.

3. Индивидуальный человеческий интеллект

Формирования индивидуального ЧИ иллюстрирует схема табл. 1. *Уровень* (левый столбец) есть категория, в основном (преимущественно) коррелирующая с возрастом человека: временной срез индивидуального развития ЧИ. Разумеется, корреляция вовсе не означает (не гарантирует) неизбежного прохождения каждым человеком всех *уровней* в течение жизни. Нумерация *уровней* — инверсная (снизу — вверх), что отображает субъективное представление о росте, развитии в процессе жизни и продвижении «от нижнего *уровня* к верхнему».

Таблица 1

Уровневость формирования индивидуального человеческого интеллекта

	Уровень	Проявление	Сущность
6	духовность	харизма	бог
5	идентичность	миссия	супервизор
4	убежденность	цель	координатор
3	способность	поступок	личность
2	поведение	воздействие	индивид
1	рефлекторность	реакция	особь

Проявление и *сущность* (средний и правый столбцы табл. 1) — категории, определяющие «наличные возможности» и «потенциальный социально-обусловленный статус» человека, соответственно. Разумеется, «наличие возможности» вовсе не означает автоматической реализации *проявления*, присущего соответствующему *уровню*, а «статус» не предполагает неизбежную реализацию конкретного человека в соответствии с его *сущностью*. *Проявление* и *сущность*, как представляется, могут реализоваться только при стечении обстоятельств, или в случае возникновения соответствующей потребности.

Человек — существо социальное, поэтому на каждом *уровне* идёт определённая конкуренция, в соответствии с которой реализуются (или не реализуются) *проявление* и *сущность*.

Первый *уровень* — *рефлекторность* — соотносится с ранним периодом жизни: ориентировочно, первым годом после рождения. При рождении человек, судя по всему, уже наделён некоторым набором безусловных рефлексов. В первые месяцы после рождения у него интенсивно складываются и закрепляются *реакции* на новые внешние воздействия. Поскольку в это же время наращивается нейронная сеть головного мозга, вновь возникающая *рефлекторность* (комплекс условных рефлексов) обретает, буквально, «аппаратную реализацию». Содержательно, данный *уровень* характеризует чисто биологический план: *особь*.

Каждый последующий *уровень* представляет собой множественность образований информационных структур предыдущего *уровня*. Вторым *уровнем* — *поведение* — реализуется как структура над наработанными комплексами условных рефлексов. Устанавливаются связи и соответствия между проявлениями *реакций* различных наработанных комплексов *рефлекторного уровня*. Проявлением *поведенческого уровня* является реагирование на внешние *воздействия* и формирование инициативных или ответных *воздействий* на объекты или процессы внешнего мира. Человек при этом проявляется как *индивид*. Его *поведенческий уровень* является строго индивидуальным, согласно *рефлекторному* комплексу, сформировавшемуся на предыдущем *уровне*.

Третий *уровень* — *способность* — формируется как совокупность структур *поведенческого уровня*. На этом *уровне* развития отрабатываются и сопоставляются различные реализованные *поведенческие* структуры (модели). Некоторые из моделей реализуются данным *индивидом* более эффективно. Их он предпочитает. Другие — менее эффективные модели — отвергаются или подлежат переработке. Предпочтение одних моделей поведения по сравнению с другими есть *поступок*, и является *проявлением способностей*. Человек при этом (на этом *уровне*) реализуется как *личность*.

Четвёртый *уровень* — *убежденность* — есть информационная структура, сформированная над комплексом *способностей*, как результат их развития и взаимосогласования. Проявлением *уровня убежденности* является целеполагание, включающее формирование *цели* или целей, поэтапное движение к *цели* и возможную корректировку *цели* по мере конкретизации пути к её достижению. *Цель* может быть личностной (внутренней, духовной), либо внешней, предполагающей межличностные взаимодействия. В обоих вариантах человек, находящийся в своём развитии на *уровне убежденности*, должен координировать свои внутренние усилия

в проявлении и дальнейшем совершенствовании своих *способностей*, либо соизмерять свои *поступки* с деятельностью (интеллектуальными проявлениями) других людей. Для обозначения реализации человека на этом уровне хорошо подходит термин *координатор*.

Пятый *уровень* — *идентичность* — является развитием, совершенствованием и взаимным согласованием проявлений четвёртого *уровня*: способность различать в себе или окружающих разные варианты *целеполагания*, воспринимать наличие *убеждённостей*, отличных от текущей собственной. К моменту проявления *идентичности*, процесс целеполагания совершает уже столько итераций, что начинает доминировать социальная обусловленность (мотивированность, масштабность) *цели*. Возникает осознание незначительности индивидуальных личных целей в контексте социальной значимости. В реализации социально обусловленной *цели* человек не может ограничиться только личностным аспектом, что является уже *миссией*.

Механизмом (инструментом) реализации *идентичности* является умение смоделировать в себе *уровень* другого человека: взглянуть на проблему «его глазами», с учётом его *способностей* или *целей*, к которым он стремится. Проявлением *идентичности* является реализация своей *миссии*, для достижения чего человек может манипулировать другими людьми, формировать или корректировать (менять) их *цели*, способствовать реализации *целей* совпадающих с *миссией* и препятствовать реализации *целей*, противоречащих *миссии*. Для обозначения «рода деятельности» человека на этом *уровне*, с учётом проявляемой им социальной сущности, подходит термин *супервизор*.

Шестой *уровень* — *духовность* — означает существенно расширенный социальный аспект развития, возможно, с игнорированием личностных аспектов (или «всего земного») и приоритетным учётом глобальных интересов и потребностей больших групп людей вплоть до человечества (вида *Homo sapiens*) в целом. *Духовность* проявляется как развитая, обобщённая и усовершенствованная *идентичность*, возможность воспринимать множественность *идентичностей*. Люди этого *уровня* очень редки, поэтому проявления из деятельности вызывают большой интерес (резонанс), формируют вокруг них «толпы адептов и почитателей», что характеризуется как *харизматичность*. Уникальность (вплоть до феноменов, воспринимаемых на обыденном уровне как сверхъестественные) и малочисленность людей, находящихся на этом *уровне* развития, позволяют применить к ним

термин *бог* (разумеется, не в мифологическом или теологическом смысле).

По мере «продвижения по жизни», как отмечалось, люди не одинаково (не синхронно) переходят с *уровня* на *уровень*. Основная масса «взрослого человечества» не продвигается выше 3-го *уровня*; 4-го *уровня* достигает лишь несколько процентов; 5-го *уровня* лишь доли процента; а 6-го — буквально единицы. Легко видеть, что 2 - 3 *уровни* — это то, чему учат (воспитывают) в основном в средней школе; 3 - 4 *уровни* охватываются (с разной степенью успешности) высшими учебными заведениями; 4 - 5 *уровни* — последующее индивидуальное развитие (саморазвитие) и совершенствование; 6-й *уровень* — удел немногих.

Объяснение причин подобного «социального неравенства» выходит далеко за пределы нашего рассмотрения. Этапность (уровневость) формирования индивидуального ЧИ интересует нас исключительно в сопоставлении с этапностью в эволюционировании МАС. Это сопоставление позволяет видеть методологическую общность путей развития и предсказывать устойчивость тенденции к дальнейшему наращиванию числа иерархических уровней.

4. Обсуждение

1. Возможны ли *уровни* организации выше 6-го, представленного в таблице 1? По-видимому возможны, но на настоящий момент — не просматриваются. По-видимому, совершенствованное и духовное (интеллектуальное) развитие не должны иметь предела. По-видимому так же интеллект, превышающий человеческий, может иметь особенности (свойства, возможности, аспекты реализации), которые не мыслятся, не усматриваются и даже не подозреваются средствами ЧИ. Нельзя отрицать, что даже 6-й *уровень* — сам по себе — тоже просматривается довольно туманно. Фактически, *уровень духовности* реконструирован, как множественное проявление *уровня идентичности*; а для обозначения *проявления* и *сущности* применены «малонаучные» термины *харизма* и *бог* (последний — с рекомендацией понимать его в «литературно-художественном» смысле).

2. Допустимо ли («научно ли») применять в качестве терминов слова *харизма* и *бог*? В этом вопросе уместно следовать примеру классиков: Э. Резерфорд (1871 - 1937) применял в качестве термина слово *эманация*, изначально употребляемое в лексиконе магии и оккультизма и означающее истечение некой сущности из духовной сферы в материальную.

3. Действительно ли *уровней* именно 6, как показано в таблице 1? Все ли *уровни* в ней

представлены? Возможна ли структуризация с иным (большим) числом уровней? Разумеется, возможна. Первичным является содержательное обнаружение (констатация наличия) сущностей или явлений. Вслед за этим идёт терминология для обозначения этих сущностей или явлений. Появление термина — сигнал, *наличия* явления (реально или предположительно).

В рассматриваемом случае (таблица 1) — по существу, изучаются явления, связанные с человеческой психикой (психологией). Описания носят качественный (гуманитарный, «литературно-художественный») характер. Числовых характеристик нет. Есть только вербальные описания процессов и состояний. Следовательно, возможен другой понятийный аппарат, иная расстановка разграничений между состояниями (*уровнями*) и другое (но в небольших пределах) число *уровней*. В принципиальном плане это дела не меняет, поскольку, как известно [5], ЧИ способен устойчиво различать 7 ± 2 разных состояния. Таким образом, число уровней в таблице 1 соответствует человеческой способности к составлению репрезентативных моделей окружения.

4. Интересен такой феномен: личность в процессе её формирования (во времени) структурирована приблизительно так же, как и человеческое общество в каждом конкретном временном срезе: иерархически. Одно из возможных объяснений, как отмечалось, — иерархическая организация есть «изобретение человека» и достаточно простая модель для описания соответствующего класса явлений. По этой же причине МАС — так же «человеческое изобретение» — эволюционируют аналогично.

Выводы

Сопоставлены этапности становления и развития мультиагентной системы и индивидуального человеческого интеллекта на протяжении индивидуальной человеческой жизни. Совпадение иерархических структур процессов эволюционирования объяснимо с учётом того, что мультиагентная система является распределённым усилителем человеческого интеллекта. В методологическом плане изучение вскрытой закономерности перспективно в связи с вопросами прогнозирования направлений совершенствования мультиагентного подхода в разработке прикладных интеллектуальных систем.

Список литературы:

1. *Михаль О.Ф.* Эволюционирование мультиагентной системы как воспроизведение процесса формирования индивидуального человеческого интеллекта // Тезисы докладов 1-й Международной научно-технической конференции «Полиграфические, мультимедийные и web-технологии» (PMW-2016). Т. 1. — Харьков: ХНУРЭ, 2016. — с. 73-74.
2. *Михаль О.Ф.* Глобально-исторический контекст развития средств вычислительной техники [Текст] / О.Ф. Михаль // Бионика интеллекта: научн. техн. журнал. — 2014. — 1 (82). — С. 55-62.
3. *Айман Найеф Ахмад Альхалайбех, Михаль О.Ф., Руденко О.Г.* Звено «техническая система — человек оператор» как модель информационного взаимодействия интеллектуальных подсистем с разной производительностью [Текст] / Айман Найеф Ахмад Альхалайбех, О.Ф. Михаль, О.Г. Руденко // Восточно-Европейский журнал передовых технологий. 2003. № 6(6). С. 1821.
4. *Дяченко В.А. Михаль О.Ф.* Интеллектуальный аспект обучения самоорганизующихся карт Кохонена [Текст] / В.А. Дяченко, О.Ф. Михаль // Бионика интеллекта: научн. техн. журнал. — 2015. — 2 (85). — С. 52-58.
5. *George A. Miller.* The Magical Number Seven, Plus or Minus Two. // The Psychological Review, 1956, vol. 63, pp. 81-97. (<http://psychclassics.yorku.ca/Miller/>).

Поступила в редакцию 29.11.2016

А.Г. Каграманян¹, Г.Г. Четвериков², В.В. Шляхов³¹ ХНУ, г. Харьков, Украина^{2,3} ХНУРЭ, г. Харьков, Украина, chetvergg@gmail.com

МУЛЬТИАЛГЕБРАИЧЕСКИЕ ДИСКРЕТНЫЕ СИСТЕМЫ ПРИ НАЛИЧИИ АЛГЕБРАИЧЕСКОЙ СТРУКТУРЫ НОСИТЕЛЯ. СООБЩЕНИЕ 1

Известно, что любое n -арное отношение индуцирует процедуру факторизации на своем носителе. Это происходит в силу того, что некоторые элементы носителя отношением “нераспознаются”. Таким образом возникают классы эквивалентностей. Частным случаем n -арных отношений являются бинарные и тернарные. Работа посвящена изучению условий, при которых бинарные и тернарные отношения на классах эквивалентностей индуцируют алгебраическую структуру в виде линейного пространства.

ПРЕДИКАТ, ДИСКРЕТНАЯ СИСТЕМА, КЛАССЫ ЭКВИВАЛЕНТНОСТЕЙ, АЛГЕБРАИЧЕСКАЯ СТРУКТУРА

Введение

Во многих практических ситуациях множество входных сигналов исследуемого объекта представляет собой некоторую алгебраическую структуру. Объясняется это тем, что, как правило, между элементами этого множества существуют определенные связи, которые можно интерпретировать как алгебраические операции [1]. Как уже говорилось выше, правильное распознавание соответствующей структуры во многом определяет адекватность математической модели в целом [2]. В рамках компараторной идентификации это распознавание должно вестись на языке экспериментально проверяемых свойств отношений или предикатов. В любом объекте совокупность его свойств выступает как структурно организованная система взаимосвязанных элементов, что говорит о многолинейности масштабной деформации. Заметим, что областью определения операторов, в данной статье, является линейное пространство.

Не останавливаясь на экспериментальной части, что выходит за рамки статьи, дадим теоретическое решение данной задачи для такой алгебраической структуры, как линейное пространство над некоторым полем. Подобная структура широко распространена на практике [3-7].

При этом в качестве принципиального основания всякого моделирования, как интуитивного аналогизирования, выступает отношение тождества [7]. Известно, что существует несколько типов отношения тождества в плане его формальных характеристик (например, изоморфизм, гомоморфизм). Предварительным условием классификации моделей по характеру воспроизведения сторон оригинала является выявление типов тождества в плане его содержательных характеристик. По своим содержательным характеристикам тождество модели и оригинала может быть: субстанциональным, функциональным и структурным.

Последнее, а именно, конкретный вид алгебраической структуры в виде линейного (гильбертового) пространства и является квинтэссенцией данной статьи.

1. Постановка задачи

Будем предполагать, что система обработки входных сигналов реализует своим поведением четыре предиката, заданных на соответствующей декартовой степени множества M (входные сигналы): один одноместный $P(x)$, один двуместный $E(x, y)$ и два трехместных $S(x, y, z), T(x, y, z)$. Символами x, y, z обозначены входные сигналы системы. Выходными сигналами системы служат элементы 0 и 1, которые являются значениями перечисленных предикатов.

2. Предикатная модель данных в виде линейного пространства

Предикат $P(x)$ формирует классы прообразов коэффициентов, которые могут быть приняты в качестве самих коэффициентов. Предикат $E(x, y)$ — это предикат эквивалентности, заданный на $M \times M$. Он формирует классы прообразов векторов, которые могут быть приняты в качестве самих векторов. Предикат $S(x, y, z)$ задан на P^3 , он определяет операцию сложения коэффициентов. Предикат $T(x, y, z)$ задан на $P \times M \times M$, он определяет операцию умножения коэффициентов на вектор. Рассмотрим множество M , на котором заданы отношения $E(x, y), S(x, y, z), P(x), T(x, y, z)$, удовлетворяющие следующим свойствам:

- 1) $E(x, x) = 1$;
- 2) $E(x, y) = 1 \Rightarrow E(y, x) = 1$;
- 3) $E(x, y) = 1, E(y, z) = 1 \Rightarrow E(x, z) = 1$;
- 4) $\forall x, y \exists z : S(x, y, z) = 1$;
- 5) $S(x, y, z) = 1, S(x, y, z') = 1 \Rightarrow E(z, z') = 1$;
- 6) $S(x, y, z) = 1, S(x, y', z) = 1 \Rightarrow E(y, y') = 1$;

- 7) $S(x, y, z) = 1, S(x', y, z) = 1 \Rightarrow E(x, x') = 1$;
- 8) $S(x, y, z) = 1 \Rightarrow S(y, x, z) = 1$;
- 9) $S(x, y, z) = 1, E(z, z') = 1 \Rightarrow S(x, y, z') = 1$;
- 10) $S(x, y, z) = 1, E(y, y') = 1 \Rightarrow S(x, y', z) = 1$;
- 11) $S(x, y, z) = 1, E(x, x') = 1 \Rightarrow S(x', y, z) = 1$;
- 12) $S(x, y, z) = 1, S(z, t, r) = 1, S(y, t, p) = 1 \Rightarrow S(x, p, r) = 1$;
- 13) $\exists 0 : S(x, y, x) = 1 \Rightarrow E(y, 0) = 1$;
- 14) $\forall x \exists (-x) : S(x, -x, y) = 1 \Rightarrow E(y, 0) = 1$;
- 15) $P(0) = 1$;
- 16) $P(x) = 1, P(y) = 1, S(x, y, z) = 1 \Rightarrow P(z) = 1$;
- 17) $P(x) = 1, E(x, y) = 1 \Rightarrow P(y) = 1$;
- 18) $\forall x, y \exists z : P(x) = 1 \Rightarrow T(x, y, z) = 1$;
- 19) $P(x) = 0, P(y) = 0 \Rightarrow T(x, y, z) = 0$;
- 20) $P(x) = 1, P(y) = 1, T(x, y, z) = 1 \Rightarrow P(z) = 1$;
- 21) $T(x, y, z) = 1 \Rightarrow T(x, y, z) = 1$;
- 22) $T(x, y, z) = 1, T(x, y, z') = 1 \Rightarrow E(z, z') = 1$;
- 23) $T(x, y, z) = 1, T(x, y', z) = 1 \Rightarrow E(y, y') = 1$;
- 24) $T(x, y, z) = 1, T(x', y, z) = 1 \Rightarrow E(x, x') = 1$;
- 25) $T(x, y, z) = 1, T(z, z') = 1 \Rightarrow T(x, y, z') = 1$;
- 26) $T(x, y, z) = 1, T(y, y') = 1 \Rightarrow T(x, y', z) = 1$;
- 27) $T(x, y, z) = 1, T(x, x') = 1 \Rightarrow T(x', y, z) = 1$;
- 28) $T(x, y, z) = 1, P(x) = 1, E(z, 0) = 1 \Rightarrow E(x, 0) = 1$;
- 29) $E(z, 0) = 1, T(x, y, z) = 1 \Rightarrow E(z, 0) = 1$;
- 30) $T(x, y, z) = 1, P(x) = 1, P(y) = 1, T(z, p, r) = 1, T(y, p, t) = 1 \Rightarrow T(x, t, r) = 1$;
- 31) $T(x, y, z) = 1, T(x', y, z') = 1, S(x, x', t) = 1, P(x) = P(x') = 1, S(z, z', p) = 1 \Rightarrow T(t, y, p) = 1$;
- 32) $P(x) = 1, T(x, y, z) = 1, T(x, y', z') = 1, S(z, z', t) = 1, S(y, y', p) = 1 \Rightarrow T(x, p, t) = 1$;
 $\exists 1 : P(1) = 1, T(y, x, x) = 1 \Rightarrow E(1, y) = 1$;
- 33) $P(x) = 1 \exists x^{-1} : T(x, x^{-1}, y) = 1 \Rightarrow E(1, y) = 1$;
- 34) $\exists \{t_i\}_{i=1}^n : \forall x \exists \{y_i(x)\}_{i=1}^n$:
 а) $P(y_i(x)) = 1$; б) $T(y_i(x), t_i, z_i) = 1$,
 $S(z_1, z_2, r_1) = 1, S(z_1, z_3, r_2) = 1, \dots, S(z_{n-2}, z_n, r_{n-1}) = 1 \Rightarrow E(x, r_{n-1}) = 1$;
- в) $\forall \{h_i(x)\}_{i=1}^n$, удовлетворяющий а) и б)
 $\Rightarrow E(h_i, y_i(x)) = 1$.
- 35) $P(x) = 1, P(z) = 1, S(x, y, z) = 1 \Rightarrow P(y) = 1$.

В этом случае множество M разбивается на классы эквивалентности отношением $E(x, y)$. Классы эквивалентности будем обозначать $A, B, C, R, T, \dots$, а все множество классов N . Тогда, как показано в работе [3], $E(x, y)$ представим в виде

$$E(x, y) = D(Fx, Fy),$$

где D — предикат равенства на $N \times N$, а $F : M \rightarrow N$ (причем $Fx = Fy \Leftrightarrow E(x, y) = 1$).

Наша задача состоит в том, чтобы показать, что заданные отношения индуцируют структуру n -мерного линейного пространства на классах эквивалентности.

Утверждение 1. Если на классах эквивалентности ввести операцию (сложения) по правилу $A + B = C$ тогда и только тогда, когда $\forall x, y, z$:

$$x \in A, y \in B, z \in C, S(x, y, z) = 1,$$

то определение будет корректным и относительно данной операции N образует абелеву группу.

Доказательство. Сначала покажем корректность определения. Выделим произвольным образом два класса эквивалентности $A, B \in N$ и два представителя каждого класса $x \in A, y \in B$. Тогда из свойства 4) вытекает, что найдется $z \in C$, для которого $S(x, y, z) = 1$. Это означает, что $A + B = C$. Таким образом, операция определена на любых парах $A, B \in N$, более того, единственным образом. Действительно, пусть $C' \neq C$, и

$$A + B = C, A + B = C'.$$

Тогда для произвольного $z' \in C'$ имеем $S(x, y, z') = 1$. С учетом того, что $S(x, y, z) = 1$, из свойства 5) получим $E(z, z') = 1$ или $z' \in C$. Значит, $C \cap C' \neq \emptyset$, а поскольку различные классы имеют пустое пересечение, то $C \in C'$. Получили противоречие. Теперь покажем, что класс Z не зависит от выбора $x \in A$ и $y \in B$. Допустим, $x, x' \in A$ и $y, y' \in B$. Тогда, так как $S(x, y, z) = 1$ и $E(x', x) = 1$, то на основании свойства 11) получим $S(x', y, z) = 1$. Далее, учитывая свойство 10) и равенство $E(y', y) = 1$, будем иметь, что $S(x', y', z) = 1$, но это и означает, что операция сложения не зависит от выбора элементов в классах A и B . Следовательно, операция, введенная нами, корректна.

Покажем, что относительно этой операции N образует абелеву группу.

Допустим, $A + B = C$. Тогда для любых $x \in A, y \in B, z \in C$ выполняется $S(x, y, z) = 1$. В этом случае из свойства 8) вытекает $S(x, y, z) = 1$ или $A + B = C$. Таким образом,

$$A + B = B + A,$$

следовательно, операция коммутативна.

Она также и ассоциативна. Пусть $(A + B) + C = R, A + B = T, B + C = G$. Тогда для

представителей классов выполняются равенства $S(x, y, t) = 1, S(y, z, g) = 1, S(t, z, r) = 1$. С учетом свойства 12) получим $S(x, g, r) = 1$. Это означает $A + G = R$ или $A + (B + C) = R$, т.е.

$$(A + B) + C = A + (B + C),$$

значит, операция ассоциативна.

Рассмотрим свойство 13). В нем утверждается, что существует $O \in M$ такой, что для любого x выполняется $S(x, O, x) = 1$. Следовательно, $A + O = A$ (O — класс эквивалентности, которому принадлежит O). Причем O — единственен, поскольку, если найдется $O' \neq O$, то для $y \in O'$ получим $S(x, y, z) = 1$ и из второй части свойства 13) будет вытекать $E(y, O) = 1$, то есть $O' = O$. Таким образом, среди N найдется единственный элемент O , который выполняет роль нуля относительно данной операции.

Наконец, остановимся на существовании обратного элемента. Выберем произвольный класс A и его представитель $x \in A$. Тогда по свойству 14) имеем: найдется $-x$, для которого из $S(x, -x, y) = 1$ вытекает $E(y, 0) = 1$. Пусть $-x \in -A$, тогда $A + (-A) = B$, где $y \in B$, но с учетом $E(y, 0) = 1$ получим $y \in 0$ или $B = 0$. Таким образом,

$$A + (-A) = 0,$$

причем $-A$ единственен. Поскольку, если равенство выполняется для какого-то другого класса C , то $S(x, z, 0) = 1, S(x, -x, y) = 1$ и $E(y, 0) = 1$. Тогда из свойства 9) получим $S(x, -x, 0) = 1$, а из свойства 6) $-E(-x, z) = 1$, т.е. $-x \in C$ или $-A = C$.

Утверждение доказано.

Утверждение 2. Отношение $P(x)$, заданное на M , определяет его подмножество M' , являющееся объединением классов эквивалентности, причем множество классов, входящих в M' , образуют подгруппу группы всех классов относительно введенной операции сложения.

Доказательство. Чтобы доказать первую часть утверждения леммы, нужно показать, что для любого класса эквивалентности S выполняется: $A \cap M'$ равно либо пустому множеству, либо A . Действительно, пусть $x \in A$, тогда, если $P(x) = 1$ и $E(x, y) = 1$, то из свойства 17) вытекает $P(y) = 1$, т.е. $A \subset M'$. Если $P(x) = 0$, то для любого $y \in A$: $P(y) = 0$, так как в противном случае при $P(y) = 1, E(x, y) = 1$ из свойства 17) получим $P(x) = 1$. Противоречие. Значит, если $P(x) = 0$, то $A \cap M' = \emptyset$. Таким образом, $M' = \{x: P(x) = 1\}$ является объединением классов эквивалентности. Набор этих классов обозначим через N' . Докажем, что $N' \subset N$ представляет собой подгруппу относительно операции сложения классов. Пусть $A, B \in N'$ и $A + B = C$. $P(x) = 1, P(y) = 1$

и $S(x, y, z) = 1$. Свойство 16) утверждает, что $P(z) = 1$, следовательно, $z \in N'$. Значит операция сложения не выводит за пределы N' . Свойство 15), утверждающее, что $P(0) = 1$, означает, что $0 \in N'$. Остается показать, что обратный элемент принадлежит N' . Для этого рассмотрим $A \in N'$ и $-A$. Допустим $-A$ не принадлежит N' . Тогда $P(-x) = 0, P(x) = 1, P(0) = 1$ и $S(x, -x, 0) = 1$. Но последний набор равенств противоречит свойству 36). Утверждение доказано.

Утверждение 3. Если на классах эквивалентности, принадлежащих N' , ввести операцию (умножения) по правилу: $AB = C$, тогда и только тогда, когда $T(x, y, z) = 1$, для $\forall x \in A, y \in B, z \in C$, то это определение будет корректным, и относительно данных операций сложения и умножения классы эквивалентности N' образуют поле.

Доказательство. Корректность определения введенной операции выясняется следующим образом. Из свойств 18) и 20) вытекает, что для любых $A, B \in N'$ и их произвольных элементов $x \in A$ и $y \in B$ найдется z , для которого $T(x, y, z) = 1$ и $P(z) = 1$. Следовательно, в силу данного определения операции умножения, существует класс $C \subset N'$, для которого $AB = C$, и $y \in B$. Действительно, пусть для какого-то z' выполняется равенство $T(x, y, z') = 1$, но тогда из свойства 22) и равенства $T(x, y, z) = 1$ вытекает $E(z, z') = 1$, то есть $z' \in C$. С другой стороны, если первоначально мы выбрали $x \neq x' \in A$ и $y \neq y' \in B$, то из равенства $E(x, x') = 1, E(y, y') = 1$ и $T(x, y, z) = 1$ и свойств 26), 27) получим $T(x', y', z) = 1$. Таким образом, класс C не зависит от первоначального выбора элементов классов A и B . Следовательно, введенное определение операции умножения корректно. Покажем теперь, что относительно операций сложения и умножения множество классов N' образует поле. Из утверждения 7.2 вытекает, что N' — абелева группа по сложению. Докажем, что по умножению N' — тоже абелева группа. Рассмотрим два произвольных класса $A, B \in N'$ и пусть $AB \in C$. Последнее равенство означает, что $T(x, y, z) = 1$ для $x \in A, y \in B, z \in C$, но из свойства 21) в этом случае вытекает $T(x, y, z) = 1$, то есть $BA = C$. Значит, операция коммутативна. Из свойства 30) следует ее ассоциативность. Действительно, рассмотрим $(AB)C$ и пусть $AB = R, RC = T$ и $BC = P$. Тогда для представителей этих классов будут выполняться равенства $T(x, y, z) = 1, T(r, z, t) = 1, T(y, z, p) = 1$, но по свойству 30) получим $T(x, p, t) = 1$, т.е. $(AB)C = A(BC)$.

Рассмотрим свойство 33). Оно утверждает, что для $\forall x \in A \subset N'$ найдется $x^{-1} \in A^{-1} \subset N'$ такой, что для любого $z \in C$ имеет место $T(x, x^{-1}, y) = 1$

и $T(y, z, z) = 1$. Если $z \in M'$, последние два равенства означают, что $(AA^{-1})C = C$, а $AA^{-1} = B \in N'$ по свойству 20). Покажем, что A^{-1} не зависит от самого класса A . Действительно, пусть $x_1 \neq x$ и $x, x_1 \in A$, тогда

$$E(x_1, x) = 1, T(x, x^{-1}, y) = 1, T(x_1, x_1^{-1}, y_1) = 1, T(y, z, z) = 1, T(y_1, z, z) = 1.$$

Из этих равенств на основании свойств 25), 27), 24) получим

$$E(y, y_1) = 1, T(x_1, x^{-1}, y) = 1, T(x_1, x^{-1}, y) = 1, T(x_1, x^{-1}, y_1) = 1$$

и $E(x_1^{-1}, x^{-1}) = 1$, то есть $x_1^{-1} \in A^{-1}$. Допустим, $A_1 = A_2$ и $A_1 A_1^{-1} = B_1$ и $A_2 A_2^{-1} = B_2$. Тогда $T(x_1, x_1^{-1}, y_1) = 1$, $T(x_2, x_2^{-1}, y_2) = 1$, и с учетом свойства 32) получим для $\forall z: T(y_1, z, z) = T(y_2, z, z) = 1$. Теперь воспользуемся свойством 24), тогда $E(y_1, y_2) = 1$, следовательно $B_1 = B_2$. Таким образом, для любого $A \subset N': AA^{-1} = B$ не зависит от A и поскольку для $\forall z \in N'$ имеем $(AA^{-1})C = C$, то AA^{-1} играет роль единицы относительно умножения. Поэтому в дальнейшем будем обозначать $AA^{-1} = E$.

Окончательно можем сделать вывод, что относительно операции умножения и сложения множества классов N' образуют группы. Эти операции еще связаны между собой так, что N' является полем. Покажем это. Фактически нами проверены все аксиомы поля, кроме дистрибутивности. А эта аксиома вытекает из свойства 31), так как для произвольных классов $A, A', B, C, C', T, P \subset N'$ из равенств $AB = C, A'B = C, A + A' = T, C = C = P$ вытекает для их представителей $T(x, y, y) = T(x', y, z') = S(x, x', t) = S(z, z', p) = 1$ и из 31) имеем $T(t, y, p) = 1$, то есть $TB = P$, следовательно $AB + A'B = (A + A')B$. Значит, дистрибутивность выполняется, что завершает доказательство утверждения.

Суммируем результаты доказанных нами утверждений. Заданные отношения:

1) разбивают исходное множество на классы эквивалентности;

2) эти классы эквивалентности образуют множество N , на котором индуцируется операция сложения и относительно нее множество N является группой;

3) во множестве N может быть выделено подмножество, $N' \subset N$, на котором исходные отношения индуцируют операцию умножения, относительно операций умножения и сложения множество N' является полем.

Теперь мы можем сформулировать и доказать теорему, являющуюся целью настоящего подраздела.

Теорема 1. Множество классов эквивалентности N является конечномерным линейным пространством над полем N' с операцией сложения векторов, определенной в утверждении 1 и с операцией умножения вектора на элемент поля, определенной в утверждении 3.

Доказательство. Для начала заметим, что операция умножения, индуцированная отношением T и введенная нами для элементов поля N' , аналогично тому, как это сделано в утверждении 7.3, может быть корректно определена для элементов N , то есть умножение векторов на элементы поля N' . Доказательство этого факта повторяет соответствующие рассуждения в доказательстве утверждения 3. Приступим к доказательству теоремы.

Нами уже показано, что относительно операции сложения множество элементов (в дальнейшем будем называть их векторами, но обозначать большими буквами, так как это классы эквивалентности) N образуют группу. Есть также поле N и операция умножения элементов поля на вектор. Покажем, что эта операция обладает следующими свойствами:

- 1) если $A \in N', B, C \in N$, то $A(B + C) = AB + AC$;
- 2) если $A, B \in N', C \in N$, то $(AB)C = A(BC)$;
- 3) если $A, B \in N', C \in N$, то $(A + B)C = AC + BC$;
- 4) для $O, E \in N'$ имеет место $OA = O$ и $EA = A$ для произвольного $A \in N$.

Первое из выше перечисленных свойств вытекает из свойства отношений 32), аналогично тому, как это сделано в утверждении 3. Третье свойство и второе по сути дела нами доказаны в утверждении 3, когда речь шла об ассоциативности и дистрибутивности операций сложения и умножения. Остановимся на свойстве 4). То, что $EA = A$ вытекает из свойства отношений 33). Для обоснования второго равенства можно использовать свойство 29), из которого следует: если рассмотреть, $OA = C$, то для элементов будет $y' \in O$ и $S(z, -z, y) = 1$ вытекает $E(y', y) = 1$, тогда, $T(y', x, t) = 1$, то есть $OA = T$ и $E(y, t) = 1$ или $T = O$. Значит, $OA = O$.

Таким образом, мы показали, что множество N является линейным пространством над полем N' . Докажем его конечномерность. Она записана в свойстве 34). Из него следует, что найдутся такие элементы $t_1 \in T_1, \dots, t_n \in T_n$ такие, что для любого $x \in A$ существуют единственные $y_1(x) \in B_1(A), \dots, y_n(x) \in B_n(A)$, для которых (справа мы будем писать то, что выполняется для классов)

А) $P(y_i) = 1$, то есть $B_i(A) \in N$ — элементы поля;

Б) $T(y_i(x), t_i, z_i) = 1$, то есть $B_i(A)T_i = Z_i$;

$S(z_1, z_2, r_1) = 1$, то есть $C_1 + C_2 = R_1$ и т.д.;

$S(r_{n-2}, z_n, r_{n-1}) = 1$, т.е. $R_{n-2} + C_n = R_{n-1}$ или

$$C_1 + C_2 + \dots + C_n = R_{n-1}.$$

Тогда по свойству 35) вытекает, что $E(x, r_{n-1}) = 1$, следовательно, $R_{n-1} = A$. Окончательно получаем разложение по базису $T_1, \dots, T_n$

$$A = B_1(A)T_1 + \dots + B_n(A)T_n.$$

Единственность классов $B_i(A)$ (в отличие от элементов $y_i(x)$, о которых говорится в свойстве 35)) вытекает из пункта в) свойства 34). Теорема доказана.

Выводы

Найдена аксиоматика существования мультиалгебраической системы в виде линейного пространства. Фактически получены необходимые и достаточные условия, выполнение которых для бинарных и тернарных отношений индуцирует структуру линейного пространства, элементами которого являются классы эквивалентности над произвольным полем.

Список литературы:

1. Мальцев А.И. Алгебраические системы. — М.: Наука, 1970. — 392 с.
2. Четвериков Г.Г., Дударь З.В., Вечирская И.Д. Дискретные структуры. — Харьков: НУРЭ, 2015. — 322 с.
3. Машталир В.П., Шляхов В.В., Яковлев С.В. Групповые структуры на фактор-множествах в задачах классификации // Кибернетика и системный анализ. — 2014. — №4. — с. 27–42.
4. Герасин С.Н., Шляхов В.В. О внешней и внутренней согласованности произвольных n -арных отношений // Доп. НАН Украины. — 2006. — №1. — с. 77–82.
5. Бондаренко Н.Ф., Машталир В.П., Шляхов В.В. Мультигруппы, индуцируемые произвольным отношением // Доп. НАН Украины. — 2012. — №4. — с. 39–42.
6. Каграманян А.Г., Машталир В.П., Шляхов В.В. Условия существования тернарной мультисистемы с единым носителем // Вісник національного університету. — Серія: Математичне моделювання. Інформаційні технології. — 2011. — №960, Вип. 16. — Харків: ХНУ імені В.Н. Каразіна. — с. 148–158.
7. Четвериков Г.Г. Формальні моделі та принципи побудови універсальних k -значних структур мовних систем штучного інтелекту // Доповіді НАН України — 2001. — №1(41). — с. 76–79.

Поступила в редколлегию 24.10.2016

УДК 519.62

А.Г. Каграманян¹, Г.Г. Четвериков², В.В. Шляхов³¹ ХНУ, г. Харьков, Украина^{2,3} ХНУРЭ, г. Харьков, Украина, chetvergg@gmail.com

МУЛЬТИАЛГЕБРАИЧЕСКИЕ ДИСКРЕТНЫЕ СИСТЕМЫ ПРИ НАЛИЧИИ АЛГЕБРАИЧЕСКОЙ СТРУКТУРЫ НОСИТЕЛЯ. СООБЩЕНИЕ 2

Рассматриваются линейные предикаты, заданные на конусе гильбертового пространства. Получены необходимые и достаточные условия их существования. Работа посвящена изучению условий, при которых бинарные и тернарные отношения на классах эквивалентностей индуцируют алгебраическую структуру в виде линейного пространства. Показана специфика этих условий в сравнении для случая полного линейного пространства.

АЛГЕБРАИЧЕСКАЯ СТРУКТУРА, ДИСКРЕТНАЯ СИСТЕМА, КОНУС, ГИЛЬБЕРТОВО ПРОСТРАНСТВО, ЛИНЕЙНОЕ ПРОСТРАНСТВО, ПРЕДИКАТ

Введение

Бинарные отношения или предикаты, которые изучаются в данной работе, в некоторых ситуациях задаются не на декартовом квадрате какого-либо множества, а на его части. Ограничения, возникающие при этом диктуются физическими условиями и характером исследуемого объекта. В случае сенсорных систем, в частности при изучении органа зрения человека, входной сигнал — это функция яркости в зависимости от длины волны, т.е. с положительными значениями. Естественно это множество не все, а часть линейного пространства.

Очень часто множество входных сигналов исследуемого объекта представляет собой некоторую алгебраическую структуру [1, 2]. Объясняется это тем, что, как правило, между элементами этого множества существуют определенные связи, которые можно интерпретировать как алгебраические операции. Как уже говорилось выше, правильное распознавание соответствующей структуры во многом определяет адекватность математической модели в целом. В рамках компараторной идентификации это распознавание должно вестись на языке экспериментально проверяемых свойств отношений или предикатов. Не останавливаясь на экспериментальной части, что выходит за рамки данной статьи, дадим теоретическое решение данной задачи для такой алгебраической структуры, как линейное пространство над некоторым полем. Подобная структура широко распространена на практике. Областью определения операторов, которые исследуются в данной работе, является линейное пространство.

1. Постановка задачи

Во многих практических задачах, решение которых связано с применением компараторной идентификации возникает ситуация, когда множество входных сигналов не является линейным или

гильбертовым пространством, а составляет какую-либо его часть или подмножество. Исследование, проведенные в работах [3–6], показывают, что наибольший интерес представляют три следующих варианта множества входных сигналов: конус K линейного пространства, выпуклое тело V и множество с непустой внутренностью Ω . Цель этой работы — изучить возможность компараторной идентификации линейных операторов на конусе линейного пространства.

В дальнейшем будем считать, что заданы предикаты на подмножествах линейного нормированного пространства L над полем действительных чисел R .

Сначала рассмотрим случай положительного конуса K и на его примере поговорим об особенностях, которые здесь возникают (они характерны и для других вариантов).

2. Линейные предикаты на конусе линейного пространства

Пусть предикат $E(x, y)$ задан на декартовом квадрате положительного конуса $K \subset L, R > 0$. Необходимо найти характеристические свойства, обеспечивающие его представимость в линейном виде. На первый взгляд может показаться, что это задача полностью совпадает с той, которая решалась в полном пространстве. Однако это не так. Сужение области определения предиката $E(x, y)$ приводит к ряду принципиальных отличий, которые не позволяют автоматически перенести условие полного линейного пространства. Сейчас мы это покажем.

Рассмотрим свойство n -мерности. Оно гласит следующее: существует набор векторов $\{e_i\}_{i=1}^n = L$ (в данном случае нам придется формулировать принадлежащий K), такой, что для любого $x \in K$ найдется единственный набор чисел $\{\alpha_i(x)\}_{i=1}^n$, для которого

$$E(x, \sum_{i=1}^n \alpha_i(x) e_i) = 1.$$

Однако здесь сразу возникает замечание, которого не было ранее. Необходимо как-то регламентировать этот набор чисел с тем, чтобы обеспечить принадлежность к положительному конусу

K линейной комбинацией $\sum_{i=1}^n \alpha_i(x) e_i$. Это можно сделать, добавив, скажем, условие, что числа $\{\alpha_i(x) > 0\}_{i=1}^n$ (последнее влечет за собой трудность в доказательстве необходимости, так как придется доказывать положительность решения системы линейных уравнений) или каким-либо другим способом. Отсюда следует, что условие n -мерности требует изменения.

Непрерывность в некоторых случаях тоже нельзя сохранить. Например, в пространствах $L_2[a, b]$ любая окрестность точки положительного конуса содержит «проколы», т.е. точки, ему не принадлежащие.

В целом ограничение на множестве входных сигналов заставляет постоянно следить за принадлежностью аргументов предикатов области его определения. Это обстоятельство вносит принципиальные изменения в формулировки аксиом и накладывает целый ряд технических трудностей в доказательство. Таким образом, возникает необходимость подробного рассмотрения отдельных типов ограничений.

Как уже говорилось выше, начнем рассмотрение с положительного конуса K пространства $\langle L, R \rangle$. Допустим, на нем задан линейный предикат

$$E(x, y) = D(F[x], F[y]), \quad (1)$$

где D — предикат равенства на

$$R^n \times R^n, F[x] = (f_1(x), \dots, f_n(x)), [f_i]_{i=1}^n$$

— линейно независимые функционалы над $\langle L, R \rangle$.

Предположим сначала, что $\dim L$ конечно, но $\dim L > n$. Изучим некоторые свойства линейного предиката. Зафиксируем систему линейно независимых векторов $\{e_i\}_{i=1}^n \in K$ и для произвольного $x \in K$ составим систему линейных уравнений относительно неизвестных $\{\alpha_i(x)\}_{i=1}^n$ вида

$$f_k(x) = \sum_{i=1}^n \alpha_i(x) f_k(e_i), k = \overline{1, n}. \quad (2)$$

Матрица этой системы, равная $A = (f_k(e_i))_{k,i=1}$ имеет определитель, равный нулю, если строки или столбцы ее линейно независимы, т.е. найдется набор чисел $\lambda_1, \dots, \lambda_n$ для которых

$$f_k(\sum_{i=1}^n \lambda_i e_i) = 0, k = \overline{1, n}, \sum_{i=1}^n \lambda_i^2 \neq 0.$$

Это может происходить только в том случае, когда для функционала $f_1, \sum_{i=1}^n \lambda_i e_i \in \text{Ker } f_1$. Однако

подобной ситуации всегда можно избежать. $\text{Ker } f_1$ представляет собой гиперплоскость пространства $\langle L, R \rangle$, а $\zeta(\{e_i\}_{i=1}^n)$ — подпространство размерности n . Всегда можно сделать выбор $\{e_i\}_{i=1}^n$ таким образом, чтобы $\text{Ker } f_1 \cap \zeta(\{e_i\}_{i=1}^n) = \emptyset$. При таком выборе $\det A \neq 0$ и система (2) будет иметь единственное решение. Это решение будет характеризоваться единственным подмножеством индексов $1, 2, \dots, n$, обозначим его $I(x)$, для которого $\alpha_i(x) \leq 0$, если $i \in I(x)$ и $\alpha_i(x) > 0$, если $i \notin I(x)$. Обозначим

$$\alpha_i(x) = \begin{cases} \alpha_i(x) & \text{при } i \notin I(x), \\ -\alpha_i(x) & \text{при } i \in I(x). \end{cases} \quad (3)$$

Тогда система (2) с учетом линейности $\{f_i(x)\}_{i=1}^n$ может быть переписана в виде

$$f_k(x + \sum_{i \in I(x)} a_i(x) e_i) = f_k(x + \sum_{i \notin I(x)} a_i(x) e_i), k = \overline{1, n}. \quad (4)$$

Заметим, что $a_i(x) \geq 0$ при любом i , поэтому

$$x + \sum_{i \in I(x)} a_i(x) e_i, \sum_{i \notin I(x)} a_i(x) e_i \in K.$$

Это означает, что для линейного предиката вида (1), заданного на $K \times K$, равенства (4) эквивалентны равенству

$$E(x + \sum_{i \in I(x)} a_i(x) e_i, \sum_{i \notin I(x)} a_i(x) e_i) = 1. \quad (5)$$

Таким образом, мы установили аналог свойства n -мерности для линейного предиката, заданного на положительном конусе. Сформулируем его.

Будем говорить, что предикат $E(x, y)$ обладает свойством n -мерности, если существует система линейно независимых векторов $\{e_i\}_{i=1}^n \in K$ такая, что для любого $x \in K$ найдется единственный набор чисел $\{a_i(x)\}_{i=1}^n$ и единственное подмножество $I(x) \subset \{1, \dots, n\}$, для которых $a_i(x) \geq 0$, при $i \in I(x)$, $a_i(x) > 0$, при $i \notin I(x)$ и имеет место равенство (5).

Сформулируем еще набор свойств, которому удовлетворяет предикат $E(x, y)$ и в справедливости которых легко убедиться непосредственной проверкой.

Однородность. Если $E(x, y) = 1$, то для любого $\lambda > 0$ $E(\lambda x, \lambda y) = 1$.

Аддитивность. Для произвольных $x, y, x', y' \in K$ из равенств $E(x, y) = 1$, $E(x', y')$ вытекают равенства $E(x + x', y + y') = 1$, $E(x + y', y + x') = 1$.

Полуаддитивность. Для произвольных $x, y, z \in K$ из равенства $E(x + z, y + z) = 1$ вытекает $E(x, y) = 1$.

Теперь мы можем сформулировать и доказать теорему об условиях существования линейных предикатов на положительном конусе.

Теорема 1. Для того, чтобы предикат $E(x, y)$, заданный на $K \times K$, был линейным, необходимо и достаточно, чтобы он обладал свойствами n -мерности, однородности, аддитивности и полуаддитивности.

Доказательство. Фактически необходимость сформулированных выше условий уже доказана. Остановимся на достаточности.

Допустим, предикат $E(x, y)$ удовлетворяет условиям теоремы. Тогда из n -мерности и аддитивности при произвольных $x, y \in K$ будем иметь

$$\begin{aligned} E(x + \sum_{i \in I(x)} a_i(x)e_i, \sum_{i \in I(x)} a_i(x)e_i) &= 1, \\ E(y + \sum_{i \in I(y)} a_i(y)e_i, \sum_{i \in I(y)} a_i(y)e_i) &= 1, \\ E(x + y + \sum_{i \in I(x) \cap I(y)} (a_i(x)e_i + a_i(y)e_i) + \\ &+ \sum_{i \in I(x) \setminus I(y)} a_i(x)e_i + \sum_{i \in I(y) \setminus I(x)} a_i(y)e_i, \\ &\sum_{i \in I(x) \setminus I(y)} (a_i(x) + a_i(y))e_i + \sum_{i \in I(y) \setminus I(x)} a_i(y)e_i + \\ &+ \sum_{i \in I(y) \setminus I(x)} a_i(x)e_i) = 1, \end{aligned} \quad (6)$$

где через $\bar{I}(x)$ мы обозначили множество индексов, равное $\{1, \dots, n\} \setminus I(x)$, и использовали равенство

$$I(x) \setminus I(y) = \bar{I}(y) \setminus \bar{I}(x).$$

Введем теперь следующие множества:

$$N_1 = \{i \in I(x) \setminus I(y) : a_i(x) \geq a_i(y)\},$$

$$N_2 = \{i \in I(x) \setminus I(y) : a_i(x) < a_i(y)\},$$

$$N_3 = \{i \in I(y) \setminus I(x) : a_i(y) \geq a_i(x)\},$$

$$N_4 = \{i \in I(y) \setminus I(x) : a_i(y) < a_i(x)\}$$

и воспользуемся полуаддитивностью. Тогда равенство (6) можно переписать в следующем виде:

$$\begin{aligned} E\left(x + y + \sum_{i \in I(x) \cap I(y)} (a_i(x)e_i + a_i(y)e_i) + \sum_{i \in N_1} (a_i(x) - a_i(y))e_i + \right. \\ \left. + \sum_{i \in N_3} (a_i(y) - a_i(x))e_i, \sum_{i \in I(x) \cap I(y)} (a_i(x) + a_i(y))e_i + \sum_{i \in N_2} (a_i(y) - \right. \\ \left. - a_i(x))e_i + \sum_{i \in N_4} (a_i(x) - a_i(y))e_i\right) = 1. \end{aligned} \quad (7)$$

Заметим, что множество $I(x) \cap I(y)$, $\bar{I}(y) \cap \bar{I}(x)$, $\{N_i\}_{i=1}^4$, не пересекающиеся и в объединении дают все множество индексов. В этом случае из n -мерности и равенства (7) вытекает

$$I(x + y) = (I(x) \cap I(y)) \cup N_1 \cup N_3,$$

$$a_i(x + y) = \begin{cases} a_i(x) + a_i(y), i \in I(x) \cap I(y) \\ a_i(x) - a_i(y), i \in N_1 \\ a_i(y) - a_i(x), i \in N_2 \\ a_i(x) + a_i(y), i \in \bar{I}(x) \cap \bar{I}(y) \\ a_i(x) - a_i(y), i \in N_4 \\ a_i(y) - a_i(x), i \in N_3. \end{cases}$$

Так как из (3) вытекает

$$a_i(x) = \begin{cases} a_i(x), i \in \bar{I}(x) \\ -a_i(x), i \in I(x), \end{cases}$$

то, если рассмотреть $a_i(x + y)$, получим

$$a_i(x + y) = a_i(x) + a_i(y). \quad (8)$$

Действительно, пусть $i \in I(x) \cap I(y)$, тогда $i \in I(x + y)$ и $a_i(x + y) = a_i(x) + a_i(y) = a_i(x) - a_i(y) = a_i(x) + a_i(y)$. Допустим, $i \in N_1$, значит, $i \in I(x + y)$ и $a_i(x + y) = -a_i(x) + a_i(y) = a_i(y) - a_i(x) = a_i(x) + a_i(y)$, так как $i \in I(x) \setminus I(y)$ в силу определения N_1 и т.д. Рассмотрев все шесть случаев, легко убедиться, что при любом индексе i выполняется равенство (8) для произвольных $x, y \in K$. Таким образом, функционалы $\{a_i(x)\}_{i=1}^n$ аддитивны. Покажем их однородность для $\lambda > 0$.

Действительно для произвольного $x \in K$ и $\lambda > 0$ из n -мерности и однородности следует

$$E(x + \sum_{i \in I(x)} a_i(x)e_i, \sum_{i \in I(x)} a_i(x)e_i) = 1,$$

$$E(\lambda x + \sum_{i \in I(x)} \lambda a_i(x)e_i, \sum_{i \in I(x)} \lambda a_i(x)e_i) = 1. \quad (9)$$

Последнее равенство означает, что $I(x) = I(\lambda x)$ и $a_i(\lambda x) = \lambda a_i(x)$, $i = \bar{1}, n$. Следовательно,

$$\lambda a_i(x) = a_i(\lambda x), i = \bar{1}, n, \lambda > 0.$$

Поскольку положительный конус в линейном пространстве воспроизводящий, то функционалы $a_i(x)$ могут быть продолжены до линейных на всем пространстве $\langle L, R \rangle$.

Докажем теперь одно вспомогательное утверждение.

Утверждение 1. Если предикат $E(x, y)$ обладает перечисленными в теореме свойствами, то он рефлексивен, симметричен, транзитивен.

Действительно, для любого $x \in K$ из n -мерности и аддитивности вытекает

$$E(x + \sum_{i \in I(x)} \lambda_i(x)e_i, \sum_{i \in I(x)} a_i(x)e_i) = 1,$$

$$E(x + \sum_{i=1}^n a_i(x)e_i, x + \sum_{i=1}^n a_i(x)e_i) = 1,$$

учитывая полуаддитивность, имеем $E(x, x) = 1$.

Далее, если $E(x, y) = 1$ и из рефлексивности

$E(y, y) = 1$, то при помощи аддитивности и полуаддитивности получим $E(2y, x + y) = 1$ и $E(y, x) = 1$.

Теперь допустим $E(x, y) = 1$, $E(y, z) = 1$, тогда ясно, что $E(x + y, y + z) = 1$ и $E(x, z) = 1$. Утверждение доказано.

Пусть $E(x, y) = 1$. Используя свойства теоремы и доказанное утверждение, нетрудно убедиться в правильности следующей цепочки равенств:

$$E(x + \sum_{i \in I(x)} a_i(x) e_i, y + \sum_{i \in I(x)} a_i(x) e_i) = 1, \quad (10)$$

$$E(x + \sum_{i \in I(x)} a_i(x) e_i, \sum_{i \in \bar{I}(x)} a_i(x) e_i) = 1, \quad (11)$$

$$E(y + \sum_{i \in I(x)} a_i(x) e_i, \sum_{i \in \bar{I}(x)} a_i(x) e_i) = 1. \quad (12)$$

Из единственности $I(x)$ и набора чисел $\{a_i(x)\}_{i=1}^n$ имеем $I(x) = I(y)$ и $a_i(x) = a_i(y)$, что означает $E(x, y) = 1$ тогда и только тогда, когда

$$\alpha_i(x) = \alpha_i(y), \quad i = 1, n, \quad (13)$$

поскольку цепочку равенств (10)–(13) можно провести и в обратном порядке. Все это означает, что предикат $E(x, y)$ представим в виде

$$E(x, y) = D(\alpha(x), \alpha(y)),$$

где $\alpha(x) = (\alpha_1(x), \dots, \alpha_n(x))$.

Таким образом, для того чтобы он был линейным, осталось показать, что набор линейных функционалов $\{\alpha_i(x)\}_{i=1}^n$ — линейно независим. Действительно, в противном случае можно считать,

без ограничения общности, что $\alpha_1(x) = \sum_{i=2}^n \lambda_i \alpha_i(x)$.

Однако из рефлексивности предиката $E(x, y)$, примененной к вектору e_1 , получим $E(e_1, e_1) = 1$. Следовательно, из n -мерности вытекает, что $\alpha_1(e_1) = 1, \alpha_2(e_1) = \dots = \alpha_n(e_1) = 0$. Тогда из предположения линейной зависимости функционалов $\{\alpha_i(x)\}_{i=1}^n$ вытекает, что $1 = 0$. Противоречие. Следовательно, функционалы $\{\alpha_i(x)\}_{i=1}^n$ линейно независимы, а предикат $E(x, y)$ линеен. Теорема доказана.

В доказанной теореме мы под линейностью оператора F понимаем его аддитивность и однородность. Когда $\dim L < \infty$, этого действительно достаточно. В противном случае однородность необходимо заменить на непрерывность. Однако в случае положительного конуса, как уже отмечалось ранее, возникают определенные сложности. Можно было бы, убрав однородность, сформулировать, что функционалы $\{\alpha_i\}_{i=1}^n$ и $\{\alpha_i\}_{i=1}^n$ непрерывны. Но вся трудность в том, что эти функционалы возникают в свойстве n -мерности и значит определены только на элементах положительного конуса. Любая же окрестность произвольного элемента x , принадлежащего пространству, например, типа

$L_2[a, b]$, содержит точки, не принадлежащие положительному конусу. Следовательно, для наших функционалов в классическом смысле не может быть сформулировано свойство непрерывности. Однако эту трудность для пространства $L_2[a, b]$, часто встречающегося на практике, можно преодолеть. Более того, в этом случае может быть усилен результат теоремы 1, а именно: при выполнении ее условий система, соответствующая линейному предикату, описывается набором положительных функционалов, то $f_i(x) > 0 \forall x \in K$ в равенстве (2).

Итак, допустим, что $\langle L, R \rangle = L_2[a, b]$. Без ограничения общности его можно считать $L_2[0, 1]$.

Утверждение 2. Пусть на положительном конусе K задан аддитивный функционал $f(x)$, удовлетворяющий свойству

$$\sup_{x \in K \cap S} |f(x)| < \infty,$$

где $S = \{x \in L_2[0, 1] : \|x\| \leq 1\}$ — единичная сфера, тогда его можно однозначно продолжить до линейного функционала на всем пространстве $L_2[0, 1]$. Предположим,

$$x_1(\tau) = \begin{cases} y(\tau) + c, & y(\tau) > 0 \\ c, & y(\tau) \leq 0, \end{cases}$$

$$x_2(\tau) = \begin{cases} -y(\tau) + c, & y(\tau) \leq 0 \\ c, & y(\tau) > 0, \end{cases}$$

где $c > 0$ — произвольная постоянная. Тогда $x_1, x_2 \in K$ и $y = x_1 - x_2$.

Положим $f(y) = f(x_1) - f(x_2)$. Это определение корректно, так как если $y = x_1 - x_2 = x'_1 - x'_2$, то $x_1 + x'_2 = x_1 + x'_2$. Оба элемента $x_1 + x'_2$ и $x_1 + x'_2$ принадлежат положительному конусу K , на котором функционал аддитивен, значит $f(x_1) + f(x'_2) = f(x'_1) + f(x_2)$ или $f(x_1) - f(x'_2) = f(x'_1) - f(x_2)$. Последнее равенство означает корректность определения. То, что построенное продолжение однозначно — очевидно.

Теперь возьмем любой элемент $y \in S$. Его можно представить в виде $y = x_1 - x_2$, где $x_1, x_2 \in K \cap S$. Для этого нужно соответствующим образом подобрать постоянную C . Рассмотрим $|f(y)|$, тогда $|f(y)| = |f(x_1) - f(x_2)| \leq |f(x_1)| + |f(x_2)| < \infty$.

$$\text{Отсюда } \sup_{y \in S} |f(y)| < \infty.$$

Значит, продолженный функционал $f(y)$ ограничен, следовательно, непрерывен, а непрерывный аддитивный функционал на $L_2[0, 1]$ является линейным. Утверждение доказано.

Свойства, сформулированные в утверждении 2 для функционала $f(x)$, заданного на положительном конусе K , будем называть непрерывностью. Нетрудно видеть, что если дан линейный предикат

на $K \times K$, то функционалы $\{a_i\}_{i=1}^n$ и $\{a_i\}_{i=1}^n$ непрерывны в этом смысле.

Более того, теорема 1 остается справедливой, если в ее формулировке изменить свойство одноднотности на непрерывность.

Таким образом, получен набор характеристических свойств линейных предикатов на $K \subset L_2[0,1]$.

Для обоснования положительности $\{f_i(x)\}_{i=1}^n$ докажем следующее утверждение.

Утверждение 3. Пусть $\{a_i(x)\}_{i=1}^n$ — система линейных линейно независимых функционалов, заданных из K и удовлетворяющих свойствам: для любого $x \in K$ существует номер k такой, что $\alpha_k(x) > 0$. Тогда найдется невырожденная n -мерная матрица A , которая переводит данную систему в систему $\{\beta_i(x)\}_{i=1}^n$, т.е. $\beta(x) = A\alpha(x)$, для которой $\beta_i(x) > 0$ при любых i и $x \in K$ или ядра $\{\beta_i(x)\}_{i=1}^n \in K$.

Доказательство. Обозначим через $L = \zeta = (\alpha_1, \dots, \alpha_n)$ линейная оболочка. Предположим $L \cap K = \emptyset$, тогда по теореме об отделимости выпуклых множеств в линейном пространстве (L и K таковыми являются) существует ненулевой линейный функционал $h(x)$ и число C такие, что $h(x) \geq C$ при $x \in K$ и $h(x) \leq C$ при $x \in L$.

Причем C не может быть меньше нуля, иначе бы функционал $h(x)$ нельзя было бы отдельно ограничить снизу на положительном конусе K . С другой стороны, так как любой линейный функционал на K может быть сколь угодно мал, то C не может быть и строго больше нуля. Значит $C = 0$. Более того, $h(x) > 0$ при $x \in K$, поскольку при существовании $x_0 \in K$, для которого $h(x_0) = 0$ нашелся бы элемент $x_1 \in K$ такой, что $h(x_1) < 0$. В итоге получаем, что $h \in K$ и в то же время $a_i(h) \leq 0$ при любом i . Это противоречит условию утверждения. Следовательно, $L \cap K \neq \emptyset$.

Возьмем элемент $e \in L \cap K$. Так как K — открытое множество, то e входит в K с какой-то окрестностью, поэтому для любого a_i найдется $\varepsilon_i > 0$ такое, что $\beta_i = \varepsilon_i \alpha_i + e \in K$. В силу линейной независимости системы $\{\alpha_i\}_{i=1}^n$ можно подобрать ε_i таким, что система $\{\beta_i(x)\}_{i=1}^n$ тоже будет линейно независимой. С другой стороны, $\{\beta_i(x)\}_{i=1}^n \in L \cap K$, т.е. существует невырожденная матрица A , для которой $\beta(x) = A\alpha(x)$. Утверждение доказано.

Справедлива следующая цепочка равенств

$$E(x, y) = D(\alpha(x), \alpha(y)) = D(A\alpha(x), A\alpha(y)) =$$

$$= D(\beta(x), \beta(y)), x, y \in K \subset L_2[0,1],$$

$$\alpha(x) = (\alpha_1(x), \dots, \alpha_n(x)), \beta(x) = (\beta_1(x), \dots, \beta_n(x)).$$

Переобозначив

$$F[x] = (f_1(x), \dots, f_n(x)), f_i(x) = \beta_i(x),$$

получим $E(x, y) = D(F[x], F[y])$.

В итоге нами получена теорема.

Теорема 2. Предикат $E(x, y)$, заданный на $K \subset L_2[0,1]$, линеен с положительными функционалами $\{f_i(x)\}_{i=1}^n$ тогда и только тогда, когда он обладает свойствами аддитивности, полуаддитивности, n -мерности и непрерывности.

Для завершения исследования покажем, что этот набор условий независим.

Утверждение 4. Условия аддитивности, полуаддитивности, n -мерности и непрерывности независимы.

Доказательство. Как и ранее, мы приведем примеры предикатов, обладающих всеми перечисленными свойствами, кроме одного. Относительно n -мерности и непрерывности можно взять соответствующие примеры из работы [3].

Полуаддитивность. Пусть

$$E(x, y) = D(\alpha(x), \alpha'(y)),$$

где $x, y \in K, \alpha_1(x), \dots, \alpha_n(x)$ набор линейных линейно независимых функционалов в $L_2[0,1]$ с ядрами, принадлежащими K и $\alpha_1(1) - \alpha_n(1) \neq 0$, а

$$\alpha'(y) = (\alpha_n(y), \alpha_2(y), \dots, \alpha_{n-1}(y), \alpha_1(y)).$$

Тогда если $E(x, y) = 1$ и $E(x', y') = 1$, то

$$\alpha_1(x) = \alpha_n(y), \alpha_k(x) = \alpha_k(y), k = 2, \dots, n-1, \alpha_n(x) = \alpha_1(y),$$

$$\alpha_1(x') = \alpha_n(y'), \alpha_k(x') = \alpha_k(y'), k = 2, \dots, n-1, \alpha_n(x') = \alpha_1(y').$$

Отсюда

$$\alpha_1(x + x') = \alpha_n(y + y'),$$

$$\alpha_k(x + x') = \alpha_k(y + y'), k = 2, \dots, n-1,$$

$$\alpha_n(x + x') = \alpha_1(y + y'),$$

$$\alpha_1(x + y') = \alpha_n(y + x'),$$

$$\alpha_k(x + y') = \alpha_k(y + x'), k = 2, \dots, n-1,$$

$$\alpha_n(x + y') = \alpha_{11}(y + x'),$$

т.е. $E(x + x', y + y') = E(x + y', y + x') = 1$. Таким образом, аддитивностью этот предикат обладает. В справедливости для него n -мерности и непрерывности можно убедиться непосредственно.

Однако свойству полуаддитивности построенный нами предикат не удовлетворяет. Действительно, допустим $C_1, C_2 = \text{const} > 0$ такие, что $E(C_1, C_2) = 1$. Выберем произвольное число λ на интервале $(0, \min(C_1, C_2))$ и положим $x = C_1 - \lambda, y = C_2 - \lambda, z = \lambda$. Тогда $x, y, z \in K$ и $E(x + z, y + z) = 1$,

Однако

$$\alpha_1(x) - \alpha_n(y) = \alpha_1(C_1) - \alpha_1(\lambda) - \alpha_n(C_2) + \alpha_n(\lambda) = \\ = \lambda(\alpha_n(1) - \alpha_1(1)) \neq 0$$

т.е. $E(x, y) = 0$.

Аддитивность. Положим

$$E(x, y) = D(\alpha'(x), \alpha(y)),$$

где $x, y \in K$, а

$$\alpha'(x) = (\alpha_1(x) + 1, \dots, \alpha_n(x)), \alpha(y) = (\alpha_1(y), \dots, \alpha_n(y)).$$

Этот предикат полуаддитивен, так как если $E(x+z, y+z)=1$, то $\alpha_1(x+z)+1=\alpha_1(y+z)$ или $\alpha_1(x)+1=\alpha_1(y)$, т.е. $E(x, y)=1$.

Он также очевидным образом непрерывен и n -мерен, однако аддитивности нет, поскольку не аддитивен функционал $\alpha_1(x)+1$. Утверждение доказано.

Выводы

Рассмотрен наиболее распространенный случай ограничения множества входных сигналов в виде конуса линейного пространства. Показано, что подобное сужение приводит к ряду принципиальных отличий, которые не позволяют автоматически перенести условия существования линейных предикатов и свести всё к обычной ситуации. Найдены новые характеристические свойства, обеспечивающие существование линейных предикатов при данном ограничении.

Список литературы:

1. Мальцев А.И. Алгебраические системы. — М.: Наука, 1970. — 392 с.
2. Четвериков Г.Г., Дударь З.В., Вечирская И.Д. Дискретные структуры. — Харьков: НУРЭ, 2015. — 322 с.
3. Машталир В.П., Шляхов В.В., Яковлев С.В. Групповые структуры на фактор-множествах в задачах классификации // Кибернетика и системный анализ. — 2014. — №4. — С. 27–42.
4. Герасин С.Н., Шляхов В.В. О внешней и внутренней согласованности произвольных n -арных отношений // Доп. НАН Украины. — 2006. — №1. — С. 77–82.
5. Бондаренко Н.Ф., Машталир В.П., Шляхов В.В. Мультигруппы, индуцируемые произвольным отношением // Доп. НАН Украины. — 2012. — № 4. — С. 39–42.
6. Каграманян А.Г., Машталир В.П., Шляхов В.В. Условия существования тернарной мультисистемы с единым носителем // Вісник національного університету. — Серія: Математичне моделювання. Інформаційні технології. — 2011. — № 960, Вип. 16. — Харків: ХНУ імені В.Н. Каразіна. — С. 148–158.

Поступила в редколлегию 24.11.2016

УДК 519.876.5

**Д.И. Чумаченко¹, С.В. Яковлев²**¹Национальный аэрокосмический университет им. Н. Е. Жуковского
«Харьковский авиационный институт»,
г. Харьков, Украина, dichumachenko@gmail.com²Национальный аэрокосмический университет им. Н. Е. Жуковского
«Харьковский авиационный институт», г. Харьков, Украина, svsyak@mail.ru

О НЕЧЕТКИХ РЕКУРРЕНТНЫХ ОТОБРАЖЕНИЯХ ПРИ МУЛЬТИАГЕНТНОМ МОДЕЛИРОВАНИИ ПОПУЛЯЦИОННОЙ ДИНАМИКИ

В статье рассмотрены проблемы анализа мультиагентных моделей популяционной динамики из-за ряда неопределенностей, связанных с переменными, граничными условиями, начальными состояниями, значениями параметров и т.д. Для решения данной проблемы разработана лингвистическая нечеткая модель, позволяющая описать системы популяционной динамики более реалистичным способом. Популяционная динамика описывается набором правил, каждое из которых предполагает вход и выход в виде нечетких множеств или нечетких функций, применяемых итеративно. Сложность описания процессов систем популяционной динамики, наличие алгоритмов фаззификации и дефаззификации, а также использование нечетких множеств и лингвистических переменных приводят к необходимости разработки новых методов анализа подобных систем.

СИСТЕМА ПОПУЛЯЦИОННОЙ ДИНАМИКИ, МУЛЬТИАГЕНТНАЯ МОДЕЛЬ, ЛИНГВИСТИЧЕСКАЯ НЕЧЕТКАЯ МОДЕЛЬ, РЕКУРРЕНТНАЯ МОДЕЛЬ, НЕЧЕТКИЕ МНОЖЕСТВА, ПРОДУКЦИОННАЯ СИСТЕМА

Введение

Динамика любой системы описывается математической моделью, которая отражает зависимости между тремя множествами переменных: входа, выхода и состояния. Основное свойство любой динамической системы заключается в том, что ее поведение в любой момент времени зависит не только от переменных, действующих на нее в данный момент времени, но и от переменных, действовавших на нее в прошлом.

Мультиагентное моделирование является одним из наиболее актуальных способов создания имитационных моделей. При решении задачи моделирования динамических систем исследователи сталкиваются с проблемой описания сложной структуры взаимодействующих элементов, которые не однородны и не подлежат типизации по своим свойствам. Аналитические модели и традиционные дискретно-событийные модели оказываются не эффективными, а возможность их применения в определенных системах отсутствует. В таких случаях прибегают к применению мультиагентного моделирования, выстраиваемого снизу вверх, т.е. глобальная динамика системы формируется за счет взаимодействия автономных моделей. Быстро увеличивающаяся вычислительная мощность персональных компьютеров, которая допускает моделирование большого числа независимых объектов, является существенным катализатором развития мультиагентного моделирования в применении к системам популяционной динамики.

Исследователи систем популяционной динамики сталкиваются с целым рядом препятствий при

попытке проверить свои модели, в частности, из-за ряда неопределенностей, связанных с переменными, граничными условиями, начальными состояниями, значениями параметров и т.д. [1]. На практике для исследования доступны только данные о зарегистрированных процессах, кроме того, такие данные демонстрируют расплывчатость в определении таких понятий, как факторы риска, опасности, силы воздействия, контактных шаблонах и т.д. [2]. Таким образом, возможным альтернативным подходом может быть комбинация методов нечеткой логики и нелинейных динамических систем с целью обеспечения всестороннего анализа и разработки инструментов прогнозирования в системах популяционной динамики.

Известно, что с помощью продукционных моделей представляется возможным естественно описать декларативный опыт человека, его интуицию и логику поведения. Зачастую целесообразно также использовать нечеткие лингвистические переменные, с помощью которых можно адекватно отразить приблизительное описание взаимодействующих элементов, в том случае, когда точное детерминированное описание отсутствует. При этом следует учесть, что многие нечеткие категории, описанные лингвистически, зачастую не менее информативны, чем точное описание.

1. Современное состояние вопроса

Основываясь на том, что нечеткие динамические системы являются относительно новой областью исследований, их основная идея заключается в расширении стандартной динамической

системы, модель которой построена при помощи дифференциальных уравнений либо другого подхода в теоретические рамки нечеткого множества. Эти методы позволяют учитывать неопределенности, связанные с переменными, параметрами, граничными условиями и начальными состояниями, моделировать их эволюцию, соблюдая основные правила и закономерности динамики системы.

Нечеткие модели основаны на концепции нечеткого разбиения информации и могут быть классифицированы в две общие группы в зависимости от того, как представлена информация: 1) лингвистические модели, у которых наиболее известным примером является модель типа Мамдани, и 2) модель Такаги-Сугено. Обе модели основаны на использовании нечетких правил и лингвистических переменных. Тем не менее, лингвистические модели являются качественным описанием поведения системы с использованием естественного языка, а модели Такаги-Сугено — это комбинация нечетких и стандартных структур.

Успешное применение нечеткой лингвистической модели в моделировании контроллеров показало, что это наиболее прикладная структура. Нечеткая лингвистическая модель основана на приближенных рассуждениях, которые обеспечивают основу для построения гипотез с неточной информацией с помощью адекватных механизмов логического вывода [3]. Данная модель может быть определена как экспертная система, так как она включает в себя базу знаний и механизм логического вывода, оба из которых основываются на экспертных знаниях человека. Большинство нечетких приложений основаны на нечетких лингвистических системах. Они широко используются в разработке нечетких контроллеров медицинских аппаратов, оценке риска и в диагностических системах [4-7]. В эпидемиологии также существуют несколько нечетких лингвистических систем [8, 9].

Нечеткие модели, основанные на правилах, имеют простую структуру и состоят из четырех основных компонентов: 1) модуль фаззификации, который переводит четкие входы (классические измерения) в нечеткие значения при помощи лингвистических переменных; 2) нечеткая база правил If-Then, которая состоит из набора условных нечетких суждений; 3) метод логического вывода, который применяет нечеткие механизмы для получения результатов или, другими словами, способ вычисления с нечеткими правилами; 4) модуль дефаззификации, который переводит нечеткие выходы обратно к четким значениям, если это необходимо.

В работе [10] была предложена структура для решения систем линейных дифференциальных

уравнений для одного класса нечетких множеств, основанная на α -уровне. Тем не менее, предложенные методы трудно применять в системах популяционной динамики, так как такие модели обладают явными нелинейностями и должны быть рассмотрены по-другому. Эти нелинейности обусловлены тем, что ход эпидемического процесса динамической системы зависит, помимо всего прочего, от доли людей, находящихся в различных состояниях, которые, по своей природе неопределенны и, следовательно, являются идеальными объектами для анализа при помощи нечеткой логики.

В работе [11, 12] предложен новый подход в описании динамики экологической модели в дифференциальных уравнениях, описывающий динамическую систему с использованием нечетких параметров. В этом случае решением системы уравнений является, так называемое, нечеткое ожидаемое значение. Применение этого подхода в системах популяционной динамики не является очевидным потому, что некоторые детали и параметры могут трактоваться неоднозначно. Несмотря на это, данный метод является возможным способом смоделировать систему популяционной динамики более реалистично.

2. Постановка задачи исследования

Существующие подходы показали некоторые теоретические трудности, в связи с которыми, учитывая особенности математики нечеткой логики, трудно показать эффективные практические результаты. Для того, чтобы исследовать, насколько нечеткая логика может описывать системы популяционной динамики более реалистичным способом, в данном исследовании разработана лингвистическая нечеткая модель, примененная к мультиагентной модели, описанной в работах [13, 14].

3. Построение нечеткой лингвистической модели популяционной динамики

Предлагаемый подход состоит в том, что системная динамика описывается набором правил, которые применяются итеративно. Каждое правило предполагает вход и выход в виде нечетких множеств или нечетких функций. Из эмпирического опыта группы экспертов мы можем создать нечеткую функцию принадлежности для каждой переменной и/или параметра, а также лингвистические правила, которые регулируют динамику системы. Таким образом, нечеткая модель состоит из набора правил и соответствующего логического вывода конечного автомата. Лингвистическая модель принимает следующую форму:

IF $U = B_1$ AND $W_1 = A_{11}$ AND ... AND $W_n = A_{1n}$
 THEN $\bar{W}_1 = \hat{A}_{11}$ AND ... AND $\bar{W}_n = \hat{A}_{1n}$ AND $V = D_1$

IF $U = B_2$ AND $W_1 = A_{21}$ AND ... AND $W_n = A_{2n}$
 THEN $\bar{W}_1 = \hat{A}_{21}$ AND ... AND $\bar{W}_n = \hat{A}_{2n}$ AND $V = D_2$

...

IF $U = B_m$ AND $W_1 = A_{m1}$ AND ... AND $W_n = A_{mn}$
 THEN $\bar{W}_1 = \hat{A}_{m1}$ AND ... AND $\bar{W}_n = \hat{A}_{mn}$ AND $V = D_m$

где U – входная переменная, W_i – переменные состояний системы, V – выходная переменная, $\bar{W}_i$ – переменные состояний системы после каждой итерации, B_i и A_{ij} – входные нечеткие множества, D_i и $\hat{A}_{ij}$ – выходные нечеткие множества.

Поэтому, выбирая подходящий метод логического вывода и, при необходимости, метод дефазификации, на каждом шаге после запуска модели рассчитывается значение переменной состояния, которое будет входным параметром системы на следующем шаге, и так далее, итеративно. Отсюда следует, следующее:

$$U(l+1) = V(l) \quad (1)$$

$$W_i(l+1) = \bar{W}_i(l) \quad (2)$$

где $(l+1)$ – следующий шаг после l .

Существуют и другие нечеткие динамические системы, и выбор конкретной модели зависит от типа имеющейся информации о системе. Иногда часть правил поведения системы известна заранее, тогда правила будут принимать следующий вид:

IF $U(l) = B_m$ AND $W_1(l) = A_{m1}$ AND ... AND

$$W_n(l) = A_{mn}$$

THEN $y(l+1) = f(U(l), W_1(l), \dots, W_n(l))$;

где $y(l+1)$ – некоторая априорная функция, известная из системной динамики.

Разработка нечетких моделей с экспертами в различных областях требует междисциплинарных отношений [15]. Для правильного применения знаний экспертов, важным является построение ими нечетких множеств. В общем виде, нечеткие множества эпидемиологических систем, построенные экспертами, не показывают правильного поведения. Они имеют тенденцию к асимметричной и нерегулярной динамике, отличающейся от поведения нечетких множеств в инженерных областях. Кроме того, у экспертов существуют проблемы с пониманием природы системной динамики, правил поведения и, как следствие, модели в целом [16]. Кроме того, создание следствий из правил поведения является более трудоемкой задачей для эксперта, чем анализ прошлой статистики

динамической системы, т.к. эксперт должен из-учить вопрос о динамике системы, учитывая все факторы, и сформулировав конкретную зависимость, соответствующую функции принадлежности. Кроме того, чтобы проводить анализ прошлой статистики, которая может влиять на дальнейшее поведение динамической системы, эксперту достаточно классифицировать переменные функции принадлежности. Поэтому, в целом, эксперт имеет больше возможностей для разработки предпосылок к развитию эпидемического процесса, чем следствий. Учитывая данные особенности, метод, позволяющий разрабатывать апостериорные лингвистические правила, может означать важный прогресс в моделировании динамических систем, которые имеют высокий уровень неопределенности, неточности или неясности в определении переменных и параметров. Кроме того, необходимо учесть, что эпидемические системы – динамические, неавтономные и открытые, и, следовательно, имеют возможности к вводу-выводу достаточно больших объемов данных для построения гибридных моделей крайне редко.

Таким образом, можно формализовать нечеткое рекуррентное отображение [17], которое определяется множеством наборов правил $R = \{r_1, r_2, \dots, r_N\}$, связывающих значения переменных состояния $(x_1, \dots, x_N)$ динамической системы в текущий τ и будущий t моменты времени:

$$\begin{pmatrix} x_1 \\ x_2 \\ \dots \\ x_N \end{pmatrix}_t = \begin{pmatrix} r_1 \\ r_2 \\ \dots \\ r_N \end{pmatrix} \circ \begin{pmatrix} x_1 \\ x_2 \\ \dots \\ x_N \end{pmatrix}_\tau \quad (3)$$

Наборы правил (1) имеют вид

$$r_i : \begin{cases} \text{IF } x_1 = A_1^1 \text{ AND } x_2 = A_2^1 \dots \\ \left[\text{AND } x_k = A_k^1 \right] \dots x_N = A_N^1 \Big|_\tau \\ \text{THEN } x_i = B^1 \Big|_t \\ \text{IF } x_1 = A_1^2 \text{ AND } x_2 = A_2^2 \dots \\ \left[\text{AND } x_k = A_k^2 \right] \dots x_N = A_N^2 \Big|_\tau \\ \text{THEN } x_i = B^2 \Big|_t \\ \dots \dots \dots \\ \text{IF } x_1 = A_1^{K_i} \text{ AND } x_2 = A_2^{K_i} \dots \\ \left[\text{AND } x_k = A_k^{K_i} \right] \dots x_N = A_N^{K_i} \Big|_\tau \\ \text{THEN } x_i = B^{K_i} \Big|_t \end{cases}, \quad (4)$$

где K_i – количество правил в наборе r_i ; элементы в квадратных скобках являются необязательными, т.е. не все переменные состояния могут быть задействованы в правиле; A_i, B – значения лингвистических переменных из соответствующих терм-множеств.

Нетрудно показать, что количество правил набора находится в диапазоне

$$0 \leq K_i \leq \prod_{i=1}^N \text{card}(S(x_i)), \quad (5)$$

где $\text{card}(S(x_i))$ — мощность терм-множества лингвистической переменной x_i . Допускаются также и пустые наборы правил.

Анализ динамических лингвистических систем, представленных в общей форме (3), (4) довольно затруднителен ввиду высокой размерности пространства состояний и нелинейного имплицитивного характера взаимосвязи событий, переводящих систему из одного состояния в другое. Поэтому актуальным представляется нахождение такой формы описания динамики процесса, которая позволяла бы решать задачи анализа и синтеза формальными методами. Обычно на практике применяют лингвистическое описание в виде правил

IF $X_k = (x_1 = nb, x_2 = pm, \dots, x_n = ze)$ AND $U_k = (u_1 = pm, u_2 = nb, \dots, u_m = nm)$,

THEN $X_{k+1} = (x_1 = pb, x_2 = ps, \dots, x_n = pb)$,
отражающих отношения изменения состояния системы в зависимости от входных воздействий

$$X_{k+1} = X_k \circ U_k, \quad (6)$$

в котором $X_k = (x_1, x_2, \dots, x_n)_k$ — обобщенный вектор состояния системы, а $U_k = (u_1, u_2, \dots, u_m)_k$ — обобщенный вектор управляющих воздействий, значения которых представляют собой лингвистические переменные из заданного терм-множества $S = \{nb, nm, \dots, ze, \dots, pm, pb\}$, где nb — *negative big*, nm — *negative middle*, ze — *zero*, pm — *positive middle*, pb — *positive big* — нечеткие множества с заданными функциями принадлежности [18].

Отношение (6) можно представить в виде сети переходов обобщенных лингвистических состояний (вершины графа) под действием обобщенных лингвистических управляющих воздействий (ребра графа). Если N — размерность вектора состояния X , P — размерность вектора управления U , M — мощность терм-множества лингвистических переменных S , то максимально возможное количество вершин сети (состояний системы) равно $M \cdot N$, а количество дуг, соединяющих эти вершины (управляющих воздействий), — $M \cdot N \cdot (MN - 1) / 2 \cdot MP$.

Анализ подобных систем на основе имитационного моделирования и синтез оптимальных правил (значения ребер сети) представляют собой комбинаторную проблему.

На практике вместо общего вида отображения (6) используют его частные формы, когда векторы X , Y , U являются скалярными лингвистическими переменными. В случае, когда консеквенты правил являются лингвистическими, рассматриваемая

модель представляет собой модель Мамдани, если функциональными, — модель Сугено.

Как правило, динамическое поведение таких систем описывается в виде таблиц лингвистических правил, связывающих управляющие воздействия U и выходы (либо состояния) объекта X . Пример такого отображения представлен в табл. 1.

Таблица 1

Таблица лингвистических правил $X_{k+1} = X_k \circ U_k$

$U_k \backslash X_k$	nb	nm	ze	pm	pb
nb	nb	nb	nb	nm	ze
nm	nb	nb	nm	ze	pm
ze	nb	nb	ze	pb	pb
pm	nm	ze	pm	pb	pb
pb	ze	pm	pb	pb	pb

Основная проблема анализа модели (6) — отсутствие формальных методов в пространстве лингвистических состояний, аналогичных числовым моделям и методам анализа и синтеза в евклидовом пространстве состояний, что затрудняет решение задач анализа устойчивости динамической модели, синтеза оптимальных систем управления и других задач. Известные методы анализа нечетких динамических систем основаны либо на исследовании функций принадлежности нечетких множеств, либо на анализе переходов в расширяющемся пространстве состояний, либо на эвристических методах лингвистической динамики [19]. Существенным фактором является также размерность набора лингвистических правил, при котором число возможных правил экспоненциально увеличивается с числом входных величин.

Сложность описания процессов систем популяционной динамики, наличие эвристических алгоритмов фаззификации и дефаззификации, а также использование нечетких множеств и лингвистических переменных приводят к необходимости разработки новых методов анализа подобных систем.

Выводы

Предложенные в статье подходы к исследованию систем популяционной динамики позволяют применять унифицированное описание процессов различной природы в виде продукционного набора правил. Дальнейшие исследования будут направлены на развитие методов идентификации хаотической динамики с произвольным количеством правил. Предполагается также исследование в направлении векторных нечетких моделей, характерных для описания динамических систем с высокой размерностью состояний.

Список литературы:

1. Osborne M. J. A Course in Game Theory [Text] / Osborne M. J., Rubinstein A. — Massachusetts, Cambridge: The MIT Press, 2014. — 352 p.
2. Santos F. S. Fuzzy Dynamical Model

- of Epidemic Spreading Taking into Account the Uncertainties in Individual Infectivity [Text] / F. S. Santos, N. R. S. Ortega, D. M. T. Zanetta, E. Massad // *Advances in Technological Applications of Logical and Intelligent Systems*. — IOS Press, 2009. — P. 180-193. **3.** *Massad E.* Fuzzy Logic in Action: Applications in Epidemiology and Beyond [Text] / E. Massad, N. R. S. Ortega, L. C. de Barros, C. J. Struchiner. — Springer-Verlag Berlin Heidelberg, 2008. — P. 181-206. **4.** *Yager R. R.* Essentials of fuzzy modeling and control [Text] / R. R. Yager, D. P. Filev. — New York: Wiley, 1994. — 408 p. **5.** *Mahfouf M.* Physiological Modelling and Fuzzy Control of Anaesthesia via Vaporisation of Isoflurane by Liquid Infusion [Text] / M. Mahfouf, A. J. Asbury, D. A. Linkens // *International Journal of Simulation- Systems, Science and Technology*, 2(1). — 2001. — P. 55-66. **6.** *Nascimento L. F. C.* Fuzzy linguistic model for evaluating the risk of neonatal death [Text] / L. F. C. Nascimento, N. R. S. Ortega // *Revista de Saude Publica*. Volume 36. Issue 6. — Faculdade de Saude Publica da Universidade de Sro Paulo, 2002. — P. 686-692. **7.** *Castanho M. A. R. B.* Membrane-active Peptides: Methods and Results on Structure and Function [Text] / M. A. R. B. Castanho. — International University Line, 2010. — 635 p. **8.** *Duarte A. P.* Cellulose acetate reverse osmosis membranes: Optimization of preparation parameters [Text] / A. P. Duarte, J. C. Bordado, M. T. Cidade // *Journal of Applied Polymer Science*. Volume 103. Issue 1. — Wiley Periodicals, 2006. — P. 134-139. **9.** *Tanaka H.* Fuzzy modeling of electrical impedance tomography images of the lungs [Text] / H. Tanaka, N. R. S. Ortega, M. S. Galizia, J. B. Borges, M. B. P. Amato // *Clinics*. Volume 63. Issue 3. — Faculdade de Medicina/USP, 2008. — P. 363-370. **10.** *Jafelice R. M.* Fuzzy modeling in symptomatic HIV virus infected population [Text] / R. M. Jafelice, L. C. de Barros, R. C. Bassanezi, F. Gomide // *Bulletin of Mathematical Biology*. Volume 66. Issue 6. — Springer, 2004. — P. 1597-1620. **11.** *Barros L. C.* A First Course in Fuzzy Logic, Fuzzy Dynamical Systems, and Biomathematics [Text] / L. C. Barros, R. C. Bassanezi, W. A. Lodwick. — Springer-Verlag Berlin Heidelberg, 2017. — 297 p. **12.** *Barros L. C.* The SI epidemiological models with a fuzzy transmission parameter [Text] / L. C. Barros, M. B. F. Leite, R. C. Bassanezi // *International Journal of Computational Mathematical Applications*. Volume 45. — 2003. — P. 1619-1628. **13.** *Чумаченко Д.І.* Інформаційна технологія епідеміологічного нагляду [Текст] / Д.І. Чумаченко, Т.О. Чумаченко // Інформаційні технології та інновації в економіці, управлінні проектами і програмами : [монографія] / за заг. ред. В.О. Тимофєєва, І.В. Чумаченко. — Харків: ФОП Панов А.М., 2016. — С.368 — 379. **14.** *Chernyshev Y.* Development of intelligent agents for simulation of hepatitis B epidemic process [Text] / Y. Chernyshev, D. Chumachenko, A. Tovstik // *Proceedings of East West Fuzzy Colloquium 2013 (20th Zittau Fuzzy Colloquium, September 25 — 27, 2013)*. — Institut fur Prozesstechnik Prozessautomatisierung und Messtechnik, 2013. — P.161 — 168. **15.** *Ortega N.* Fuzzy gradual rules in epidemiology [Text] / N. Ortega, L. C. Barros, E. Massad // *Kybernetes*. Vol. 32. Iss. 3. — 2015. — P. 460-477. **16.** *Ortega N. R. S.* Fuzzy dynamical systems in epidemic modeling [Text] / N. R. S. Ortega, P. C. Sallum, E. Massad // *Kybernetes*. Vol. 29. Iss. 2. — 2000. — P. 201-218. **17.** *Соколов А. Ю.* Анализ хаотической динамики в нечетких рекуррентных моделях [Текст] / А. Ю. Соколов, М. Вагенкнехт // *Радіоелектронні і комп'ютерні системи*. № 5. — 2007. — С. 54—65. **18.** *Кофман А.* Введение в теорию нечетких множеств [Текст] / А. Кофман. — М.: Радио и связь, 1982. — 432 с. **19.** *Поспелов Д. А.* Ситуационное управление. Теория и практика [Текст] / Д. А. Поспелов. — М.: Наука, 1986. — 288 с.

Поступила в редколлегию 28.10.2016

УДК 004.82



И.В. Левыкин

ХНУРЭ, г. Харьков, Украина ihor.levykin@nure.ua

РАЗРАБОТКА ОБОБЩЕННОЙ МОДЕЛИ ПРОЦЕССА РЕШЕНИЯ ЗАДАЧИ С ИНТЕРВАЛЬНЫМ ПРЕДСТАВЛЕНИЕМ ВРЕМЕНИ

В статье предложена обобщённая модель процесса решения задачи как элемента прецедента. Модель охватывает уровни событийного, дискретного и интервального представления процесса. Предложенная модель обеспечивает возможность построения интервальной модели процесса путем дополнения дискретного уровня представления интервалами выполнения действий.

БИЗНЕС-ПРОЦЕСС, ИНТЕРВАЛ ОЖИДАНИЙ, ПРОЦЕССНОЕ УПРАВЛЕНИЕ

Введение

Общие модели представления, методы построения и корректировки прецедентов в рамках задачи процессного управления полиграфическим предприятием, а также интервальные операции на процессах являются теоретическим базисом для разработки методов и технологий построения и адаптации прецедентов бизнес-процессов управления полиграфическим производством [1,2].

Основное преимущество прецедентного подхода состоит в формализации и последующем использовании имеющегося практического опыта для построения прецедента – аналога решения таких классов задач, для которых применение традиционных аналитических методов сопряжено со значительными трудностями [3].

В процессно-ориентированных информационных системах каждая реализация бизнес-процесса (БП) записывается в виде последовательности событий в логе процесса. Поэтому при построении прецедента бизнес-процесса для данного класса задач целесообразно использовать методы интеллектуального анализа процессов [4,5].

Традиционные методы анализа логов направлены на построение событийной модели бизнес-процесса с жестко предопределенной последовательностью действий [6,9], а записанные в логах события отражают начало либо завершение действий в такой workflow – модели.

1. Постановка задачи

Целью данной работы является формализация взаимосвязей между структурными элементами каждого уровня, а также связей между уровнями описания процесса с последующей разработкой трехуровневой обобщенной модели бизнес-процесса решения задачи, включающая в себя событийное, дискретное и интервальное представления последовательности действий процесса.

2. Проблематика процессного управления полиграфическим производством

В данной статье рассматривается задача «разработка обобщенной модели процесса решения

задачи с интервальным представлением времени» процессного управления полиграфическим производством.

В то же время общая задача процессного управления полиграфическим производством включает в себя, помимо построения модели действий процесса, подзадачи выделения подпроцессов с учетом имеющихся ресурсов, поиска узких мест процесса и прогнозирования длительности его выполнения, а также адаптации процесса во времени его выполнения с учётом требований к вновь поступившим заказам.

При решении данных задач необходимо учитывать продолжительность как отдельных действий процесса, так и подпроцессов, и всего процесса в целом [10]. Продолжительность выполняемых процессом одинаковых действий может отличаться в различных реализациях процесса в зависимости от состояния внешней среды. Из этого следует, что для выявления узких мест и последующей оценки продолжительности выполнения БП необходимо использовать интервальную оценку времени выполнения действий процесса, а так же интервальные операции над событиями, записанными в его лог.

При построении модели каждого процесса необходимо учитывать его взаимодействие с другими выполняющимися БП. Иными словами, в качестве базовых составляющих внешней среды для текущего процесса в полиграфии являются остальные выполняющиеся бизнес-процессы.

Отличие бизнес-процессов полиграфического производства (БПП) состоит в необходимости обеспечения своевременного выполнения асинхронно поступающих и параллельно выполняющихся заказов. Поэтому необходимо в первую очередь рассматривать не отдельные процессы, а систему взаимодействующих процессов в целом.

Существует общее множество операций (действий, процедур процесса), которые используются различными производственными процессами и, возможно, в различной последовательности. Это означает, что каждый производственный процесс в

данной сфере использует подмножество из общего набора операций в зависимости от типа обрабатываемого объекта (книга, журнал, буклет и т.п.).

Аналогично адаптируются процессы планирования, административные, сервисного обслуживания. В зависимости от вида обрабатываемого объекта изменяется требуемый отбор показателей качества продукции и процесса, а также последовательность процедур процессов управления качеством.

Для реализации идентичных операций в различных процессах используется одно и то же или однотипное оборудование. Поэтому для БПП характерны специфические ресурсные ограничения: «узкие места» процесса не полностью зависят от его workflow — последовательности. Они могут возникать в различных местах процесса в зависимости от текущих заявок и имеющихся ресурсных ограничений (в первую очередь по полиграфическому оборудованию).

Иными словами, при априорном анализе модели лишь одного процесса до его запуска очень сложно выявить операции, при выполнении которых возникнет недостаток производственных мощностей. Они будут динамически меняться, проявляться в новых фрагментах БПП во время его выполнения в зависимости от текущей загрузки отдельных единиц оборудования и сроков выполнения заказов для одновременно выполняющихся процессов.

Таким образом, окружение каждого производственного процесса в полиграфии непосредственно влияет на эффективность использования ресурсов и, следовательно, возможность достижения целей процесса.

Локальные цели производственных процессов в полиграфии могут быть связаны со сроками поставки продукции и с ее качеством, снижением затрат и т.п. Однако в условиях конкуренции набор локальных целей обычно с одной глобальной — своевременной поставке продукции при и заданных ограничениях по качеству, затратам и времени.

Из изложенного выше можно сделать вывод, что динамически изменяющиеся ресурсные ограничения непосредственно влияют на достижение глобальной цели процесса, поскольку изменяется продолжительность отдельных его операций. Поэтому в рамках технологии производства последовательность действий БПП может быть частично изменена с тем, чтобы достичь цели процесса при текущих ограничениях на характеристики оборудования.

Следовательно, БПП целесообразно рассматривать как процессы с переменной длительностью и изменяемой последовательностью операций, действующих в окружении таких же процессов и

конкурирующие с ними за ресурсы во время выполнения.

Новые процессы могут запускаться одновременно с выполнением существующих БПП. При этом последовательность операции зависит от типа обрабатываемого объекта, а количество таких процессов на предприятии ограничено возможностью существующего оборудования.

Рассмотренные свойства бизнес-процессов полиграфического производства затрудняют применение как аналитических методов, так и методов имитационного моделирования при планировании совокупности БПП.

В связи с этим необходимы эффективные модели и методы, позволяющие практически в реальном времени оценить возможность достижения темпоральных целей процессов при пополнении портфеля заказов во время выполнения уже имеющихся БПП и, по результатам такой оценки, принять новый заказ (отказаться от заказа, изменить сроки заказа и т.п.).

Таким образом, общая проблематика процессного управления полиграфическим производством связана с несоответствием между существующими подходами к управлению бизнес-процессами и выявленными свойствами БПП, а именно:

- переменная длительность и изменяемая последовательностью операций процесса;
- функционирование набора процессов с однотипными операциями;
- динамически изменяющиеся «узкие места», связанные с ограничениями по оборудованию.

Указанные выше особенности БПП, то есть использование упорядоченных подмножеств операций из общего набора возможных операций полиграфического производства в зависимости от типа выпускаемой продукции требуют разработки модели процесса решения задачи с интервальным представлением времени.

3. Модель процесса решения задачи с интервальным представлением времени

При решении задач применения и корректировки прецедента-аналога, в особенности при параллельном решении нескольких однотипных задач с использованием одного прецедента возникает проблема оценки продолжительности решения функциональной задачи с тем, чтобы она удовлетворяла ограничениям предметной области. Для решения этой проблемы необходимо перейти от событийного описания процесса, а также моделей процесса с дискретным временем к моделям с интервальным представлением времени, содержащим оценку времени выполнения отдельных операций процесса.

На событийном уровне данный процесс представляется набором трасс, для каждой из которых заданы отношения перехода. Данное отношение обладает свойствами рефлексивности и транзитивности.

Следовательно, на любой трассе π_k задается отношение строгого порядка. Тогда трасса процесса определяется следующим образом.

Трасса процесса π_k представляет собой упорядоченное подмножество E_k множества событий E процесса решения задачи:

$$\begin{aligned}\pi_k &= \langle E_k, \succ \rangle, \\ E_k &= \{e_{k,i}\}, \\ \forall e_{k,i} \exists e_{k,i+1} \neq e_{k,i} : e_{k,i} \succ e_{k,i+1}, \quad E_k &\subseteq E, i = \overline{1, I-1}.\end{aligned}\quad (1)$$

Для произвольной пары несовпадающих событий $e_{k,j}, e_{k,i}$ трассы π_k , связанных последовательностью переходов, будет выполняться отношение совершенного строгого порядка $\gg$:

$$\begin{aligned}\forall (e_{k,j} \neq e_{k,i}) \in \pi_k \\ (e_{k,i} \gg e_{k,j} \mid e_{k,i} \succ \dots \succ e_{k,j}) \vee (e_{k,j} \gg e_{k,i} \mid e_{k,j} \succ \dots \succ e_{k,i}),\end{aligned}\quad (2)$$

где $\gg$ – транзитивное замыкание отношения перехода

Покажем, что транзитивное замыкание отношения перехода, является отношением строгого порядка.

Отношение $\gg$ обладает свойствами транзитивности, антирефлексивности и асимметричности. Первые два свойства вытекают из свойств исходного отношения перехода $\succ$ и выражения (2).

Третье свойство определяется на основе свойства антирефлексивности для отношения перехода. Действительно, $(e_{k,i} \succ e_{k,j}) \wedge (e_{k,j} \succ e_{k,i}) \Rightarrow e_{k,j} = e_{k,i}$, что противоречит (1).

Для любой пары неравных событий $(e_{k,i}, e_{k,j})$ либо i – событие будет предшествовать j – событию, либо наоборот, поскольку временные метки (и соответственно индексы) у них упорядочиваются по мере выполнения процесса:

$$\forall e_{k,i}, e_{k,j} \in \pi_k \quad e_{k,i} \gg e_{k,j} \vee e_{k,j} \gg e_{k,i}. \quad (3)$$

Следовательно, транзитивное замыкание отношения перехода определяется для любой пары событий на трассе π_k и представляет собой отношение совершенно строгого порядка.

Таким образом, на трассе π_k процесса решения задачи R заданы отношения строгого порядка $\succ$ и совершенного строгого порядка $\gg$, т.е. π_k является дважды упорядоченным множеством:

$$\pi_k = \langle E_k, R \rangle, R = \{\succ, \gg\}, \quad (4)$$

где E_k – подмножество событий процесса, зафиксированных на трассе π_k .

Известно что исходный лог Π содержит набор трасс процесса, на которых могут возникать одни и те же события. Поэтому лог является частично упорядоченным множеством:

$$\Pi = \langle E, R \rangle, R = \{\succ, \gg, \geq\}, \quad (5)$$

где $E = \bigcup_k E_k$ – множество всех событий, которые зафиксированы на всех трассах лога.

Поскольку одни и те же события возникают на различных трассах лога в различные моменты времени, то на уровне лога такие события считаются различными, что не позволяет задать отношения строгого порядка на данном уровне описания процесса решения.

Действительно, лог обладает свойствами транзитивности, антисимметричности и рефлексивности. Первые два свойства относятся к отдельным трассам процесса. Свойство рефлексивности является очевидным для одной трассы и позволяет получить качественную модель путем сопоставления событий, зафиксировавших идентичные действия в различных трассах лога:

$$\forall \pi_k \exists e'_k \in \pi_k, e'_l \in \pi_l : e'_k = e'_l \mid \pi_k \neq \pi_l, \quad (6)$$

где $\pi_k, \pi_l \in \Pi$ – различные трассы лога.

Теперь формализуем условия перехода от событийного к качественному описанию процесса на основе использования дискретного времени и к количественному описанию на основе интервального времени.

При решении этой задачи необходимо выполнить согласование порядков событий на различных трассах лога. Для этого докажем однозначность отображения упорядоченной последовательности событий, а также последовательности интервалов, отражающих выполнение процесса на отрезок натурального ряда.

Представление процесса в дискретном времени на втором качественном уровне предлагаемой модели должно объединять все возможные последовательности решения задачи, которые возникают в силу следующих факторов:

– влияние окружающей среды, представленное в форме одного из альтернативных наборов ограничений $\bigwedge \Pi_{m,l}$;

– дополнительные ограничения на процесс, задающие один из альтернативных результатов решения задачи Π_s .

Поэтому представления процесса на втором уровне задается упорядоченным множеством:

$$P_{dis} = \langle E_{dis}, R \rangle, E_{dis} = \{e_{k,i} \mid \forall (k \neq l) \ e_{k,i} \neq e_{l,i}\},$$

$$R = \{\succ, \gg, \geq\}, \quad (7)$$

где E_{dis} — набор событий процесса, который получен из множества E исключением совпадающих событий из разных трасс π_k ; $\succ$ — отношение перехода, которое связывает два последовательно возникших события; $\gg$ — транзитивное замыкание на отношении перехода; $\geq$ — отношение частичной упорядоченности по времени.

Отношение частичной упорядоченности по времени позволяет упорядочить события из разных трасс, которые не связаны между собой отношением перехода или отношением $\gg$. Более детально данное отношение будет рассмотрено далее.

Таким образом, для того, чтобы построить представление процесса (7) из лога (5) необходимо определить отношение равенства событий из разных трасс.

При определении равенства событий учитываются их дополнительные атрибуты, в частности наименование связанного с событием действия и предшествующие события, так как с одним действием может быть связано несколько событий. Для того, чтобы можно было сопоставлять несколько событий, определим понятие интервала событий.

В наиболее простом случае продолжительность текущего действия процесса определяется через разность времени возникновения пары событий, фиксирующих окончание предыдущего и текущего действия.

Тогда каждая пара вида (предшествующее событие e_i , текущее событие e_j), фиксирующая завершение предшествующего i и текущего j — действия определяют интервал ее выполнения.

Сопоставление нескольких таких числовых последовательностей для одного процесса позволяет найти узкие места и выполнить адаптацию процесса.

Таким образом, интервальное представление должно содержать набор подмножеств упорядоченных интервалов событий.

$$P_{int} = \langle \{ \langle A, \{\succ, \gg\} \rangle \}, \{ \geq \} \rangle, \quad (8)$$

где A — подмножество интервалов, для которых задан строгий порядок выполнения действий процесса, определяемый отношениями $\succ, \gg$; $\{ \langle A, \{\succ, \gg\} \rangle \}$ — множество подмножеств строго упорядоченных интервалов.

Каждое подмножество упорядоченных интервалов $\{ \langle A, \{\succ, \gg\} \rangle \}$ представляет такую последовательность действий процесса, которая

удовлетворяет соответствующему набору ограничений внешней среды.

Адаптация процесса решения задачи на основе данной модели выполняется с использованием отношения нестрогого порядка между подмножествами упорядоченных интервалов.

На практике это означает следующее. Каждый упорядоченный фрагмент процесса представляется в виде последовательности соответствующих интервалам рациональных чисел (например, меток относительного времени). Фактически эти фрагменты представляют собой подпроцессы процесса решения задачи. При изменении ограничений предметной области указанные фрагменты могут быть смещены по временной линии с тем, чтобы получить требуемый результат процесса в установленные сроки.

Иными словами, предлагаемая модель позволяет адаптировать подпроцессы. на основе их выделения и последующего комбинирования. При адаптации искомого подпроцесса остальные выступают в роли структурных ограничений.

Преобразование уровней представления процесса решения основывается на обоснованных выше правилах перехода от последовательности событий и интервалов к числовым последовательностям. Обобщенная модель процесса решения охватывает уровни упорядоченных последовательностей событий (5), дискретного (7) и интервального (8) представлений, а также правила преобразований между уровнями:

$$P = (\Pi, P_{dis}, P_{int}, O), \quad (9)$$

где O — операции преобразования элементов различных уровней.

Выводы

Предложена трехуровневая модель процесса решения задачи как элемента прецедента. Модель охватывает уровни событийного, дискретного и интервального представления процесса.

На первом уровне процесс представляется в виде набора строго упорядоченных трасс, каждая из которых отражает однократное выполнение процесса в виде последовательности событий.

На втором уровне процесс представляется в виде частично упорядоченного множества, объединяющего события из различных трасс процесса. Подмножества событий лога, фиксирующие выполнение идентичных действий при одинаковых ограничениях предметной области реализованы в модели в виде одного события.

На третьем уровне модели процесс отображается моделью с интервальным представлением времени.

В практическом аспекте предложенная модель обеспечивает возможность построения интервальной модели процесса путем дополнения интервалами выполнения действий процесса дискретного уровня представления. Последний может быть получен с помощью существующих средств process mining. Интервалы действий процесса формируются на основе анализа событийного представления.

Список литературы:

1. Richter M. M. Case-Based Reasoning. A Textbook [Text] / Michael M. Richter, Rosina O. Weber // Springer. — 2013. — 546 p.
2. Kolodner J. Case-Based Reasoning [Text] / J. Kolodner // Magazin Kaufmann. San Mateo. 1993. — 386 p.
3. Carbonell, J. G. Learning by analogy: Formulating and generalizing plans from past experience. [Text] / R.S. Michalski, J.G. Carbonell, & T.M. Mitchell (Eds.) // Machine learning, an artificial intelligence approach (Vol. 1), Palo Alto, CA: Tioga Press. — 1983. — P. 137—162.
4. Van der Aalst W.M.P. Process Mining: Discovery, Conformance and Enhancement of Business Processes [Text] / W.M.P. Van der Aalst // Springer-Verlag, Berlin — 2011. — 352 p.
5. Чалый С.Ф. Разработка обобщенной процессной модели прецедента, метода его формирования и использования [Текст] / С.Ф. Чалый, И. В. Левыкин // Журнал управляющие системы и машины. — 2016. — №3 — С. 23—29.
6. Николайчук О.А. Прототип интеллектуальной системы для исследования технического состояния механических систем [Текст] / О.А. Николайчук, А.Ю. Юрин // Искусственный интеллект. — 2006. — №4 — С. 459—468.
7. Николайчук О.А. Применение прецедентного подхода для автоматизированной идентификации технического состояния деталей механических систем [Текст] / О.А. Николайчук., А.Ю. Юрин // Автоматизация и современные технологии. — 2009. — №5 — С. 3—12.
8. Берман А.Ф. Интеллектуальная система для анализа отказов сложных технических систем [Текст] / А.Ф. Берман, О.А. Николайчук, А.И. Павлов, М.А. Грищенко, А.Ю. Юрин // Труды Тринадцатой национальной конференции по искусственному интеллекту с международным участием КИИ-2012 (16-20 сентября 2012 г., г. Белгород, Россия): Труды конференции. — М.: Физматлит. — 2012. — Т.3. — С. 146—154.
9. Aamodt A., Plaza E. Case-Based Reasoning: Foundational issues, methodological variations, and system approaches [Text] / A. Aamodt, E. Plaza // AI Communications. 1994. IOS Press, Vol.7:1, P. 39—59.
10. Van der Aalst W.M.P. Process Mining in the Large: A Tutorial. [Text] / W.M.P. Van der Aalst // In E. Zimnyi, editor, Business Intelligence (eBISS 2013), volume 172 of Lecture Notes in Business Information Processing, Springer-Verlag, Berlin, — 2014 — P. 33—76.

Поступила в редколлегию 03.11.2016

УДК 681.513



О.Ф. Михаль, О.Г. Лебедев

ХНУРЭ, г. Харьков, Украина, oleg.mikhal@gmail.com

МУЛЬТИАГЕНТНАЯ ИНТЕРПРЕТАЦИЯ ГЕНЕРАЦИИ СОБЫТИЙ В МОДЕЛИ СИСТЕМЫ МАССОВОГО ОБСЛУЖИВАНИЯ

Системы массового обслуживания (СМО) и мультиагентность могут эффективно дополнять друг друга в компьютерном моделировании. Мера соответствия модели предметной области определяется варьируемыми параметрами, что приводит к росту размерности. На примере «доски Гальтона» дан вариант модели СМО с сокращением размерности при локально-параллельном (ЛП) представлением данных. Предложены мультиагентные ЛП алгоритмы генерации потоков событий с межагентным информационным обменом.

ЛОКАЛЬНО-ПАРАЛЛЕЛЬНАЯ ОБРАБОТКА, МУЛЬТИАГЕНТНОСТЬ, СИСТЕМЫ МАССОВОГО ОБСЛУЖИВАНИЯ, СТАТИСТИЧЕСКОЕ МОДЕЛИРОВАНИЕ

Введение

При изучении (моделировании) процессов, происходящих в системах обработки информации (в частности, процессов в компьютерных системах и сетях) широко применяется парадигма *системы массового обслуживания* (СМО) [1]. Элементарным звеном СМО является *устройство*, на которое в случайные (псевдослучайные) моменты времени подаются *события* (задания), каждое из которых должно быть обработано. Для обработки (преобразования входной информации в выходную) требуется определённый конечный промежуток времени. Отдельные *устройства* скоммутированы между собой в соответствии со структурой моделируемой системы (предметной области). Выходы одних *устройств* поданы на входы других. Практический интерес представляют случаи, когда интервалы между входными *событиями* меньше времени обработки, вследствие чего образуются *очереди*, наглядно демонстрирующие *потоки событий* (ПС). Конечной целью моделирования СМО является изучение эффективности работы системы по переработке ПС, в частности, с учётом регулирования объёмов *очереди*. Для моделирования ПС применяются генераторы *событий*, обеспечивающие заданные (требуемые) статистические характеристики ПС. В компьютерном моделировании обычно используются *генераторы псевдослучайных чисел* (ГСЧ), основанные на преобразовании исходного равномерного закона распределения в требуемый, наперёд заданный закон. Ограниченность такого подхода – в его жёсткой «вычислительной детерминированности». По своей сущности, компьютеры (компьютерные системы) являются *усилителями человеческого интеллекта* [2]; поэтому реальные ПС в компьютерных системах и сетях являются *антропоморфными* («человекоподобными»), т.к. в конечном счёте они являются порождениями человеческого интеллекта и отражениями коллективного человеческого поведения.

Нет оснований предполагать, что все свойства (статистические особенности) «человеческих» ПС исчерпываются поддержанием определённых законов распределения. Поэтому представления о том, что ПС, создаваемые ГСЧ, адекватны реальным «человеческим», «человекопорождённым» ПС – является *изначально ограниченными*. В связи с этим, перспективны пути поиска *интеллектуальных* вариантов генерации событий для модели СМО, которые явились бы в большей степени адекватными «человекопорождённым» ПС. Один из перспективных путей – *мультиагентные системы* (МАС).

Цель настоящего сообщения – мультиагентная интерпретация (интерпретация с позиций МАС) процесса генерации случайного ПС, а также рассмотрение вариантов программной реализации этого процесса применительно к использованию в прикладных моделях, типа СМО. Структура предлагаемого программного решения основывается на *локально-параллельном* (ЛП) представлении информации [3]. Принцип генерации случайных ПС иллюстрируется (рассматривается) на примере классической механической модели – «доски Гальтона» [4], работа которой реализуется ЛП-алгоритмом. Реализация допускает ряд обобщений, которыми иллюстрируется мультиагентная интерпретация, включая *интеллектуальную мультиагентность*.

1. Доска Гальтона

Достаточно странно, что это устройство (демонстрационная модель) было изобретено и описано Гальтоном (Francis Galton 1822-1911) лишь в конце XIX века, хотя соответствующие вероятностные и комбинаторные представления в основном были разработаны ещё в трудах Паскаля (Blaise Pascal 1623-1662) и Ньютона (Isaac Newton 1642-1727), т.е. к XVII веку. Достаточно странно также, что в авторском описании «доски Гальтона»

(рис. 1, воспроизводит авторскую иллюстрацию из [4]) имеется *неточность* (комментируется ниже), говорящая о том, что сам автор скорее всего не реализовывал модель «в металле».

Устройство (механическая модель) представляет собой плоский лоток (ящик прямоугольной формы с низкими закраинами), в котором выделяются четыре функционально различные области: загрузочный бункер (1), воронка (2) с отверстием (щелью) шириной D , рабочее поле (3) с регулярно расположенными рядами «гвоздиков» и приёмный сепаратор (4). В бункер (1) загружаются шарики (на рис. 1 не показаны), калиброванные по диаметру d . Лоток располагается наклонно, чтобы обеспечивалось самопроизвольное скатывание шариков под действием силы тяжести в направлении от бункера (1) к сепаратору (4). Узкое (D «немного больше» d , обозначение: $D \geq d$) отверстие воронки (2) обеспечивает одиночное (последовательное, без взаимных столкновений) прохождение шариков через рабочее поле (3). «Гвоздики» расположены в N рядов с шагом $D \geq d$; при этом $(i+1)$ -й (где: $i \in \{1, 2, \dots, (N-1)\}$) ряд «гвоздиков» смещён относительно i -го ряда в горизонтальном направлении на $D/2$. Каждый шарик (очередное событие из ПС) попадает на k -й (где: $k \in \{2, 3, \dots, (m-1)\}$; m – число «гвоздиков» в i -м ряду) «гвоздик» i -го ряда и отскакивает от него; проходя с вероятностью $1/2$ влево (между $(k-1)$ -м и k -м «гвоздиками»), или вправо (между k -м и $(k+1)$ -м «гвоздиками»); попадая далее, соответственно, на k -й или $(k+1)$ -й «гвоздики» $(i+1)$ -го ряда. Процесс отскока шарика интерпретируется как случайный (в механической модели не отслеживается и специально не организуется), равновероятный и не обратимый. Поэтому каждый шарик движется в рабочем поле (3) индивидуальным случайным путём. По прохождении рабочего поля (3), шарик попадает в одну из M ячеек сепаратора (4). Вероятность попадания в j -ю (где: $j \in \{1, 2, \dots, M\}$) ячейку – пропорциональна числу возможных исходов статистических экспериментов. Эти числа, как известно [5], есть соответствующие строки в «треугольнике Паскаля»; они же – коэффициенты в разложении бинома Ньютона соответствующей степени.

Назначение «доски Гальтона» демонстрация прохождения *некоторого* случайного процесса. В действительности в макромире (с объектами макромира) процессы не столь (не в такой степени) случайны. Вспомним, что модель была разработана в «до-квантово-механический» период науки, когда детерминистская Вселенная «работала исключительно по Ньютону и Лапласу». Поэтому *главная* идея «доски Гальтона» не в конкретной (чисто игровая) ситуации с «гвоздиками»

и шариками. Модель демонстрирует тот факт, что, при наличии большого числа независимых внешних влияющих факторов, полный учёт обстоятельств поведения объекта чрезвычайно затруднителен. Содержательный результат реально может быть получен лишь в статистическом смысле. Влияющими факторами в модели являются ряды «гвоздиков» рабочего поля (3), а статистический результат – картина распределения шариков в ячейках сепаратора (4).

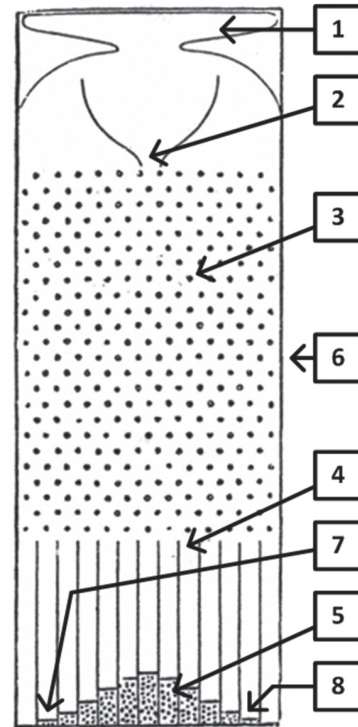


Рис. 1. Эскиз модели «доски Гальтона», заимствованный из [4], дополненный нумерационными обозначениями (1)–(8) для комментирования

Неточность, допущенная Гальтоном в изображении модели рис. 1, состоит в том, что распределение (5) шариков в сепараторе может иметь правильную «колоколообразную» форму только в случае если сепаратор (4), имеющий M ячеек, помещён после M -й строки «гвоздиков». На рис. 1 этот уровень обозначен (6). Таким образом, на оригинальном рисунке в [4] рабочее поле (3) в два раза длиннее, чем нужно. Это значит, что редкие (низковоероятные) шарики, направляющиеся в «хвосты» статистического распределения, будут отскакивать от боковых стенок и перенаправляться к центру распределения. Результат – «хвосты» (7) и (8) распределения в реальной физической модели, выполненной по эскизу из [4], будут «излишне утолщены» и реальная картина будет отличаться от правильной колоколообразной (5).

С учётом представленного уточнения, описанная Гальтоном модель является «стандартным вариантом». Физические реализации модели

(например [6]) позволяют выявить ряд факторов, помимо стенок лотка, существенно влияющих на картину распределения шариков в сепараторе (5). Так, существенно важны конфигурация «гвоздиков» и соотношение между D и d . В [6] демонстрируются «гвоздики» ромбоидального сечения (рис. 2) и три соотношения: $D \geq d$, $D \sim 2d$ (т.е. D «порядка» $2d$) и $D \sim 3d$. Первое даёт более-менее правильную колоколообразную форму распределения; второе — существенно стягивает и укрупняет центральную часть распределения (получается нечто похожее на шляпу «котелок»); третье — напротив, растягивает распределение почти что в равномерное.

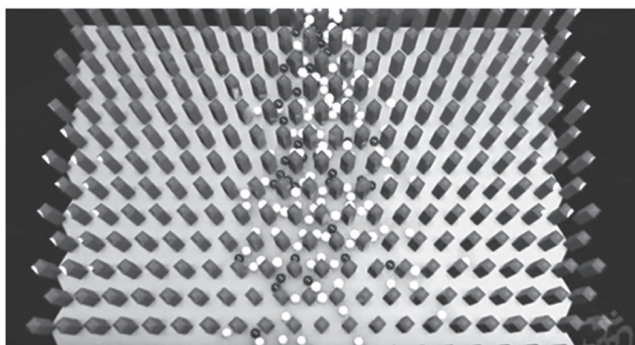


Рис. 2. Картина движения шариков между «гвоздиками» ромбоидального сечения

Ретроспективно, может быть предложено следующее объяснение. Имеется два дополнительных влияющих фактора. *Первый фактор*: наклонные плоскости верхних частей ромбоидальных «гвоздиков» придают шарикам момент вращения. При этом, с уменьшением диаметра, шарик успевает сделать большее число оборотов, катясь по наклонной плоскости «гвоздика», т.е. получить больший момент вращения. Наличие момента вращения влияет на выбор направления отскока шарика: способствует сохранению направления, полученного при столкновении с первым «гвоздиком». *Второй фактор*: с уменьшением диаметра шарика всё в большей степени нарушается «принцип отсутствия взаимовлияния между шариками». Шарики начинают «толкаться», двигаться группами и «мешать друг другу» сохранять момент вращения. При $D \sim 2d$ по-видимому преобладает *второй фактор*. Сохранение момента вращения сказывается не столь значительно, поэтому только «проседает» вершина колокола распределения. При $D \sim 3d$ по-видимому начинает преобладать *первый фактор*. Существенная часть шариков набирает значительные моменты вращения, так что «толчки» движущихся с ними «соседей» в меньшей степени дезориентируют их вращение. Результат: существенная часть шариков попадает в «хвосты» распределения, а «центральная вершина» оказывается «совершенно просевшей».

Разумеется, возможные факторы влияния не исчерпываются описанным. Просто, данная картина была продемонстрирована в модели [6]. Вероятно, варьируя конфигурацию расположения «гвоздиков», размер и вес шариков, шероховатость их поверхностей и др., можно получить ещё более разнообразные картины распределений. Вероятно, дополнительные возможности разнообразия обеспечит введение «дальнодействия» (дистанционного притяжения-отталкивания) между шариками. Разумеется, изобретение способов усложнения модельной ситуации не является предметом настоящего сообщения. Приведенное ретроспективное объяснение представлено единственно с целью демонстрации широких концептуальных возможностей такой простой модели как «доска Гальтона» для генерации событий в более сложных (составных) моделях типа СМО.

Далее рассмотрим вопросы программной реализации модели типа «доски Гальтона» для использования её в качестве генератора ПС, с мультиагентной интерпретацией.

2. Локальная параллельность

Числа в компьютере хранятся и обрабатываются в двоичном представлении. При реализации математических операций, числа (операнды) находятся в соответствующих регистрах процессора. Двоичная разрядность N регистров процессора — фиксированная.

(Здесь оставлена переменная N , применявшаяся выше, в п.1. Из дальнейшего будет понятно, что она имеет сходный смысл.)

В настоящее время в компьютерах общего назначения распространены, преимущественно, 32- и 64-разрядные процессоры. Если обрабатываемые числа не слишком велики, они занимают не весь регистр, а, возможно, только несколько младших двоичных разрядов. Старшие разряды заполнены нулями, и при проведении вычислительных операций — процессор «гоняет эти нули впустую».

В ЛП представлении регистр с двоичной разрядностью N подразделяется на m не пересекающихся соседствующих k -разрядных (двоичные разряды) сегментов, где $mk \leq N$. Будем называть такую структуру *регистровым представлением* (РгП). В каждый из сегментов РгП помещается обрабатываемое число. Разумеется, оно ограничено значением 2^k . Бинарные операции над РгП производятся по специально разработанным ЛП-алгоритмам [3], так, что i -й ($i \in 1, 3, \dots, m$) сегмент одного операнда взаимодействует только с i -м сегментом другого операнда. Побочные межсегментные взаимодействия (перенос разряда) — алгоритмически устраняются. В результате, при каждой

бинарной операции над РгП, реализуется параллельно (одновременно) m бинарных операций над обрабатываемыми данными.

Для демонстрации сказанного, рассмотрим следующий упрощённый пример. Будем считать, что разряды РгП — десятичные. Пусть имеются следующие «примеры на сложение»: $2+3$, $5+4$, $6+2$, $8+3$ и $7+6$. Требуется выполнить эти вычисления.

В классическом последовательном варианте «примеры» должны решаться один за другим. Числа, которые нужно складывать, изначально помещены в некоторых ячейках памяти. Результаты сложения так же должны быть помещены в некоторые ячейки памяти. Процессор будет использовать три регистра: A (для первого операнда), B (для второго операнда) и C (для результата). Для каждого из «примеров» процессор выполняет следующие 4 простые операции (шага):

- извлечь из памяти значение первого операнда, поместить его в регистр A ;
- извлечь второй операнд и поместить в B ;
- произвести операцию «+», при этом регистры A и B освобождаются (очищаются), а результат помещается в регистр C ;
- извлечь результат из регистра C (при этом он очищается) и поместить в соответствующую ячейку памяти.

Таким образом, для решения всех 5 «примеров» требуется произвести 20 указанных простых операций.

В ЛП варианте операнды вначале конкатенируются. Получаются числа $A=0205060807$ и $B=0304020306$. Затем производится операция суммирования. В нашем простом случае это обычная операция сложения, поскольку введены разделительные нули для помещения переноса разряда. Результат $C=0509081113$ интерпретируется как конкатенация результатов решения отдельных «примеров». Деконкатенация (разделение сегментов) даёт набор результатов (5, 9, 8, 11, 13), соответствующий искомым значениям. Легко видеть, что, не считая конкатенации и деконкатенации, результаты получены за 4 шага, что соответствует параллельному решению всех 5 «заданных примеров».

Разумеется, конкатенация и деконкатенация, сами по себе, так же требуют расхода определённого временного ресурса. В результате, данный представленный ЛП вариант (данный конкретный пример), возможно, окажется даже проигрышным по времени по сравнению с последовательным вариантом. Но если задача представляет собой значительный комплекс вычислительных операций и конкатенация (подготовка данных) осуществляется только в самом начале, а деконкатенация (интерпретация результатов) только в самом конце, то

выигрыш в ЛП производительности может оказаться существенным, близким к значению, пропорциональному числу сконкатенированных сегментов.

Далее рассмотрим процедуру моделирования СМО. Для представления отдельных событий (элементов) используются отдельные биты. Т.о., размер сегмента — $k = 1$ — один бит.

3. Моделирование

При генерировании в 32-разрядном варианте, в качестве исходного, берётся число:

$$A = 32768_{(10)} = 1000000000000000_{(2)}. \quad (1)$$

Подстрочный индекс — основание системы счисления. Единица изображает собой шарик. Положение (позиция) единицы — номер двоичного разряда, в котором она находится, — есть текущее положение шарика при движении в рабочем поле. Легко видеть, что в A единица стоит по центру 32-разрядного числа, в 16-м разряде. Далее следует «прохождение рядов гвоздиков», воспроизводимое регистровыми сдвигами ($A < 1$) (влево, в сторону старших разрядов) и ($A > 1$) (вправо, в сторону младших разрядов), управляемыми от ГСЧ. В результате, «шарик выкатывается» к одной из «ячеек сепаратора», где накапливается распределение.

Могут быть предложены следующие 3 варианта реализации «конструкции сепаратора», т.е., процесса построения (накопления) результирующего модельного распределения.

1. Применение оператора *case*. Число A является степенью двойки. Оператор *case* должен содержать все возможные варианты, и для каждого из них выдавать соответствующее значение B . Указанный B -й счётчик в «сепараторе» увеличивается на 1.

2. Подсчёт (учёт) числа регистровых сдвигов. Для этого в 32-разрядном варианте, в начале устанавливается $B = 16$ (исходная позиция шарика). Далее в цикле по числу N рядов гвоздиков реализуются $B += 1$, (при $A > 1$) или $B -= 1$, (при $A < 1$), т.е. значение B изменяется согласно текущей позиции шарика. При завершении цикла, B -й счётчик в «сепараторе» увеличивается на 1.

3. Полностью ЛП вариант с «вертикальным размещением информации». Вводится массив из k чисел $C(j)$, $j \in (1, 2, \dots, k)$, являющихся хранителями j -х двоичных разрядов чисел, накапливаемых в «распределителе». Организуется ЛП алгоритм модифицированного поразрядного накопления (МПН) единиц из сгенерированных чисел A , при котором i -е разряды не взаимодействуют ни с $(i+1)$ -ми, ни с $(i-1)$ -ми. Сущность МПН в том, что по каждому из i -х разрядов происходит (по мере поступления новых единичек) «в вертикальном направлении» (по индексу j) двоичный пересчёт, как в электронном счётчике (триггерной цепочке). Процедура

завершается ЛП транспонированием матрицы, после которого имеется m штук k -разрядных чисел — сгенерированное моделью результирующее распределение.

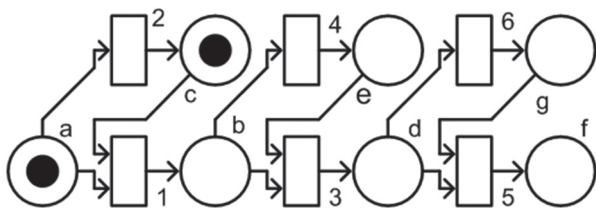


Рис. 3. Сеть Петри, иллюстрирующая алгоритм МПН

Принцип действия МПН иллюстрируется приоритетной сетью Петри (СП) (рис. 3). СП воспроизводит работу одного сегмента ЛП представления данных, что соответствует работе одной ячейки сепаратора. Места обозначены $a, b, \dots, g$; переходы пронумерованы в порядке их опроса и срабатывания.

Место a — шарик, поступающий на вход i -го «сепаратора». Размер (глубина) сегмента $k = 3$, что соответствует трём строкам мест (b, c), (d, e) и (f, g). В изображённом случае метки из мест a и c «зарядят» переход 1. После его срабатывания метка поступит в b . Далее сработает переход 4 и метка поступит в e . Таким образом, после поступления метки из a , запись в i -м сегменте изменится с $1_{(2)}$ на $10_{(2)} = 2_{(10)}$. При поступлении следующей метки в a , место c будет пустым, поэтому переход 1 не сработает, но сработает переход 2. Метка поступит в c , что будет соответствовать записи в i -м сегменте $11_{(2)} = 3_{(10)}$. Таким образом, всякий раз после поступления на вход a очередного шарика, вся система срабатывает, как триггерная пересчётная цепочка. В результате, расстановка меток в линейке (c, e, g) представляет собой число в двоичной системе счисления (метка изображает 1; c — младший разряд), равное количеству шариков, поступивших в данную ячейку.

Полная процедура завершается ЛП транспонированием матрицы [7], после которого получается m штук k -разрядных чисел — сгенерированное моделью результирующее распределение. Рассмотренный вариант интересен тем, что в нём обработка всецело, включая транспонирование матрицы [7], выполнена в ЛП представлении информации. При этом, содержательная ЛП часть работы «сепаратора» выдаёт результат сразу, минуя этап деконкатенации. Вместо деконкатенации, работает ЛП транспонирование матрицы.

4. Антропоморфные представления

Как отмечалось, в моделировании для создания входных (исходных, изначальных) информационных потоков традиционно используются ГСЧ.

Стандартный ГСЧ обеспечивает равномерный закон распределения. Далее равномерный закон преобразуется в требуемый, наперёд заданный. Ограниченность этого подхода в его «вычислительной детерминированности». Создаваемый информационный поток может иметь любые требуемые статистические характеристики, но не соответствовать моделируемой *антропоморфной реальности*. Принципиальное различие состоит в следующем.

Возможно, не корректно будет говорить, что человек полностью создаёт (контролирует) окружающую его реальность. Но та часть реальности, в которой человек живёт, является понятной ему, «уже отмоделированной» и потому относительно безопасной. В потенциально опасную (неизвестную, непонятную) часть реальности человек «выходит» постепенно, моделируя её, приспособливая её к своим представлениям и одновременно расширяя круг своих представлений. Расширенный круг представлений может содержать ошибки (человеческие заблуждения). Их выявление — естественный элемент процесса познания реальности. Но расширение круга *антропоморфных представлений* происходит также не бесконтрольно, а в рамках *антропоморфных представлений* о том, как должен расширяться круг *антропоморфных представлений*. В результате, та реальность, в которой живёт человек, сама по себе, уже является *антропоморфной*. Существенная её часть есть пересечение с *реальной реальностью*. Но есть ещё *непознанная реальная реальность* (реальная, но лежащая вне пересечения с *антропоморфной реальностью*), а также область человеческих фантазий и заблуждений. Парадоксально, но последние — в не меньшей степени определяют человеческую (разумную, мотивированную) деятельность. Здесь следует вспомнить религию.

Таким образом, хотя критерием истинности и является практика, но «человек является мерилем всех вещей» (Протагор), поскольку человек применяет (реализует) практику, как критерий истинности, в меру наличествующей созданной им *антропоморфной реальности*. Вещи, не подходящие под этот критерий, находятся вне человеческой реальности. Как результат, человеческие действия могут выглядеть как случайные, но в действительности они содержат (могут содержать) некую осознанную структуру. Пример такой структуры — отсутствие, либо наоборот — наличие, повторений.

Пример с отсутствием повторений. Требуется написать последовательность (строку) со случайным расположением цифр от 1 до 9. В машинном варианте (в рамках стандартных подходов на основе ГСЧ) в этой последовательности будут достаточно часто встречаться рядом стоящие одинаковые

цифры; в человеческом варианте — нет. Человек будет стараться избегать повторов, интуитивно полагая, что таким образом случайная последовательность будет «более случайной».

Пример с наличием повторений. Требуется ввести пароль и войти в некий ресурс. Неприятность в том, что пароль «полузабыт». В человеческом варианте возможно «вспоминание полузабытого», с повторным многократным вводом одних и тех же комбинаций. Ход мысли такой: «...Да что такое? Неправильно ввёл, что ли?...» В машинном варианте (взлом пароля) — полный последовательный или выборочный перебор без повторов, с полной фиксацией опробованных вариантов.

Попутно отметим, что данные примеры годятся для реализации алгоритма *теста Тьюринга*. Предположительно, они *сработают* для интеллектуальных программ «общего назначения»; и *провалятся* для интеллектуальных программ, в которых будет *специально имитироваться* описанная особенность человеческого восприятия реальности. Т.е. такой тест окажется не эффективным для специально разработанных программ — «имитаторов человеческих заблуждений». Возможны ли (могут ли потребоваться) такие интеллектуальные программы? Работа (не алгоритм) и гипотетические последствия работы одной из них достаточно реалистично описаны в [8].

Статистическое описание «скрытой структуры» моделируемой реальности, само по себе, является *одним из* антропоморфных представлений. Т.е. в нём изначально заложена его неполнота и возможная ошибочность применительно к реальности. Так, в реальной жизни (в своей «антропоморфной реальности») человек редко руководствуется исключительно постановкой статистических экспериментов и статобработкой. *Утрированный пример:* при покупке хлеба в магазине человек не станет набирать полную статистически достоверную выборку результатов измерения (50 раз переспрашивать продавца о наличии хлеба).

Резюмируем сказанное. По своей сущности, компьютеры (компьютерные системы) являются *усилителями человеческого интеллекта* [2]; поэтому реальные ПС в компьютерных системах и сетях являются *антропоморфными*, т.к. в конечном счёте, они являются порождениями человеческого интеллекта и (или) отражениями коллективного человеческого поведения. Поэтому нет оснований полагать, что все свойства (статистические особенности) «человеческих» ПС исчерпываются поддержанием определённых законов распределения. Поэтому представления о том, что ПС, создаваемые ГСЧ, адекватны реальному «человеческим», «человекопорождённым» ПС — является

изначально ограниченными. В связи с этим, перспективны пути поиска «интеллектуальных» вариантов генерации событий для моделей СМО, которые явились бы в большей степени адекватными «человекопорождённым» ПС. Один из перспективных путей — *мультиагентные системы* (МАС).

5. Мультиагентность

Интеллектуальный агент (ИА) есть некоторая программная (алгоритмическая) сущность. МАС — образование из нескольких взаимодействующих ИА, предназначенное для решения таких проблем, которые сложно или невозможно решить с помощью одного ИА. Характерные особенности МАС — автономность ИА; ограниченность представлений каждого отдельного ИА (ни у одного из ИА нет представления о *всей* системе; система слишком сложна, чтобы знание о ней имело практическое применение для данного конкретного ИА); децентрализация (нет ИА, управляющих всей системой); самоорганизация и сложное поведение системы (при том, что стратегия поведения каждого ИА достаточно проста); роевой интеллект (ищется оптимальное решение задачи, на которое потрачено наименьшее количество энергии в условиях ограниченных ресурсов). Для реализации указанных свойств (возможностей), ИА могут обмениваться полученными знаниями, используя некоторый специальный язык и подчиняясь установленным правилам «межагентского общения» (протоколам). При этом, ключевым преимуществом МАС становится гибкость: МАС может быть дополнена (модифицирована) без переписывания значительной части программного кода.

Интересно, что в рамках МАС реализуются *антропоморфные представления*, поскольку человеческое общество в целом соответствует перечисленным выше характеристикам МАС. В связи с этим, парадигма МАС представляется продуктивной применительно к моделированию процессов и ситуаций *антропоморфной реальности*.

Рассмотрим вариант МАС, демонстрирующий аналог *антропоморфного поведения* модели. В данном варианте мультиагентность реализуется только на уровне генерации ПС. Т.е., рассматриваются только ИА — генераторы ПС. На рис. 4 показаны *моделируемый объект* (МО) и N интеллектуальных агентов (ИА-1, ИА-2,..., ИА-N), выделенных пунктирными рамками. Для сохранения наглядности, только два ИА раздетализованы, остальные (аналогичные) — свёрнуты. Для обозначения блоков ИА применены следующие аббревиатуры: *анализатор состояния модели* (АСМ), *анализатор других потоков событий* (АДПС), *генератор потока событий* (ГПС), *корректировщик генератора* (КГ).

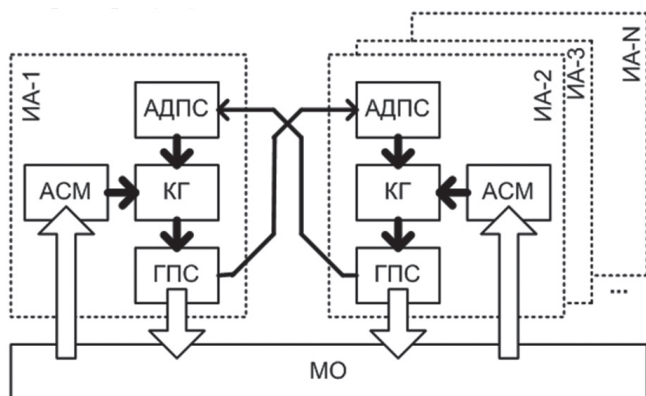


Рис. 4. Модель с мультиагентной организацией генерации потоков событий

МО – единый для всей МАС. Отдельные ИА – автономны в принятии решений, и дополнительно общаются между собой по вопросам согласования отдельных параметров. Деятельность ИА состоит в генерации ПС для МО. Сами ПС показаны белыми стрелками ГПС → МО. Кроме того, каждый из ИА независимо контролирует МО. Соответствующие информационные потоки показаны белыми стрелками МО → АСМ. Тёмными жирными стрелками на рис. 4 показаны межблочные внутриагентные связи, а тёмными тонкими стрелками – межагентные связи по принципу «каждый с каждым». Для сохранения наглядности, на рис. 4 изображены только связи, касающиеся ИА-1 и ИА-2; остальные – условно опущены.

Система работает следующим образом. ГПС генерируют ПС для работы МО. Генерируемые события имеют определённые параметры (в частности это могут быть статистические характеристики, дополненные «для интеллектуальности» некоторыми решающими правилами). Конкретные значения параметров устанавливаются КГ. Состояние (поведение) МО отслеживается АСМ. Имеется в виду, что следствием сгенерированного ПС данного ИА и ПС остальных ИА является изменение некоторых характеристик МО. Эти характеристики отслеживаются (измеряются, регистрируются) АСМ данного ИА. Возможно (допустимо), что каждый из ИА отслеживает свой набор характеристик МО. Возможно (допустимо) так же, что эти наборы пересекаются. Уникальность наборов характеристик или наличие пересечений – в данной (рис. 4) модели не принципиальны. Важно, что в АСМ имеются критерии сравнения (оценки), на основе которых вырабатывается и передаётся (АСМ → КГ) некоторая *рекомендательная информация* для изменения последующей генерации ПС.

ИА общаются между собой. Это может быть интерпретировано как «язык общения» или «протоколы информационного обмена». Сущность в том, что каждый из ИА «интересуется вопросом:

а как ведут себя другие ИС?». (Отметим, что подобная антропоморфная трактовка межагентного информационного обмена – по существу не более антропоморфна, чем использование терминов «язык» или «протокол»). Для этого в рассматриваемом варианте МАС, у ИА имеются АДПС, куда поступает информация с ГПС других (остальных (N-1)) ИА. (Как отмечено выше, на рис. 4 показаны связи ГПС → АДПС только для одной пары ИА) Характер этой информации – не смотря на аббревиатуру АДПС – не обязательно полные ПС с остальных ГПС. С точки зрения разработки конкретной прикладной модели, более экономно, напротив, чтобы это были только значения некоторых *ключевых параметров* ГПС (что и составляет, собственно, «язык» или «протоколы обмена»). Тогда в ИС не нужны будут собственные (локальные) системы анализа ПС, аналогичные той, что реализуется в МО.

Полагается, что в АДПС, так же как и в АСМ, имеются свои «пороговые уровни» или «критерии отбора», согласно которым вырабатывается и подаётся (АДПС → КГ) *рекомендательная информация*, обобщающая «опыт работы» остальных ИА применительно к улучшению работы данного ИА. Далее, в КГ информация с АСМ и АДПС обобщается и на её основе вводятся коррективы в работу ГПС.

6. Варианты реализации

Одна из особенностей МАС – *малоресурсность*. Применительно к моделированию, ИА в МАС «берут не умением, а числом». Как отмечено выше, ИА автономны, «не слишком разумны» (ни один из них не охватывает «замысла системы» в целом) и децентрализованы (ни один из ИА не управляет системой в целом). Смысл указанных ограничений – простота, низкая затратность в программировании и активизация самоорганизации по типу роевого интеллекта. В аспекте моделирования, парадигма МАС хорошо соответствует технологии (математическому аппарату) СМО, что рассмотрено выше (п. 2) на примере «доски Гальтона»; а в аспекте реализации – ЛП-подходу, описанному выше в целом (п. 3) и применительно к реализации «доски Гальтона» (п. 4). Рассмотрим далее вопросы ЛП-реализации структуры МАС рис. 4

В качестве базового варианта, избрана упрощённая структура СМО, по типу «доски Гальтона». Рассмотрен ЛП вариант движения шарика, как последовательности регистровых сдвигов начального числа (1). Направление сдвига задаётся ГСЧ с равномерным законом распределения. Отмечено, что закон распределения можно менять, но *антропоморфными* случайные числа от этого не становятся. Приведён пример одной из возможных особенностей

антропоморфности: наличие или отсутствие (в зависимости от контекста) соседствующих повторений. В развитие этих идей, предложим, *как пример*, следующее.

1. Не сложно сгенерировать с помощью ГСЧ w_i — последовательность нулей и единиц, определяющую «судьбу i -го шарика» — последовательность его отскоков влево или вправо при движении по «доске Гальтона». Не сложно так же составить ЛП алгоритм, который будет «реализовывать эту судьбу» — брать начальное число (1) и выполнять последовательность регистровых сдвигов согласно «предписанию» w_i .

2. Числа w_i могут быть изначально сгенерированы в требуемом количестве. Тогда могут быть определены дополнительные логические операции, которыми массив чисел w_i может быть откорректирован так, чтобы соответствовать любым наперед заданным (формализованным) представлениям об антропоморфности. В частности, для исключения соседствующих повторяющихся чисел, массив должен быть проверен на соответствие правилу $w_i = w_{i+1}$.

3. Некоторые обобщённые свойства (в зависимости от конкретного моделируемого контекста) массивов чисел w_i для разных ИА, могут быть использованы как *ключевые параметры*, передаваемые «по межагентской связи» ГПС → АДПС. Объёмы массивов чисел w_i соответствуют объёмам ПС, но их обобщения — *ключевые параметры* — существенно меньше. При надлежащем уплотнении и формализации (в зависимости от контекста моделируемой системы), они могут быть сведены до коротких записей протоколов межагентного обмена в виде чисел, интерпретируемых как РгП. Далее они могут обрабатываться ЛП алгоритмами.

4. Выявление взаимного соответствия указанных *ключевых параметров*, принадлежащих разным ИА, сводится так же к ЛП алгоритмам, конкретный вид (содержание) которых определяется контекстом моделируемой системы.

Реализация представленной ЛП концепции основывается на изначальной ЛП реализуемости СМО и на принципиальной возможности ЛП описания информации протоколов межагентского обмена. Достигаемые при этом преимущества — низкоресурсность и выигрыш во времени обработки — всецело соответствуют идеологии МСО.

Выводы

В статистическом моделировании широкое распространение получили *системы массового обслуживания* (СМО) и концепция *мультиагентности*. Эти две парадигмы эффективно дополняют друг друга в поддержании разнообразия модельных сущностей. Но качество модели (мера её соответствия реалиям предметной области) определяется множественностью параметров и возможностями их варьирования, что приводит к росту размерности модели. На примере «доски Гальтона» показан простой вариант модели СМО, демонстрирующий, что эффективным методом преодоления роста размерности системы является *локально-параллельное* (ЛП) представление данных. В плане практической реализации, предложены ЛП алгоритмы в мультиагентной интерпретации по генерации потоков событий, а так же рассмотрены ЛП варианты поддержки самоорганизующихся систем с межагентным информационным обменом.

Список литературы:

1. Михаль О.Ф., Лебедев О.Г. Мультиагентная интерпретация модельного представления системы массового обслуживания // Современные направления развития информационно-коммуникационных технологий и средств управления: Материалы шестой международной научно-технической конференции. — Полтава: ПНТУ; Баку: ВА ЗСАР; Кировоград: КЛА НАУ; Харьков: ДП «ХНИИ ТМ», 2016 — С. 29.
2. Михаль О.Ф. Глобально-исторический контекст развития средств вычислительной техники // Бионика интеллекта : науч.-техн. журн. — Х. : Изд-во ХНУРЕ, 2014. — Вып. 1 (82). — С. 55–62.
3. Михаль О.Ф. Локально-параллельные однородные алгоритмы нечетких операций, основанных на процедурах сложения, вычитания и сравнения // Вестник Харьковского государственного политехнического университета. 2000. Вып. 118. С. 6-9.
4. Википедия. Доска Гальтона https://ru.wikipedia.org/wiki/Доска_Гальтона
5. Выгодский М.Я. Справочник по элементарной математике. Гос. издат. ф.-м. лит., М.- 1959 — С. 247.
6. Galtonboard / Galtonbrett Simulation (or Bean machine or quincunx or Galton box) <https://www.youtube.com/watch?v=3m4bxse2JEQ>
7. Михаль О.Ф., Хрипкова М.Ю. Локально-параллельный алгоритм транспонирования матрицы нечеткого отношения в однородном регистровом представлении // Вестник Харьковского государственного политехнического университета. 2000. Вып. 119. С. 9-12.
8. Кларк А. Космическая одиссея 2001 года. М.: Мир, 1991 г. — 616 с.

Поступила в редакцию 15.11.2016

УДК 519.6



О.М. Литвин, О.В. Ярмош, Т.І. Чорна

Українська інженерно-педагогічна академія, Харків, Україна

tanya_chorna@ukr.net

МЕТОД СПЛАЙН-ІНТЕРЛІНАЦІЇ ПРИ ЗНАХОДЖЕННІ НАЙБІЛЬШИХ (НАЙМЕНШИХ) ЗНАЧЕНЬ ФУНКЦІЇ ДВОХ ЗМІННИХ В ЗАМКНУТІЙ ОБЛАСТІ

Пропонується для розв'язання задачі знаходження найбільшого та найменшого значень неперервної функції двох змінних в замкнутій області $D = [0, 1]^2$ використовувати послідовність оптимізаційних задач на заданій системі взаємоперпендикулярних прямих. Побудовано оператори сплайн-інтерлінації та охарактеризовано їх властивості. Розглянуто приклади. Наведено аналіз обчислюваного експерименту.

НЕПЕРЕРВНА ФУНКЦІЯ, ІНТЕРВАЛИ МОНОТОННОСТІ, ІНТЕРПОЛЯЦІЙНИЙ СПЛАЙН, СИСТЕМА ВЕРТИКАЛЬНИХ І ГОРИЗОНТАЛЬНИХ ПРЯМИХ

Вступ

Приймаючи ті чи інші рішення у різних проблемних ситуаціях наукової або практичної діяльності виникають найрізноманітніші завдання оптимізації. Розв'язуючи ці завдання у більшості випадків намагаються знайти найбільше або найменше значення досліджуваної функції при заданих обмеженнях, оскільки технічні чи економічні процеси моделюються функцією або декількома функціями, які залежать від змінних факторів, що впливають на стан моделюючого процесу.

Класичне формулювання алгоритму знаходження найбільшого (найменшого) значень функції двох змінних в замкнутій області D полягає в наступному.

Крок 1. Знаходимо стаціонарні точки першого роду (X_k, Y_k) , $k = \overline{1, M}$, $M \geq 0$ в середині області D .

В даних точках $\frac{\partial f(x, y)}{\partial x} = 0$ і $\frac{\partial f(x, y)}{\partial y} = 0$ або можуть не існувати.

Крок 2. Знаходимо стаціонарні точки на границі Γ області D . Позначаємо їх $(X_p^*, Y_p^*) \in \Gamma$, $p = \overline{1, N}$, $N \geq 0$.

Крок 3. Обчислюємо значення $Z_k = f(X_k, Y_k)$, $k = \overline{1, M}$, $(X_k, Y_k) \in \overline{D} \setminus \Gamma$ досліджуваної функції в точках (X_k, Y_k) (якщо $M = 0$, то такі значення відсутні) та значення $Z_p^* = f(X_p^*, Y_p^*)$ в точках границі $(X_p^*, Y_p^*) \in \Gamma$, $p = \overline{1, N}$ (якщо $N = 0$, то таких значень немає). Обов'язково підраховуємо значення функції і в кутових точках, якщо вони існують.

Крок 4. Визначивши

$Z_k = f(X_k, Y_k)$, $k = \overline{1, M}$ та $Z_p^* = f(X_p^*, Y_p^*)$, $p = \overline{1, N}$, знаходимо

$$\max \{Z_1, \dots, Z_M, Z_1^*, \dots, Z_N^*\} \text{ або} \\ \min \{Z_1, \dots, Z_M, Z_1^*, \dots, Z_N^*\}.$$

Якщо $M = 0$, то відповідних елементів у фігурних дужках не буде.

Якщо в області D число критичних точок, в яких може бути найбільше чи найменше значення, велике $(M + N) > 10^3$, то знаходження їх координат вимагає великої кількості обчислень (взагалі кажучи слабо формалізованих операцій — диференціювання, знаходження точок недиференційовності тощо). Тому ми пропонуємо метод редукції поставленої задачі на найбільше та найменше значення функції двох змінних в замкнутій області до послідовності задач на знаходження найбільшого і найменшого значення $f(x, y)$ на системі прямих $x = X_p = \frac{p}{m}$, $0 \leq p \leq m$, $m \geq 2$ та $y = Y_q = \frac{q}{n}$, $0 \leq q \leq n$, $n \geq 2$.

Тобто в даній роботі пропонується метод знаходження найбільшого та найменшого значення неперервної функції $f(x, y)$ шляхом зведення двовимірної задачі до послідовності одновимірних задач на знаходження найбільшого та найменшого значення на системі паралельних горизонтальних та вертикальних прямих, які перетинають область дослідження D . Для тестування використовується $D = [0, 1]^2$.

1. Актуальність використання методу сплайн-інтерлінації

Задача полягає в наступному: на квадраті $D = [0, 1]^2$ знайти найбільше або найменше значення функції $f(x, y) \in C(D)$ за допомогою методу сплайн-інтерлінації функції двох змінних.

Серед численних методів знаходження глобального екстремуму увагу привертають наступні.

Класичні методи знаходження глобального екстремуму, які базуються на диференціальному численні. Вони передбачають визначення всіх точок можливих екстремумів на основі рівності частинних похідних функції в даних точках нулю, обчислення значення функції в цих точках, знаходження

серед них максимальних або мінімальних значень [1, 2, 3].

Але, як справедливо зазначав Н.С. Бахвалов, формальний підхід вирішення задач оптимізації потребує дуже велике число операцій та знаходиться під сильним впливом округлень при обчисленнях [4, с. 323].

Окрім того, класичний метод дає можливість знайти екстремум, якщо він знаходиться усередині, а не на границі можливих значень аргументів.

Тобто, класичний метод, який вивчається на першому курсі університетів та вищих технічних закладів, має деякі недоліки. Вони були усунені у відомих методах знаходження найбільшого і найменшого значення в методі лінійного програмування для того частинного випадку, коли область D є випуклим багатокутником, а функція $f(x, y)$ є лінійною функцією двох змінних [2, 3].

Відмітимо також метод квадратичного програмування, в якому функція, що мінімізується, може бути квадратичною функцією, і обмеження, що виділяють область, в яких знаходиться екстремум можуть бути записані як у вигляді лінійних нерівностей, так і у вигляді квадратичних нерівностей [5, 6]. Задачу знаходження глобального екстремуму квадратичних функцій багатьох змінних детально проаналізовано в [7, Див. також огляд робіт в 7].

В роботах [2] досліджується для розв'язання задачі оптимізації випуклих функцій випадок недиференційовних функцій і вводиться поняття субградієнта для розв'язання задачі оптимізації випуклих функцій.

Чисельні методи знаходження найбільшого та найменшого значення в замкнутій області обґрунтовуються різними науковцями [8, 9, 10, 11, 12].

Серед цих методів і метод градієнтного спуску, метод яру («овражный метод»), симплекс — метод, метод статистичної обробки результатів випадкових випробувань (метод Монте-Карло) та багато інших.

Та на даний час авторам невідомі загальні конструктивні методи мінімізації функції двох змінних на довільній замкнутій однозв'язній області.

Тому, очевидно, що використання методу сплайн-інтерлінації для знаходження найбільшого та найменшого значення функції в замкнутій області є доцільним і цілком актуальним.

2. Сутність методу сплайн-інтерлінації при знаходженні найбільшого (найменшого) значення функції двох змінних

Будемо вважати, що функція $f(x, y)$ є неперервною і має для кожного x скінченну кількість інтервалів монотонності для $y \in [0, 1]$, а також для кожного y функція $f(x, y)$, як функція тільки

змінної y , теж має скінченну кількість інтервалів зростання та спадання (інтервалів монотонності) при $x \in [0, 1]$.

Розіб'ємо область D на прямокутні елементи $\Pi_{p,q} = \{X_{p-1} \leq x < X_p, Y_{q-1} \leq y < Y_q\}$, $p = \overline{1, m}$; $q = \overline{1, n}$ прямими $x = X_p, 0 = X_0 < X_1 < \dots < X_m = 1; y = Y_q, 0 = Y_0 < Y_1 < \dots < Y_n = 1, p = \overline{0, m}, q = \overline{0, n}, m \geq 2, n \geq 2$.

Побудуємо оператор сплайн-інтерлінації у вигляді $O_{m,n}f(x, y) = (O_{1,m} + O_{2,n} - O_{1,m} \times O_{2,n})f(x, y)$.

$$O_{2,n}f(x, y) = \sum_{q=0}^n f(x, Y_q) \times h_{2,q}(y)$$

$$O_{1,m}f(x, y) = \sum_{p=0}^m f(X_p, y) \times h_{1,p}(x)$$

$$h_{1,p}(x) = \begin{cases} 0, & \text{якщо } x \leq X_{p-1} \\ \frac{x - X_{p-1}}{X_p - X_{p-1}}, & \text{якщо } X_{p-1} < x \leq X_p \\ \frac{x - X_{p+1}}{X_p - X_{p+1}}, & \text{якщо } X_p < x < X_{p+1} \\ 0, & \text{якщо } x \geq X_{p+1} \end{cases}$$

$X_{-1} = X_0 - X_1, X_{m+1} = 2 \times X_m - X_{m-1}$. Аналогічно визначаються допоміжні функції $h_{2,q}(y), q = \overline{0, n}$.

Оператори $O_{m,n}f(x, y)$ задовольняють умови: $f(x, y) \in C(\overline{D}) \Rightarrow O_{m,n}f(x, y) \in C(\overline{D})$ де $C(\overline{D})$ — клас неперервних функцій в $\overline{D}$ [13].

$$O_{m,n}f(X_k, y) = f(X_k, y), \text{ якщо } 0 \leq k \leq m$$

$$O_{m,n}f(x, Y_l) = f(x, Y_l), \text{ якщо } 0 \leq l \leq n.$$

Тобто оператор $O_{m,n}f(x, y)$ має ті ж самі сліди на прямих $x = X_k$ та $y = Y_l$, що й функція $f(x, y)$.

Ці властивості покладені нами в основу розробки і дослідження методу наближеного знаходження найбільшого та найменшого значень функції двох змінних шляхом редукції до знаходження найбільшого та найменшого значень функції до $(m+1) + (n+1) = (m+n+2)$ одновимірних задач на знаходження найбільшого та найменшого значень функції в $\overline{D}$.

Тобто будемо зводити задачу знаходження найбільшого та найменшого значень неперервної функції двох змінних в замкнутій області до $(m+n+2)$ одновимірних задач на знаходження найбільшого та найменшого значень функції.

Необхідно знайти $X_{q,e} \in [0, 1], q = \overline{0, n}$ з умови $f(X_{q,e}, Y_q) \geq f(x, Y_q), \forall x \in [0, 1]$ (якщо знаходимо максимальне значення $f(x, y)$) або $f(X_{q,e}, Y_q) \leq f(x, Y_q), \forall x \in [0, 1]$ (якщо знаходимо мінімальне значення $f(x, y)$), та знайти $Y_{p,e} \in [0, 1]$, яке задовольняє нерівність

$$f(X_p, Y_{p,e}) \geq f(X_p, y) \text{ або } f(X_p, Y_{p,e}) \leq f(X_p, y), \\ \forall p = \overline{0, m}.$$

Для знаходження розв'язків одновимірних задач оптимізації будемо використовувати один із методів, зазначених у [8].

Важливою особливістю операторів $O_{m,n}f(x,y)$ є вказане у наступних лемі та теоремах. Введемо позначення $\Delta_1 = \max_{1 \leq p \leq m} (X_p - X_{p-1})$; $\Delta_2 = \max_{1 \leq q \leq n} (Y_q - Y_{q-1})$.

Лема. Нехай $g(x) \in C[0,1]$ і має на цьому інтервалі скінченне число інтервалів монотонності (позначимо його m), тоді для кожного $\varepsilon > 0$ можна знайти таке розбиття інтервалу $[0,1]$ при $0 = x_0 < x_1 < \dots < x_m = 1$, що $|g(x) - Sp(x)| \leq \varepsilon$, де $Sp(x)$ — інтерполяційний сплайн 1 степеню для функції $g(x)$:

$$Sp(x) = \sum_{k=0}^m g(X_k) h(x, X_{k-1}, X_k, X_{k+1}),$$

де $X_{-1} = X_0 - X_1, X_{m+1} = 2X_m - X_{m-1}$.

Це твердження витікає з того, що пряма лінія сплайну, яка з'єднує початкову ($x=a$) і кінцеву ($x=b$) точки інтервалу зростання (спадання) функції $g(x)$ не може відхилитися від $y=f(x)$ більше ніж на $(b-a)$. Розбиваючи цей інтервал інтерполювання ab на Q -частин точками

$$X_0 = 0, X_1 = 0 + \frac{b-a}{Q}, X_2 = 0 + 2\frac{b-a}{Q}, \dots,$$

$$X_Q = 0 + Q\frac{b-a}{Q} = b$$

і вибираючи Q так, щоб максимум

$$|g(X_k) - g(X_{k-1})| \leq \varepsilon,$$

отримаємо, що сплайн 1 степені на цьому інтервалі при такому виборі Q буде задовольняти вимогам леми. Повторюючи сказане для кожного інтервалу монотонності $g(x)$ побудуємо інтерполяційний сплайн 1 степені, який буде наближувати функцію $g(x)$ з потрібною точністю $\varepsilon, \varepsilon \rightarrow 0$.

Теорема 1. Для похибки $Rf(x,y)$ наближення кожної неперервної функції $f(x,y)$ справедливе співвідношення

$$Rf(x,y) = \max_{(x,y) \in \overline{D}} |f(x,y) - O_{m,n}f(x,y)| \rightarrow 0,$$

якщо $\Delta_1 \rightarrow 0, \Delta_2 \rightarrow 0, m \rightarrow \infty, n \rightarrow \infty$.

Доведення. Очевидно, що функції $O_{1,m}f(x,y), O_{2,n}f(x,y) \in C(\overline{D})$ за побудовою будуть неперервними функціями. Згідно з доведеною лемою, існують такі m і n та відповідне їм розбиття, що

$$\max_{(x,y) \in \overline{D}} |f(x,y) - O_{1,m}f(x,y)| \leq \varepsilon,$$

$$\max_{(x,y) \in \overline{D}} |f(x,y) - O_{2,n}f(x,y)| \leq \varepsilon$$

У відповідності з побудовою оператора

$$O_{m,n}f(x,y) = (O_{1,m} + O_{2,n} - O_{1,m} \times O_{2,n})f(x,y)$$

для похибки

$$Rf(x,y) = \max_{(x,y) \in \overline{D}} |f(x,y) - O_{m,n}f(x,y)|$$

можна написати, що

$$R_{m,n}f(x,y) = f(x,y) - O_{1,m}f(x,y) - O_{2,n}f(x,y) + \\ + O_{1,m} \times O_{2,n}f(x,y) = (I - O_{1,m})(I - O_{2,n})f(x,y),$$

де I — тотожний оператор. Тому

$$R_{m,n}f(x,y) = \max_{(x,y) \in \overline{D}} |(I - O_{1,m})(I - O_{2,n})f(x,y)| \leq$$

$\leq \varepsilon^2 \rightarrow 0$, якщо $\varepsilon \rightarrow 0$.

Теорема 1 доведена.

Введемо позначення $g_2(x)$ та $g_1(y)$, де

$$g_2(x) = \sum_{k=0}^m Y_{2k} \times h(x, X_{2k-1}, X_{2k}, X_{2k+1}),$$

$$g_1(y) = \sum_{l=0}^n X_{1l} \times h(y, Y_{1l-1}, Y_{1l}, Y_{1l+1}).$$

Теорема 2. Для кожної точки $(x^*, y^*) \in \overline{D}$, в якій досягається найбільше (найменше) значення функції $f(x,y) \in C(\overline{D})$, яка задовольняє вимогам теореми 1, знайдуться такі m і n , і таке розбиття області D прямими $x = X_k$ та $y = Y_l$, що точка A^* — точка перетину кривих $y = g_2(x)$, $x = g_1(y)$ і точка $A_{\max}$ будуть знаходитися в одному прямокутнику Π_{k^*, l^*} , $k = \overline{0, m}, l = \overline{0, n}$, тобто в ε — околі точки $A_{\max}$, де $\varepsilon \leq \max\{\Delta_1, \Delta_2\}$.

Для доведення використаємо наступний рис.1.

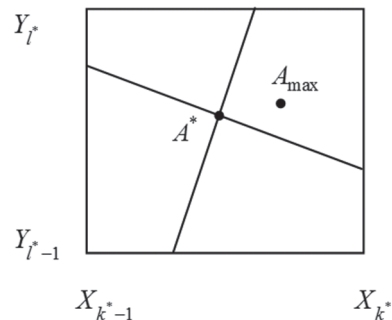


Рис. 1. Прямокутник Π_{k^*, l^*}

Оскільки функція неперервна, то знайдуться такі значення k^* та l^* , при яких $A^* \in \Pi_{k^*, l^*}$, і якщо в середині цього прямокутника функція $f(x,y)$ досягає найбільшого (найменшого) значення, то в точці $A_{\max}$:

$$f(A_{\max}) > f(X_{k^*-1}, y), f(X_{k^*}, y) \forall y \text{ та}$$

$$f(A_{\max}) > f(x, Y_{l^*-1}), f(x, Y_{l^*}) \forall x.$$

Тоді знайдуться такі m і n , що відповідне їм розбиття забезпечить виконання таких нерівностей для k^* і

$$l^* : |x^* - X_{k^*-1}|, |X_{k^*} - x^*|, |y^* - Y_{l^*-1}|, |Y_{l^*} - y^*| \leq \varepsilon.$$

є задане таке, що відповідні максимуми на лініях $x = X_{k^*-1}$ та $x = X_{k^*}$ будуть знаходитись в точках

$$(X_{k^*-1}, Y_{k^*-1,e}), (X_{k^*}, Y_{k^*,e}), (X_{l^*-1,e}, Y_{l^*-1,e}), (X_{l^*,e}, Y_{l^*,e}).$$

за умови, що

$$Y_{l^*-1} \leq Y_{k^*-1,e} \leq Y_{l^*}, Y_{l^*-1} \leq Y_{k^*,e} \leq Y_{l^*},$$

$$X_{k^*-1} \leq X_{l^*-1,e} \leq X_{k^*}, X_{k^*-1} \leq X_{l^*,e} \leq X_{k^*}.$$

Це означає, що $|f(x^*, y^*) - f(X_{k^*-1}, Y_{k^*-1,e})|$ та $|f(x^*, y^*) - f(X_{l^*-1,e}, Y_{l^*-1,e})| \leq \varepsilon \times c$, де c — деяка стала, яка залежить від функції $f(x, y)$.

Наприклад, якщо $f(x, y)$ є функцією, яка задовольняє умови

$$f(x', y') - f(x'', y'') \leq |(x' - x'')^\alpha \times (y' - y'')^\beta| \times M,$$

$$0 < \alpha, \beta < 1, M > 0$$

то вказані нерівності будуть виконуватись.

Таким чином якщо $\Delta 1 \rightarrow 0$ і $\Delta 2 \rightarrow 0$, то вказаний прямокутник стискується в точку (x^*, y^*) . Очевидно, що найбільше відхилення буде лише якщо A^* знаходиться в центрі прямокутника Π_{k^*, l^*} . Оскільки доведено, що при зменшенні кроку розбиття області наближення системою вертикальних і горизонтальних прямих в силу неперервності функції $f(x, y)$ різниця між слідами $|f(X_{k^*-1}, y) - f(X_{k^*}, y)|, |f(x, Y_{l^*-1}) - f(x, Y_{l^*})| \rightarrow 0$.

Тобто похибка наближується до 0, що й треба було довести. Теорема 2 доведена.

3. Обчислюваний експеримент

Приклад 1. Визначити найменше значення функції

$$f(x, y) = \frac{(\frac{2(x+y)+1}{4} - a)^2}{c^2} + \frac{(\frac{2(y-x)+1}{4} - b)^2}{d^2} - 1$$

за умови $a = 0,5; b = 0,5; c = 0,5; d = 1$.

Розв'язання. На рис. 2 представлено лінії рівня даної функції.

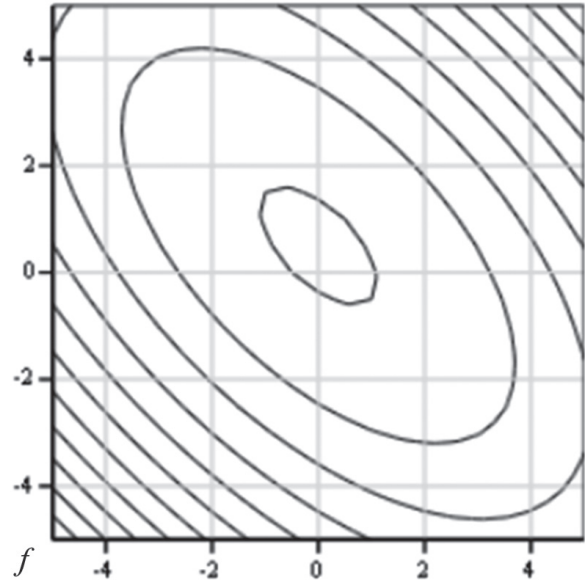


Рис. 2. Лінії рівня функції $f(x, y)$

Область D задання функції розбиваємо на прямокутні елементи лініями, паралельними осям OX та OY .

На системі паралельних вертикальних ліній, які перетинають область дослідження з кроком 0,2, знаходимо мінімальні значення даної функції $f(X_p, y), 0 \leq y \leq 1$. Позначимо значення координати y , які відповідають найменшим значенням $f(X_p, y), 0 \leq y \leq 1$ через $Y2_{pe}$. Таким чином ми отримуємо набір точок $(X2_p, Y2_{pe}), p = \overline{0, m}$ за умови, що $\forall x = X2_p$ існує лише одне значення $Y2_{pe}$. В цьому прикладі вказана умова виконується. Якщо ж значення $Y2_{pe}$ не єдине, то дії, описані нижче, потрібно буде виконати для кожного значення $Y2_{pe}$. За допомогою цих точок будемо функцію $y = g2(x)$, де

$$g2(x) = \sum_{k=0}^m Y2_k \times h(x, X2_{k-1}, X2_k, X2_{k+1}).$$

Аналогічно на системі паралельних горизонтальних ліній отримуємо набір точок $(X1_{qe}, Y1_q), Y1_q = Y_q, q = \overline{0, n}$. За допомогою цих точок будемо функцію $x = g1(y)$, де

$$g1(y) = \sum_{l=0}^n X1_l \times h(y, Y1_{l-1}, Y1_l, Y1_{l+1}),$$

де $Y1_{-1} = Y1_0 - Y1_1, Y1_{n+1} = 2Y1_n - Y1_{n-1}$.

Конкретні значення цих даних наводимо в табл. 1.

Таблиця 1

m, n	Точки $(X1_{le}, Y1_l)$ найменших значень функції $f(x, y), y \in [0, 1], l = \overline{0, 1}$	Точки $(X2_k, Y2_{ke})$ найменших значень функції $f(x, y), x \in [0, 1], k = \overline{0, m}$
5	$X1 = [0, 3; 0, 18; 0, 06; 0; 0; 0]^T$ $Y1 = [-0, 2; 0; 0, 2; 0, 4; 0, 6; 0, 8; 1; 1, 2]^T$	$X2 = [-0, 2; 0; 0, 2; 0, 4; 0, 6; 0, 8; 1; 1, 2]^T$ $Y2 = [0, 5; 0, 38; 0, 26; 0, 14; 0, 02; 0]^T$
10	$X1 = [0, 3; 0, 24; 0, 18; 0, 12; 0, 06; 0; 0; 0; 0; 0]^T$ $Y1 = [-0, 1; 0; 0, 1; 0, 2; 0, 3; 0, 4; 0, 5; 0, 6; 0, 7; 0, 8; 0, 9; 1; 1, 1]^T$	$X2 = [-0, 1; 0; 0, 1; 0, 2; 0, 3; 0, 4; 0, 5; 0, 6; 0, 7; 0, 8; 0, 9; 1; 1, 1]^T$ $Y2 = [0, 5; 0, 44; 0, 38; 0, 44; 0, 26; 0, 2; 0, 14; 0, 08; 0, 02; 0; 0]^T$

Графіки функцій $x = g1(y)$ та $y = g2(x)$ є ламаними, що з'єднують написані вище набори точок (рис. 3).

Дані лінії перетнулися в точці $A(0; 0,5)$. Координати точки знайшли за допомогою розрахунків. Значення функції в даній точці становить -1 .

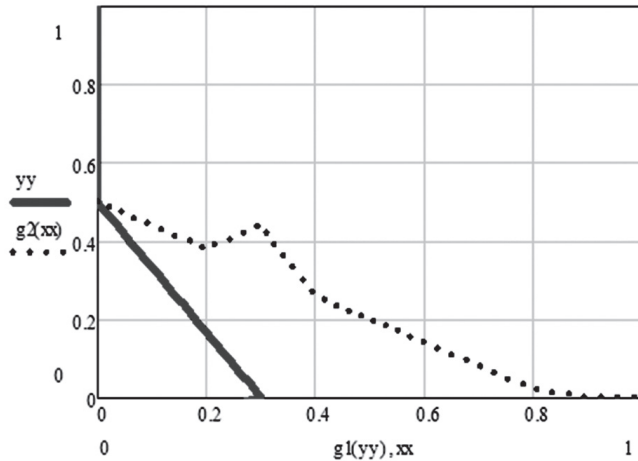


Рис. 3. Графік ліній, які поєднують точки з найменшими значеннями функції $Z = f(x, y)$ на вертикальних та горизонтальних прямих

Точні значення координат точки $A_{\text{точ}}$ такі $(6,816 \times 10^{-10}; 0,5)$. Значення функції в даній точці становить -1 .

Приклад 2. Визначити найбільше значення функції

$$f(x, y) = \frac{1}{\sqrt{c_1 \times c_2}} e^{-\frac{(x-a_1)^2}{c_1^2} - \frac{(y-b_1)^2}{c_2^2}} + \frac{1}{\sqrt{c_3 \times c_4}} e^{-\frac{(x-a_2)^2}{c_3^2} - \frac{(y-b_2)^2}{c_4^2}},$$

якщо $a_1 = 0,25; b_1 = 0,29; a_2 = 0,61; b_2 = 0,75; c_1 = 0,4; c_2 = 0,5; c_3 = 0,12; c_4 = 0,15$.

Розв'язання. На рис. 4 представлено лінії рівня даної функції. Область D задання функції розбиваємо на прямокутні елементи лініями, паралельними осям OX та OY .

На системі паралельних горизонтальних і вертикальних ліній, які перетинають область дослідження з кроком $0,1$, знайшли максимальні значення даної функції за кожною із змінних (табл. 2).

Ці дані представили у вигляді точкового графіка (рис. 5).

Проводимо лінії через точки максимуму на вибраних лініях (рис. 6). Ці лінії описані нами як графіки функцій $y = g1(x)$ та $E = g2(y)$:

$$\begin{aligned} g1(x) &= Y2_0 \times h(x, X2_0 - X2_1, X2_0, X2_1) + \\ &+ \sum_{k=1}^{m-1} (Y2_k \times h(x, X2_{k-1}, X2_k, X2_{k+1})) + \\ &+ Y2_m \times h(x, X2_{m-1}, X2_m, 2X2_m - X2_{m-1}); \\ g2(y) &= X1_0 \times h(y, Y1_0 - Y1_1, Y1_0, Y1_1) + \\ &+ \sum_{k=1}^{m-1} (X1_k \times h(y, Y1_{k-1}, Y1_k, Y1_{k+1})) + \\ &+ X1_m \times h(y, Y1_{m-1}, Y1_m, 2Y1_m - Y1_{m-1}). \end{aligned}$$

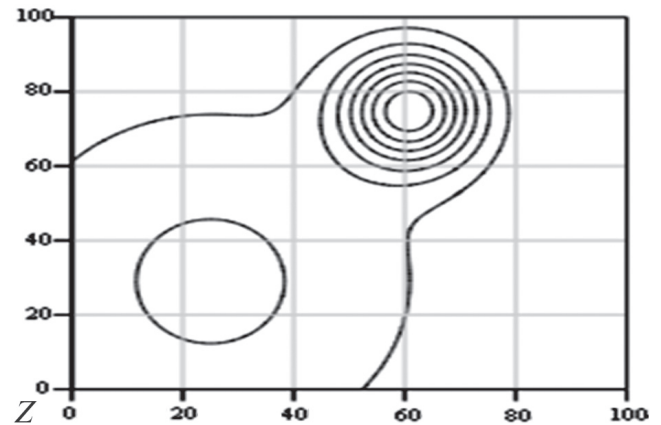


Рис. 4. Лінії рівня функції $f(x, y)$

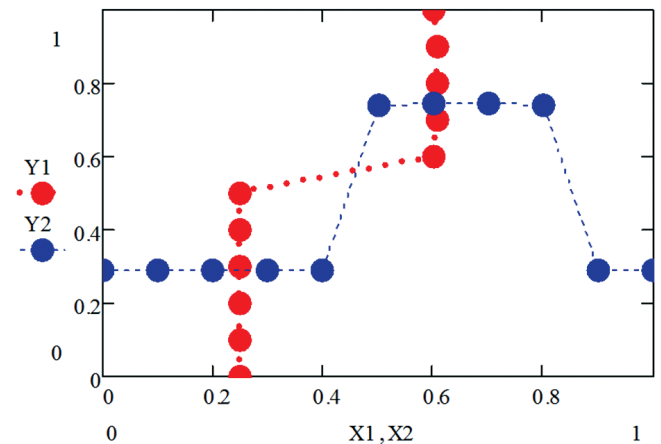


Рис. 5. Точковий графік

Отримані лінії перетнулися в трьох точках (див. рис. 6, зліва направо) $A1(0,25; 0,29)$, $A2(0,46; 0,561)$, $A3(0,608; 0,747)$. Обчислювальний експеримент показав, що значення заданої функції $f(x, y)$ в точці $A1$ дорівнює $2,236$, в точці $A2$ — $1,585$, в точці $A3$ — $7,884$. Тобто максимальне значення цієї функції в області D досягається в точці $A3(0,608; 0,747)$. При цьому значення функції $f(x, y)$ становить $7,884$.

Таблиця 2

m, n	Точки $(X1_l, Y1_l)$ найбільших значень функції $f(x = X1_k, y)$, $y \in [0, 1]$, $l = 0, n$	Точки $(X2_k, Y2_{ke})$ найбільших значень функції $f(x, y = Y2_k)$, $x \in [0, 1]$, $k = 0, m$
10	$X1 = [0,25; 0,25; 0,25; 0,25; 0,25; 0,25; 0,602; 0,608; 0,608; 0,607; 0,601]^T$ $Y1 = [0; 0,1; 0,2; 0,3; 0,4; 0,5; 0,6; 0,7; 0,8; 0,9; 1]^T$	$X2 = [0; 0,1; 0,2; 0,3; 0,4; 0,5; 0,6; 0,7; 0,8; 0,9; 1]^T$ $Y1 = [0,29; 0,29; 0,29; 0,29; 0,29; 0,742; 0,747; 0,747; 0,74; 0,29; 0,29]^T$

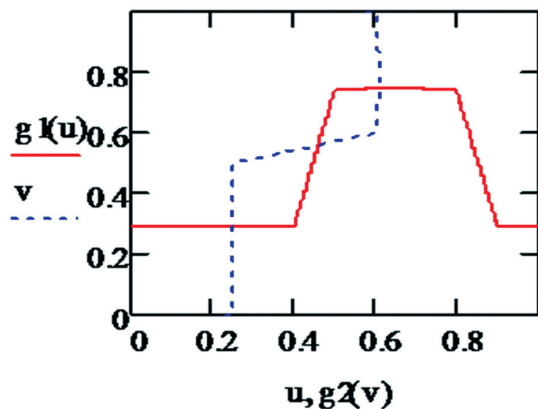


Рис. 6. Графік ліній, які поєднують точки з найменшими значеннями функції $Z = f(x, y)$ на вертикальних та горизонтальних прямих

Висновки

Таким чином в даній роботі запропоновано і досліджено метод знаходження найбільшого (найменшого) значення функції $f(x, y)$ в замкнутій області за допомогою знаходження максимальних (мінімальних) значень функції $f(x, y)$ на заданій системі взаємоперпендикулярних прямих паралельних системі координат. Результати обчислювального експерименту підтвердили його ефективність і обумовлюють можливість використання даного методу при розв'язанні багато-екстремальних задач. Автори впевнені, що даний метод може бути застосований при знаходженні найбільших або найменших значень функцій не лише двох змінних, а й трьох змінних в замкнутій області.

Список літератури:

1. *Никольский, С. М.* Курс математического анализа [Текст] / С. М. Никольский, 6-е изд., стереотип. — М. : ФИЗМАТ-ЛИТ, 2001. — 592 с.
2. *Шор, Н. З.* Методы минимизации недифференцируемых функций и их приложения [Текст] / Н. З. Шор. — К. : Наук. думка, 1979. — 200 с.
3. *Пшеничный, Б. Н.* Необходимые условия экстремума [Текст] / Б. Н. Пшеничный. — 2-е изд., перераб. и доп. — М. : Наука, 1982. — 144 с.
4. *Бахвалов, Н. С.* Численные методы (анализ, алгебра, обыкновенные дифференциальные уравнения) [Текст] / Н. С. Бахвалов. — М. : Главная редакция физико-математической литературы изд-ва "Наука", 1975. — 632 с.
5. *Акулич, И. Л.* Математическое программирование в примерах и задачах: Учеб. пособие для студентов эконом. спец. вузов / И. Л. Акулич. — М. : Высш. шк., 1986. — 319 с.
6. *Капустин, В. Ф.* Практические занятия по курсу математического программирования [Текст] / В. Ф. Капустин. — Л., Изд-во Ленингр. ун-та, 1976. — 192 с.
7. *Косолап, А. И.* Выпуклый анализ и многоэкстремальные задачи [Текст] : Моногр. / А. И. Косолап. — Днепропетровск, Изд-во ДНУ, 2007. — 280 с.
8. *Гаврилюк, І. П.* Методи обчислень [Текст] : Підручник: У 2ч. / І. П. Гаврилюк, В. Л. Макаров. — К.: Вища шк., 1995. — Ч.1., 367 с.
9. *Гаврилюк, І. П.* Методи обчислень [Текст] : Підручник: У 2ч. / І. П. Гаврилюк, В. Л. Макаров. — К.: Вища шк., 1995. — Ч.2., 431 с.
10. *Крылов, В. И.* Вычислительные методы [Текст] / В. И. Крылов, В. В. Бобков, П. И. Монастырный — М. : Главная редакция физико-математической литературы изд-ва "Наука", 1977. — Т.2. — 400 с.
11. *Демидович, Б. П.* Основы вычислительной математики [Текст] / Б. П. Демидович, И. А. Марон. — М.: Наука, 1970. — 664 с.
12. *Демидович, Б. П.* Численные методы анализа [Текст] / Б. П. Демидович, И. А. Марон, Э. З. Шувалова. — М., 1967 — 368 с.
13. *Литвин, О. М.* Інтерлінація функцій та деякі її застосування [Текст] / О. М. Литвин. — Х. : Основа, 2002. — 544 с.

Поступила до редколегії 23.11.2016

В.А. Гороховатский¹, В.С. Столяров²¹Харьковский институт ГВУЗ «Университет банковского дела»,

г. Харьков, Украина, e-mail: gorohovatsky.vl@gmail.com,

²ХНУРЭ, г. Харьков, Украина, e-mail: huccah@gmail.com

КЛАССИФИКАЦИЯ ИЗОБРАЖЕНИЙ НА ОСНОВЕ КЛАСТЕРНОГО ПРЕДСТАВЛЕНИЯ СТРУКТУРНЫХ ОПИСАНИЙ

Обсуждаются методы классификации в структурном распознавании изображений на основе преобразования их описаний к кластерному виду. Применение кластеризации существенно снижает время обработки, в то время как качественные показатели распознавания остаются на высоком уровне. Приведены результаты экспериментов по классификации на моделях структурных описаний.

РАСПОЗНАВАНИЕ ИЗОБРАЖЕНИЙ, СТРУКТУРНЫЕ МЕТОДЫ, ХАРАКТЕРНЫЕ ПРИЗНАКИ, КЛАСТЕРИЗАЦИЯ, ПРОСТРАНСТВО ВЕКТОРОВ, РЕЛЕВАНТНОСТЬ ОПИСАНИЙ, ВЕКТОР ЗНАЧИМОСТИ КЛАССОВ, КЛАССИФИКАЦИЯ

Введение

Структурное распознавание объектов в системах компьютерного зрения основано на вычислении уровня соответствия (степени релевантности) для описаний распознаваемых объектов и базы эталонов [1,2]. Описание объектов при этом представляют как множество векторов, зафиксированных в отдельных специфических точках изображения и инвариантных к допустимым геометрическим преобразованиям объектов [3-5]. Примеры изображений с координатами характерных точек, полученных методом SURF, приведены на рис. 1. Традиционно релевантность определяют как степень подобия, вычисляемую для структур данных типа множество-множество. Величина релевантности является основой для решения таких практических задач, как классификация путем сравнения с эталоном, поиск объектов на изображении, выбор идентичных объектов в базе данных и т.д.

Одним из наиболее эффективных по быстродействию способов сопоставления структурных описаний есть вычисление релевантности данных, представленных в виде интегрированных векторных описаний, формируемых путем трансформации множества признаков к обобщенному вектору, основанному на кластерном преобразовании множества структурных признаков (СП) базы эталонов [2,3].

Кластерное представление есть результат предварительного обучения системы распознавания на конкретном множестве образцов, которые составляют базис для распознавания. Векторную трансформацию структурного описания — множества можно трактовать как аппроксимацию пространства признаков системой кластеров. В результате такого представления становится возможным более продуктивно вычислять релевантность на структурах данных типа вектор-множество и вектор-вектор. В то же время возникает необходимость

детального изучения особенностей и параметрического управления такой трансформацией с точки зрения результативности построенных на её основе процедур распознавания.

Цель работы — разработка и исследование свойств методов структурного распознавания изображений на основе компрессионного кластерного представления данных. За счет применения кластерных характеристик на множестве структурных элементов обеспечивается векторное представление, значительно сокращается объем вычислительных затрат и улучшается быстродействие распознавания.

Задачи — исследование особенностей и усовершенствование информационных технологий распознавания в пространстве структурных признаков — описаний как множеств дескрипторов характерных точек изображений, а также оценивание результативности распознавания путем моделирования.



Рис. 1. Примеры изображений с координатами характерных признаков

1. Классификация на основе вычисления релевантности

На предварительном этапе обработки эталонное множество $Z = \{Z^j\}_{j=1}^J$ СП, включающее все образцы для распознавания (Z^j — эталон, J — число эталонных классов), разбивается на конечное число k кластеров $M = \{M_i\}_{i=1}^k$, так что $M_i \cap M_d = \emptyset$, $M = Z$, кластеры заданы множеством их

центроидов $m = \{m_i\}_{i=1}^k$. Кластеризация осуществляет отображение конечного множества СП в себя $Z \rightarrow Z$, причем каждый СП принадлежит одному из кластеров. В общем случае в результате кластеризации может оказаться, что $m_i \notin M_i$. Множества Z и Z^j есть мультимножества, где близкие между собой СП считаются эквивалентными [2,7]. После завершения кластеризации эталонного множества осуществляют «просеивание» множеств СП каждого из эталонов, в результате чего описание j -го эталона приобретает вид целочисленного вектора

$$H[Z^j] = (h_1, h_2, \dots, h_i, \dots, h_k)^j, \quad (1)$$

где $h_i = \text{card}\{z \mid z \in Z^j \& z \in M_i\}$, $h_i \in C$ — число элементов эталона Z^j , отнесенных к кластеру M_i , C — множество целых чисел.

Представление (1) — это образ эталона $Z^j \subseteq Z$ в кластерном отображении. Вычисление релевантности для пары образов — числовых векторов вида (1) в вычислительном плане значительно (в десятки раз [3,7]) менее затратное, чем сопоставление описания входного изображения с множествами $\{Z^j\}$. Наиболее практичные способы — вычисление расстояния или коэффициента корреляции в пространстве векторов [2].

При определении релевантности применимо также нормализованное кластерное описание объекта $O = \{o_i\}$ в виде вектора

$$H[O_n] = (h_1^n, h_2^n, \dots, h_k^n)^o, \quad h_i^n = h_i / s^o, \quad (2)$$

где $s^o = \text{card}(O) = \sum_{i=1}^k h_i$ — мощность множества O . Нормализация (2) необходима из-за разного в общем случае числа элементов в эталонах и распознаваемом объекте. Однако, в результате нормировки (2) описания, в которых значения h_i^n близки, могут быть признаны идентичными.

Фактически в результате осуществления кластеризации мы имеем дело с двумя разными представлениями эталонного множества Z : в виде набора эталонных описаний $Z = \{Z^j\}_{j=1}^J$ и в виде конечного числа кластеров $Z = M = \{M_i\}_{i=1}^k$. Целью распознавания есть идентификация объекта с одним из эталонов, а кластерное представление применяется для упрощения обработки и повышения быстродействия.

Создадим теперь метод распознавания для отнесения описания $O = \{o_i\}$ неизвестного визуального объекта к одному из классов $\{Z^j\}_{j=1}^J$ на основе кластерного описания M базы эталонов. Метод осуществляет отображение из множества описаний объектов $O \in \{O\}$ в конечное множество номеров классов $\{1, \dots, J\}$ и может быть представлен последовательностью шагов.

1. Причислим каждый элемент $o_i \in O$ распознаваемого объекта к одному из кластеров $M_i \subseteq M$ в

соответствии с правилом оптимальности

$$o_i \rightarrow M_i \mid \arg \min_d \rho(o_i, m_d) = i, \quad (3)$$

где $\rho(o_i, m_d)$ — метрика на множестве СП.

Правило (3) выполняет конкурентное отнесение СП o_i по принципу оптимальной близости в системе кластеров. Характеристикой кластера выступает его центр $m_i \in M_i$. Здесь же реализуют фильтрацию непохожих (ложных) элементов. Для этого после (3) дополнительно выполняют пороговую верификацию установленного минимума m_i : $\rho(o_i, m_i) \leq \delta$, где δ — некоторый порог значимости для минимума расстояния. Если условие не выполнено, распознаваемый элемент o_i объекта не относят ни к одному из кластеров и считают ошибочным.

2. Реализуем шаг 1 $\forall o_i \in O$, в результате формируются описания вида (1), (2) объекта как $O = (h_1, h_2, \dots, h_k)^o$, где h_i — число элементов, отнесенных в кластер M_i .

3. Вычислим степень r_j релевантности между нормализованными кластерными описаниями (2) объекта и каждого из эталонов как расстояние $r_j = \beta(H[O_n], H[Z_n^j])$ в векторном пространстве R^k .

4. Объект O отнесем к классу $d \in \{1, \dots, J\}$ по правилу $d = \arg \min_j r_j$.

Расстояния $\rho(\dots)$ для элементов и $\beta(\dots)$ для векторных описаний в общем случае различны, вместо них применимы также меры подобия, например, коэффициент корреляции, а порог δ определяется уровнем различимости в конкретной базе эталонов и допустимым уровнем помех. На этапе 4 также применима проверка значимости полученного минимума релевантности r_d : $r_d \leq \epsilon$, где ϵ — порог. При неудовлетворении неравенства класс объекта считаем неопределенным, т.к. отсутствует значимое соответствие в имеющемся пространстве эталонов.

Рассмотренный метод не учитывает важные для результативного распознавания соотношения характеристик построенных кластеров со свойствами набора взаимно различаемых эталонов в составе базы образцов. Для учета этого ключевого обстоятельства применим в процессе распознавания векторный показатель значимости классов для каждого из признаков $o_i \in O$. Значимость для СП o_i определим в виде вектора α^i весов классов, приписанного отдельному кластеру. Будем формировать α^i на этапе предварительного обучения системы распознавания, учитывая особенности конкретной базы эталонов в аспекте её кластерного представления. Определим α^i для кластера M_i в виде вектора $\alpha^i = (\alpha_1^i, \alpha_2^i, \dots, \alpha_j^i)$, компоненты которого зададим соотношением

$$\alpha_d^i = c_d^i / s_i, \quad (4)$$

где $c_d^i = \text{card}\{z \mid z \in Z^d \ \& \ z \in M_i\}$ — число элементов эталонного класса Z^d , определенных за результатами обучения в кластер M_i , $s_i = \text{card}(M_i)$ — число элементов кластера M_i .

Значение (4) соответствует нормированному представлению (2) для кластерного разложения эталона по множеству классов. Из (4) видно, что $\sum_{d=1}^J \alpha_d^i = 1$, так что α_d^i можно считать оценкой вероятности события, что элемент кластера M_i относится к классу Z^d . Величины α_d^i есть относительные веса СП кластера в плане его отнесения к эталонам классов. В частности, максимальное значение среди компонентов α^i соответствует наиболее вероятному классу, к которому может принадлежать СП. Очевидно, что значения характеристик (4) непосредственно зависят как от содержания базы эталонов, так и от применяемого метода кластеризации, который устанавливает отображение эталон-кластер.

Модифицируем метод распознавания включением вектора весов α^i в процесс принятия решения. Отнесение элемента o_i к одному из кластеров на шаге 1 теперь сопровождаем формированием суммы значений априорных векторов α^i , в результате чего после анализа всех значимых $o_i \in O$ на шаге 2 будет сформирован суммарный для элементов объекта вектор характеристик классов

$$\Sigma = \sum_O \alpha^i = (\Sigma_1, \Sigma_2, \dots, \Sigma_J), \quad \Sigma_d = \sum_i \alpha_d^i. \quad (5)$$

Суммирование в (5) осуществляется по всем значимым соответствиям кластерам для множества СП объекта.

Модификация п. 4 метода приобретает вид: объект относим к классу $d \in \{1, \dots, J\}$ по правилу:

$$O \rightarrow d \mid d = \arg \max_j \Sigma_j. \quad (6)$$

В модифицированном методе распознавания также могут быть применены проверки на значимость при отнесении к кластеру (на шаге 1) и на значимость максимума подобия в (6) (шаг 4). В частности, значимость максимума в (6) может быть подтверждена степенью преобладания глобального максимума над ближайшим из локальных [2].

Как видим, первый метод базируется на процедуре голосования (3), в основе которой лежит конкурентный выбор центра кластера для каждого СП объекта и анализ подобия кластерного описания объекта на множестве эталонов. Здесь использовано разложение в разрезе «эталон по кластерам». Идея же модифицированного метода основана на построении интегрального вероятностного распределения кластеров на множестве СП объекта в имеющемся диапазоне классов. Здесь информационной основой есть разложение «кластер по эталонам». Это, по нашему разумению,

должно обеспечить более глубокое согласование со значениями СП эталонных образцов, что в целом улучшит результативность распознавания. Модифицированный метод также можно считать голосованием СП объекта, взвешенным распределением имеющихся классов базы изображений по построенному множеству кластеров.

В прикладных исследованиях при наличии помех целесообразно рекомендовать схему анализа на шаге 1 рассмотренных методов, когда каждый кластер базы эталонов «отбирает» соответствующие ему элементы объекта, т.е. поиск соответствия для СП осуществляется «от эталона», а не «от объекта». Нормализация при этом может быть выполнена на основе числа отобранных признаков анализируемого объекта.

2. Анализ качества кластерного представления

Обобщая проведенный анализ, отметим, что исходными данными для распознавания есть целочисленная матрица кластерного представления H размерами $J \times k$, где число строк равно количеству эталонов, а число столбцов — количеству кластеров. Строка матрицы — это разложение имеющихся данных вида «эталон по кластерам», а столбец — вида «кластер по эталонам». Каждый из срезов (строки, столбцы) допускает проведение дополнительной обработки в целях улучшения свойств распознавания. Однако при этом нужно вводить изменения и в метод распознавания. Другим направлением усовершенствования есть изменение результатов кластеризации за счет применения других методов.

Матрица кластерного представления H есть основой процедуры распознавания. Проанализируем теперь возможности изучения ее качественных свойств в целях улучшения распознавания, а также возможность ее дополнительной логической обработки. Заметим, что изменение матрицы повлияет и на весовые коэффициенты α^i . Такую обработку осуществляют на предварительном этапе, поэтому время распознавания не возрастает. Дополнительная логическая обработка может быть необходима в целях дальнейшего сжатия и уточнения трансформированных описаний, исключения кластеров незначительного объема, т.е. для формирования конструктивной системы кластеров, в рамках которой можно реализовать качественное распознавание. Другим приемом обработки может быть исключение (обнуление) в столбцах матрицы незначимых по сравнению с остальными элементов.

Для анализа или логической обработки значений h_i столбцов матрицы H , характеризующих кластер M_i , применим критерий [2]

$$\gamma_i = \sum_{a=1}^{J-1} \sum_{b>a} |h_i[a] - h_i[b]|, \quad (7)$$

где $h_i[a]$ – компонента i -го столбца.

Величина (7) равна нулю в случае полного совпадения значений в столбце и возрастает с увеличением их различия между собой. Вместо (7) можно применить также дисперсионные характеристики значений столбцов. Обработка может состоять в вычислении величины (7) для столбцов матрицы H и отбрасывании как непригодных для распознавания столбцов с наименьшими значениями критерия. В результате получим усеченную по числу столбцов матрицу H . Заметим, что введение логической обработки приводит к трансформации имеющихся данных и к необходимости в связи с этим некоторого изменения алгоритма распознавания. Окончательная эффективность таких эвристических обработок может быть оценена только на основе экспериментов по распознаванию в прикладной задаче. В то же время, величина компонент вектора γ_i может выступить критерием качества кластерного представления: чем более удалены значения (7) от нуля, тем лучше качество кластеризации в аспекте распознавания по сформированной системе кластеров.

3. Результаты экспериментов

В ходе моделирования в ряде наших экспериментов сформированы равномерно распределенные на множестве $\{0, \dots, 9\}$ дискретные значения СП в составе структурных описаний, где число кластеров взято равным $k=10$. В качестве $p(\dots)$ для элементов использована метрика изолированных точек (СП равны в случае их совпадения), а в качестве меры $\beta(\dots)$ для кластерных описаний выбрана евклидова метрика. Для объема эталонных описаний в 100 элементов многократное вычисление релевантности показало, что значение r для разных реализаций образцов в евклидовой метрике было не меньше величины 0,16, что говорит о достаточной надежности и работоспособности разработанного метода сопоставления описаний даже в условиях одинаково распределенных значений СП разных эталонов. В целом полученный диапазон значений релевантности в этом эксперименте определяется отрезком $[0,16-0,34]$.

Для этих исходных данных для 4-х равномерно распределенных эталонных множеств при $k=10$ вычислены представления $H[Z^j]$ (табл. 1) и веса α^i кластеров (табл. 2).

Как видим, веса кластеров в столбцах табл. 2, значения которых непосредственно используются во втором методе, оказались достаточно близкими из-за того, что эталоны имеют одинаковое распределение. Наиболее благоприятные для

распознавания весомые различия между эталонами наблюдаются в кластерах 1, 3, 5, 7, а в кластере 9 имеем полное совпадение значений, что говорит о невозможности различения эталонов в рамках этого кластера. Эти факторы в целом затрудняют процесс распознавания. Примером логической обработки здесь может быть удаление из рассмотрения кластера 9 (одинаковые элементы), а также обнуление значения 0,08 в кластере 5 табл. 2 (незначимость). Значения критерия γ в табл. 1 позволяют установить полезность кластеров: кластеры 8, 9, 10 наименее пригодны, а кластеры 1, 4, 5, 7 позволяют в наилучшей степени различать классы. В других экспериментах со случайно выбранными значениями данных величина компонент вектора (7) достигала значений 70.

Таблица 1

Кластерные представления $H[Z^j]$ и значения (7)

H	M_1	M_2	M_3	M_4	M_5	M_6	M_7	M_8	M_9	M_{10}
Z^1	7	9	9	9	14	15	6	11	11	9
Z^2	8	11	8	16	3	14	13	9	11	7
Z^3	15	7	3	13	9	18	7	8	11	9
Z^4	10	8	8	13	13	12	7	11	11	7
γ	26	13	18	21	37	19	21	11	0	8

Таблица 2

Весовые коэффициенты α^i

α^i	M_1	M_2	M_3	M_4	M_5	M_6	M_7	M_8	M_9	M_{10}
Z^1	0.17	0.26	0.32	0.18	0.36	0.25	0.18	0.28	0.25	0.28
Z^2	0.20	0.31	0.29	0.31	0.08	0.24	0.39	0.23	0.25	0.22
Z^3	0.38	0.20	0.11	0.25	0.23	0.31	0.21	0.21	0.25	0.28
Z^4	0.25	0.23	0.29	0.25	0.33	0.20	0.21	0.28	0.25	0.22

Отметим, что в реальных условиях структурного распознавания ситуация с низкими значениями критерия (7) возникает достаточно редко [6], потому что структурные описания изображений все-таки значимо отличаются между собой. В то же время преобладание отдельных значений в столбцах матрицы H свойственны СП изображений.

Для второго метода с использованием весов α^i для тех же исходных данных в исследовании также достигнута безошибочная классификация. Для входных векторов эталонов с номерами 1, 4, где получена релевантность $r=0,16$ по первому методу, второй метод также показал незначительное отличие в значениях вектора Σ : для эталонов 1 и 4 он равен 25,48, а при сравнении эталона 4 с самим собой – 25,57. Такие ничтожные различия подчеркивают более высокую чувствительность второго метода.

Проведены эксперименты с учетом влияния помех, когда каждая компонента исходного

описания (случайный вектор из 100 элементов) изменялась с заданной вероятностью θ на другое случайное значение, выбранное из того же диапазона. При значениях θ на промежутке $[0, \dots, 0,7]$ вероятность правильного распознавания случайных векторов уменьшалась от 1 до 0,47. При этом для малого уровня помех $\theta < 0,3$ большая вероятность наблюдалась у первого метода (вероятность более 0,85), в то время как при значительном уровне помех $\theta \geq 0,3$ преимущество имеет второй метод (вероятность равна 0,7). Этот факт можно объяснить тем, что первый метод строит и сопоставляет интегральные по кластерной системе описания, в то время как второй метод классифицирует каждый элемент в отдельности, что делает этот подход более конкурентоспособным.

В целях более глубокого изучения свойств кластерного представления выбраны другие реализации эталонов, обладающие существенным различием в построенном пространстве. Конкретно для первого метода расстояние между выбранными описаниями составило более 0,42. Соответственно усилились различия и для второго метода. По результатам экспериментов в условиях помех представлена табл. 3.

Таблица 3

Значения вероятности правильного распознавания
в условиях помех

Уровень помех θ	0,1	0,2	0,3	0,4	0,5	0,6
Первый метод	1,0	1,0	0,98	0,95	0,86	0,73
Второй метод	1,0	1,0	0,99	0,97	0,91	0,80

Как видим из табл. 3, оба метода обладают высокой помехоустойчивостью: при искажении до 30% числа СП из описаний они обеспечивают практически безошибочное распознавание с вероятностью выше 0,98. В сравнительном плане второй метод обеспечивает более устойчивое функционирование в условиях помех, т.к. при уровне $\theta \in [0, \dots, 0,3]$ вероятности для обоих методов примерно одинаковы, а при $\theta > 0,3$ значения вероятности для второго метода несколько выше.

Выводы

Рассмотренные методы реализуют процедуру распознавания визуального объекта на основе кластерного представления в пространстве структурных признаков в форме двух срезов: «эталон по кластерам» и «кластер по эталонам». Критерий (7) при этом описывает качество представления второго типа. Аналогично (7) можно применить некоторый критерий и к строкам матрицы табл. 1. Однако, разнообразие внутри строки для первого метода не так существенно влияет на распознавание, как

различия между строками. Эти различия отражаются в матрице расстояний между эталонами в построенном пространстве признаков.

Каждый из предложенных методов имеет свои особенности в плане результативности. Оба метода опираются на использованную процедуру кластеризации. В примере сложной ситуации равномерного распределения признаков эталонов по классам более эффективным оказался второй метод. В сравнительном плане значения вероятности правильного распознавания второй метод обеспечивает более устойчивое функционирование в условиях помех.

Научная новизна исследования состоит в синтезе метода структурного распознавания изображений на основе кластерного представления описаний. Переход к векторному виду существенно увеличивает быстродействие распознавания за счет упрощения обработки.

Практическая ценность работы — получение прикладных программных моделей для модификаций структурного распознавания и подтверждение результативности предложенной обработки в конкретных примерах модельных сигналов.

Рассмотренные методы достаточно универсальны и могут быть использованы для анализа и классификации произвольных наборов данных, представленных в виде множеств.

Список литературы:

1. Берестовский А.Е. Нейросетевые технологии самообучения в системах структурного распознавания визуальных объектов / А.Е. Берестовский, А.Н. Власенко, В.А. Гороховатский // Реєстрація, зберігання і обробка даних. — 2015. — Т. 17, № 1. — С. 108—120.
2. Гороховатский В.А. Структурный анализ и интеллектуальная обработка данных в компьютерном зрении: монография / В.А. Гороховатский. — Х.: Компания СМІТ, 2014. — 316 с.
3. Gorokhovatsky V.A. Efficient Estimation of Visual Object Relevance during Recognition through their Vector Descriptions / V.A. Gorokhovatsky // Telecommunications and Radio Engineering. — 2016, Vol. 75, No 14. — P. 1271—1283.
4. Bay H. Surf: Speeded up robust features / H. Bay, T. Tuytelaars, L. Van Gool // European Conference on Computer Vision. — 2006. — P.404—417.
5. Szeliski R. Computer Vision: Algorithms and Applications / R. Szeliski. — London: Springer, 2010. — 979 p.
6. Гороховатский В.А. Структурное распознавание изображений с применением моделей интеллектуальной обработки и самоорганизации признаков / В.А. Гороховатский, А.В. Гороховатский, А.Е. Берестовский // Радиоэлектроника, информатика, управление.—2016. — №3 (38). — С. 39—46.
7. Гороховатский В.А. Изучение свойств методов кластеризации применительно к множествам характерных признаков изображений / В.А. Гороховатский, М.Д. Дунаевская, В.А. Струненко // Системи обробки інформації. — 2016. — Вип. 5 (142). — С. 124—127.

Поступила в редколлегию 12.10.2016

УДК 519.6

О.М. Литвин¹, О.О. Литвин¹, Г.Д. Лісний², О.В. Славик¹¹ Українська інженерно-педагогічна академія, м. Харків, Україна,² Київський університет імені Т. Шевченка, м. Київ, Україна

ВІДНОВЛЕННЯ ЗОБРАЖЕНЬ В ЗОНАХ ВІДСУТНОСТІ ПОПІКСЕЛЬНОЇ ІНФОРМАЦІЇ З ВИКОРИСТАННЯМ ІНТЕРСТРІПАЦІЇ ФУНКЦІЙ

Задача відновлення зображення в зонах відсутності попиксельної інформації є вкрай важливою. Такі задачі виникають в машинобудуванні, сейсмографії, обробці даних дистанційного зондування Землі тощо. В даній роботі проведено огляд існуючих методів відновлення пошкоджених цифрових зображень. Наведено стандартний метод інтерстріпації функції двох змінних. Запропоновано новий модифікований метод текстурної інтерстріпації для відновлення зображення поверхні за неповною інформацією про неї.

ЗОБРАЖЕННЯ, ВІДНОВЛЕННЯ ЗОБРАЖЕНЬ, ІНТЕРСТРІПАЦІЯ, ІНТЕРЛІНАЦІЯ, ТЕКСТУР-НА ІНТЕРСТРІПАЦІЯ

Вступ

Інколи в файлах, які містять графічну інформацію виявляються дефекти. Наприклад, в процесі передачі по мережі вони можуть бути пошкоджені в результаті помилок при передачі даних або перенавантаженню мережі. Оцінка справжніх значень втрачених пікселів необхідна в більшості задач цифрової обробки зображень або, наприклад, в задачах обробки архівних документів у вигляді зображень, що мають різноманітні спотворення (подряпини, плями, пил, непотрібні написи, лінії згину тощо). Тому актуальною є розробка методів для відновлення зображення в тих його частинах, де інформація з тих чи інших причин відсутня або якщо вона відома неповністю. Більшість методів відновлення зображень можна умовно поділити на наступні групи [1]: текстурні, шаблонні, основані на рівняннях в частинних похідних, гібридні та швидкі напівавтоматичні. В роботі [11] було запропоновано метод інтерстріпації для відновлення функції двох змінних в точках між смугами за допомогою інформації про цю функцію, яка відома лише в точках заданої системи смуг.

Аналіз літературних джерел. Позначимо множину пікселів в невідомій області через D , а множину коректних пікселів через D . Позначимо піксель зображення I , який знаходиться на перетині i -ї строки та j -го стовпця через $I(i, j)$, а квадратну область розміром $m \times m$, ($m \geq 1$) на зображенні I , центральний піксель якого має координату (i, j) , через $B(i, j)$.

Вище вже було зазначено, що більшість методів відновлення зображень можна умовно поділити на групи [1]. Розглянемо їх більш детально.

Текстурні методи відновлення зображень для заповнення невідомої області $\bar{D}$ безпосередньо використовують пікселі з відомої області зображення D . Головна відмінність між цими алгоритмами полягає в забезпеченні неперервності на

границі області D [1]. В роботах [2] та [3] запропоновані методи текстурного відновлення зображення, які відрізняються способом відновлення різних кольорів, інтенсивності, градієнта та навіть статистичних характеристик.

Основна ідея роботи шаблонних методів відновлення зображень заснована на припущенні про наявність повторюваних фрагментів даних на зображенні, які зазвичай називаються шаблонами. Відновлення області $\bar{D}$ проводиться частинами шляхом копіювання значень яскравості з найбільш схожого шаблону. Невідома область може містити як текстурну, так і структурну інформацію. В роботі [4] було показано, що для досягнення більш високого рівня відновлення, необхідно вміти знаходити та відрізняти структурну та текстурну складові з метою відновлення в першу чергу саме структурної складової. В роботі [5] було запропоновано ітераційний алгоритм для заповнення області D . Особливо виділяється робота [6], де на відміну від всіх вищеописаних робіт, для заповнення пошкодженої області використовується база даних зображень, яка містить мільйони зображень-шаблонів для відновлення.

Основу для методів відновлення зображень, основаних на рівняннях з частинними похідними, було закладено в роботі [7]. В цій роботі відновлення даних області $\bar{D}$ проводиться за допомогою даних, що є природним продовженням інформації, яка міститься в D . Цей підхід став основою для наступних робіт. Так, наприклад, в роботі [8] було запропоновано алгоритм анізотропної вектор-регуляризації.

Гібридні методи відновлення зображень являють собою поєднання двох класів методів. А саме текстурних методів та методів, основаних на використанні диференціальних рівнянь з частинними похідними. Основна ідея алгоритму полягає в тому, що перш за все виділяють текстурну та структурну

складову зображення, які потім заповнюються відповідними алгоритмами [1].

Недоліком більшості представлених вище методів є їх висока обчислювальна складність, тому в деяких працях застосовують швидкі напівавтоматичні методи відновлення зображень для прискорення обчислень. В роботі [9] було наведено метод відновлення зображення за допомогою виділеної структури. Автори роботи [10] запропонували відновлення зображення з використанням ітеративної згортки зображення з дифузним ядром.

Робота [11] була присвячена розробці методу інтерстріпації неперервних функцій.

Інтерстріпацією (від англ. *inter* – між, від англ. *stripe* – смуга) функції двох змінних називається відновлення цієї функції між системою смуг, якщо інформація про цю функцію відома лише в точках вказаних смуг.

Даний метод дозволяє відновлювати поверхню за відомою інформацією про неї на смугах у випадку якщо границі смуг є неперетинними прямими паралельними осям координат.

Постановка задачі. Необхідно відновити пошкоджене зображення деякої поверхні Σ . Вважаємо, що зображення поверхні Σ відоме лише на системі m ($m \geq 2$) вертикальних смуг вигляду:

$$D_{1,k} = \{\alpha_k \leq x \leq \beta_k\}, k = \overline{1, m}$$

та на системі n ($n \geq 2$) горизонтальних смуг вигляду:

$$D_{2,l} = \{\gamma_l \leq y \leq \delta_l\}, l = \overline{1, n}$$

Невідомі смуги зображення задаються наступним чином для невідомих вертикальних смуг:

$$\overline{D}_{1,k,k+1} = \{\beta_k \leq x \leq \alpha_{k+1}\}, k = \overline{1, m-1}$$

та для невідомих горизонтальних смуг:

$$\overline{D}_{2,l,l+1} = \{\delta_l \leq y \leq \gamma_{l+1}\}, l = \overline{1, n-1}$$

Тоді об'єднання множин $\overline{D}_{1,k,k+1}, k = \overline{1, m-1}$ та $\overline{D}_{2,l,l+1}, l = \overline{1, n-1}$ дає область $\overline{D}$ незаповнених ділянок зображення. В точках зображення D , які не потрапили до $\overline{D}$ зберігається вся наявна інформація про зображення.

Поверхня $\Sigma: z = f(x, y)$, $f(x, y) \in C^{N,N}(R^2)$, яку ми хочемо відновити, вважається відомою лише на вказаних смугах, тобто

$$f(x, y)|_{\alpha_k \leq x \leq \beta_k} = f_{1,k}(x, y), \alpha_k \leq x \leq \beta_k; \gamma_l \leq y \leq \delta_{l+1}$$

$$f(x, y)|_{\gamma_l \leq y \leq \delta_l} = f_{2,l}(x, y), \gamma_l \leq y \leq \delta_l; \alpha_1 \leq x \leq \beta_{m+1}$$

При цьому

$$\alpha_k < \beta_k < \alpha_{k+1} < \beta_{k+1}, k = \overline{1, m},$$

$$\gamma_l < \delta_l < \gamma_{l+1} < \delta_{l+1}, l = \overline{1, n}.$$

$C^{N,N}(R^2)$ – клас функцій, які мають неперервні похідні $f^{(p,q)}(x, y)$ для $0 < p, q \leq N$.

1. Інтерстріпація на системі вертикальних смуг

Вважаємо, що зображення поверхні Σ відоме лише на системі m ($m \geq 2$) вертикальних смуг $D_{1,k}, k = \overline{1, m}$.

Ведемо до розгляду такий оператор

$$T_{1,k,k+1}f(x, y) = \frac{x - \beta_k}{\alpha_{k+1} - \beta_k} \Delta_{1,k,k+1}f(x, y) + \frac{x - \alpha_{k+1}}{\beta_k - \alpha_{k+1}} \Delta_{2,k,k+1}f(x, y).$$

Оператори $\Delta_{1,k,k+1}f(x, y)$ та $\Delta_{2,k,k+1}f(x, y)$ обираються наступним чином:

$$\Delta_{1,k,k+1}f(x, y) = \frac{1}{(2\rho + 1)^2} \sum_{i=-\rho}^{\rho} \sum_{j=-\rho}^{\rho} f(\beta_k - \rho + i, y + j),$$

$$\Delta_{2,k,k+1}f(x, y) = \frac{1}{(2\rho + 1)^2} \sum_{i=-\rho}^{\rho} \sum_{j=-\rho}^{\rho} f(\alpha_{k+1} + \rho + i, y + j).$$

де $\rho = \rho_{1,k,k+1}(x, \beta_k, \alpha_{k+1}) = \min(x - \beta_k, \alpha_{k+1} - x)$ – мінімальна відстань від точки $I(x, y)$ до границь невідомої вертикальної смуги $\overline{D}_{1,k,k+1}$, яка знаходиться між смугами $D_{1,k}$ та $D_{1,k+1}$.

Оператори $\Delta_{1,k,k+1}f(x, y)$ та $\Delta_{2,k,k+1}f(x, y)$ обчислюють середнє значення освітленості зображення в пікселях, які попадають в квадратні області $B(\beta_k - \rho, y)$ та $B(\alpha_{k+1} + \rho, y)$ відповідно. розміром $(2\rho + 1) \times (2\rho + 1)$ кожна.

Зауваження 1. В залежності від області точки з квадратних областей $B(\beta_k - \rho, y)$ та $B(\alpha_{k+1} + \rho, y)$ можуть братися з невідомих областей зображення або поза ним, тому слід перевіряти їх на належність відомим областям $D_{1,k}, k = \overline{1, m}$. Аналогічне зауваження справедливе і для операторів $\Delta_{3,l,l+1}f(x, y)$ та $\Delta_{4,l,l+1}f(x, y)$, наведених нижче.

Зауваження 2. Для зменшення обчислювальної складності оператори $\Delta_{1,k,k+1}f(x, y)$ та $\Delta_{2,k,k+1}f(x, y)$ можна обрати наступним чином:

$$\Delta_{1,k,k+1}f(x, y) = f(\beta_k - \rho, y),$$

$$\Delta_{2,k,k+1}f(x, y) = f(\alpha_{k+1} + \rho, y).$$

Тобто, замість обчислення середнього значення освітленості в відповідній квадратній області брати значення освітленості в центрі цієї квадратної області.

Ведемо до розгляду такий оператор

$$\Theta_1 f(x, y) = \begin{cases} f_{1,k}(x, y) & (x, y) \in D_{1,k}, k = \overline{1, m}; \\ T_{1,k,k+1}f(x, y) & (x, y) \in \overline{D}_{1,k,k+1}, k = \overline{1, m-1}; \end{cases}$$

Поверхня $z = \Theta_1 f(x, y)$ є наближеною математичною моделлю освітленості поверхні Σ , яка на кожній смузі $D_{1,k}, k = \overline{1, m}$ точно відновлює поверхню, а між смугами зображує поверхню за допомогою оператора $T_{1,k,k+1}f(x, y)$.

Зауваження 3. Якщо $\rho = 0$, то отримаємо стандартний метод інтерстріпації на системі вертикальних смуг, описаний у [11].

2. Інтерстріпація на системі горизонтальних смуг

Вважаємо, що зображення поверхні Σ відоме лише на системі n ($n \geq 2$) горизонтальних смуг $D_{2,l}$, $l = \overline{1, n}$.

Ведемо до розгляду такий оператор

$$T_{2,l,l+1}f(x,y) = \frac{y - \delta_l}{\gamma_{l+1} - \delta_l} \Delta_{3,l,l+1}f(x,y) + \frac{y - \gamma_{l+1}}{\delta_l - \gamma_{l+1}} \Delta_{4,l,l+1}f(x,y).$$

Оператори $\Delta_{3,l,l+1}f(x,y)$ та $\Delta_{4,l,l+1}f(x,y)$ обираються наступним чином:

$$\Delta_{3,l,l+1}f(x,y) = \frac{1}{(2\rho+1)^2} \sum_{i=-\rho}^{\rho} \sum_{j=-\rho}^{\rho} f(x+i, \delta_l - \rho + j),$$

$$\Delta_{4,l,l+1}f(x,y) = \frac{1}{(2\rho+1)^2} \sum_{i=-\rho}^{\rho} \sum_{j=-\rho}^{\rho} f(x+i, \gamma_{l+1} + \rho + j),$$

де $\rho = \rho_{2,l,l+1}(y, \delta_l, \gamma_{l+1}) = \min(y - \delta_l, \gamma_{l+1} - y)$ — мінімальна відстань від точки $I(x,y)$ до границь невідомої горизонтальної смуги $\overline{D}_{2,l,l+1}$, яка знаходиться між смугами $D_{2,l}$ та $D_{2,l+1}$.

Оператори $\Delta_{3,l,l+1}f(x,y)$ та $\Delta_{4,l,l+1}f(x,y)$ обчислюють середнє значення освітленості зображення в пікселях, які попадають в квадратні області $B(x, \delta_l - \rho)$ та $B(x, \gamma_{l+1} + \rho)$ відповідно. розміром $(2\rho+1) \times (2\rho+1)$ кожна.

Зауваження 4. Для зменшення обчислювальної складності оператори $\Delta_{3,l,l+1}f(x,y)$ та $\Delta_{4,l,l+1}f(x,y)$ можна обрати наступним чином:

$$\Delta_{3,l,l+1}f(x,y) = f(x, \delta_l - \rho),$$

$$\Delta_{4,l,l+1}f(x,y) = f(x, \gamma_{l+1} + \rho).$$

Ведемо до розгляду такий оператор

$$\Theta_2f(x,y) = \begin{cases} f_{2,l}(x,y) & (x,y) \in D_{2,l}, l = \overline{1, n}; \\ T_{2,l,l+1}f(x,y) & (x,y) \in \overline{D}_{2,l,l+1}, l = \overline{1, n-1} \end{cases}$$

Поверхня $z = \Theta_2f(x,y)$ є наближеною математичною моделлю поверхні Σ , яка на кожній смузі $D_{2,l}$, $l = \overline{1, n}$ точно відновлює поверхню, а між смугами зображує поверхню за допомогою оператора $T_{2,l,l+1}f(x,y)$.

Зауваження 5. Якщо $\rho = 0$, то отримаємо стандартний метод інтерстріпації на системі горизонтальних смуг, описаний у [11].

3. Інтерстріпація на системі взаємоперпендикулярних смуг

Вважаємо, що зображення поверхні Σ відоме лише на системі m ($m \geq 2$) вертикальних смуг $D_{1,k}$, $k = \overline{1, m}$, та на системі n ($n \geq 2$) горизонтальних смуг $D_{2,l}$, $l = \overline{1, n}$.

В результаті їх об'єднання отримаємо набір прямокутних областей

$\overline{\Pi}_{k,l} = [\beta_k, \alpha_{k+1}] \times [\delta_l, \gamma_{l+1}]$, $k = \overline{1, m-1}$, $l = \overline{1, n-1}$, інформацію в яких треба відновити.

Ведемо до розгляду такий оператор

$$T_{1,2,k,l}f(x,y) = (T_{1,k,k+1} + T_{2,l,l+1} - T_{1,k,k+1}T_{2,l,l+1})f(x,y) =$$

$$= \frac{x - \beta_k}{\alpha_{k+1} - \beta_k} \Delta_{1,k,k+1}f(x,y) + \frac{x - \alpha_{k+1}}{\beta_k - \alpha_{k+1}} \Delta_{2,k,k+1}f(x,y) +$$

$$+ \frac{y - \delta_l}{\gamma_{l+1} - \delta_l} \Delta_{3,l,l+1}f(x,y) + \frac{y - \gamma_{l+1}}{\delta_l - \gamma_{l+1}} \Delta_{4,l,l+1}f(x,y) -$$

$$- \frac{x - \beta_k}{\alpha_{k+1} - \beta_k} \frac{y - \delta_l}{\gamma_{l+1} - \delta_l} \Delta_{1,k,k+1} \Delta_{3,l,l+1}f(x,y) -$$

$$- \frac{x - \beta_k}{\alpha_{k+1} - \beta_k} \frac{y - \gamma_{l+1}}{\delta_l - \gamma_{l+1}} \Delta_{1,k,k+1} \Delta_{4,l,l+1}f(x,y) -$$

$$- \frac{x - \alpha_{k+1}}{\beta_k - \alpha_{k+1}} \frac{y - \delta_l}{\gamma_{l+1} - \delta_l} \Delta_{2,k,k+1} \Delta_{3,l,l+1}f(x,y) -$$

$$- \frac{x - \alpha_{k+1}}{\beta_k - \alpha_{k+1}} \frac{y - \gamma_{l+1}}{\delta_l - \gamma_{l+1}} \Delta_{2,k,k+1} \Delta_{4,l,l+1}f(x,y).$$

Оператори $\Delta_{1,k,k+1}f(x,y)$, $\Delta_{2,k,k+1}f(x,y)$, $\Delta_{3,l,l+1}f(x,y)$ та $\Delta_{4,l,l+1}f(x,y)$ наведені вище.

Введемо до розгляду такий оператор

$$\Theta_{12}f(x,y) = \begin{cases} f_{1,k}(x,y) & (x,y) \in D_{1,k}, k = \overline{1, m}; \\ f_{2,l}(x,y) & (x,y) \in D_{2,l}, l = \overline{1, n}; \\ T_{1,2,k,l}f(x,y) & (x,y) \in \overline{\Pi}_{k,l}, \\ & k = \overline{1, m-1}, l = \overline{1, n-1}. \end{cases}$$

Поверхня $z = \Theta_{12}f(x,y)$ є наближеною математичною моделлю поверхні Σ , яка на кожній із смуг $D_{1,k}$, $k = \overline{1, m}$ та $D_{2,l}$, $l = \overline{1, n}$ точно відновлює поверхню, а на невідомих прямокутних областях $\overline{\Pi}_{k,l} = [\beta_k, \alpha_{k+1}] \times [\delta_l, \gamma_{l+1}]$, $k = \overline{1, m-1}$, $l = \overline{1, n-1}$ зображує поверхню за допомогою оператора $T_{1,2,k,l}f(x,y)$.

Зауваження 6. Оператор $T_{1,2,k,l}f(x,y)$ — оператор інтерлінації (від англ. *inter* — між, від англ. *line* — лінія) функції двох змінних, який при $\rho = 0$, є класичним оператором інтерлінації між чотирма сторонами довільного прямокутника

$$\overline{\Pi}_{k,l} = [\beta_k, \alpha_{k+1}] \times [\delta_l, \gamma_{l+1}], k = \overline{1, m-1}, l = \overline{1, n-1},$$

наведеним у [12] та [13].

4. Інтерстріпація на системі кадрів

Як і у випадку текстурної інтерстріпації вважаємо, що зображення поверхні Σ відоме лише на системі m ($m \geq 2$) вертикальних смуг $D_{1,k}$, $k = \overline{1, m}$, та на системі n ($n \geq 2$) горизонтальних смуг $D_{2,l}$, $l = \overline{1, n}$.

В результаті їх перетину отримаємо набір відомих прямокутних областей зображення (кадрів)

$$D = \Pi_{k,l} = [\alpha_k, \beta_k] \times [\gamma_l, \delta_l], k = \overline{1, m}, l = \overline{1, n}.$$

Невідома область зображення в цьому випадку являє собою систему вертикальних смуг

$$\overline{D}_{1,k,k+1}, k = \overline{1, m-1}$$

та горизонтальних смуг

$$\overline{D}_{2,l,l+1}, l = \overline{1, n-1}.$$

Ведемо до розгляду наступні оператори

$$\Theta_3 f(x, y) = \begin{cases} T_{1,k,k+1}(x, y) & (x, y) \in \overline{CD}_{1,2}; \\ T_{2,l,l+1}(x, y) & (x, y) \in \overline{CD}_{2,1}; \\ f(x, y) & (x, y) \in \Pi_{k,l}, k = \overline{1, m}, l = \overline{1, n}. \end{cases}$$

$$\Theta_4 f(x, y) = \begin{cases} \Theta_3(x, y) & (x, y) \notin \overline{UD}_{1,2,k,l}; \\ T_{1,2,k,l}[\Theta_3 f(x, y)] & (x, y) \in \overline{UD}_{1,2,k,l} \end{cases}$$

де

$$\overline{CD}_{1,2} = \overline{D}_{1,k,k+1} \setminus \overline{D}_{2,l,l+1}, k = \overline{1, m-1}, l = \overline{1, n-1}$$

$$\overline{CD}_{2,1} = \overline{D}_{2,l,l+1} \setminus \overline{D}_{1,k,k+1}, k = \overline{1, m-1}, l = \overline{1, n-1}$$

$$\overline{UD}_{1,2,k,l} = \overline{D}_{1,k,k+1} \cap \overline{D}_{2,l,l+1}, k = \overline{1, m-1}, l = \overline{1, n-1}$$

Відновлення зображення в цьому випадку відбувається в два етапи. Спершу до зображення застосовується оператор $\Theta_3 f(x, y)$, До отриманої поверхні $z = \Theta_3 f(x, y)$ застосовується оператор $\Theta_4 f(x, y)$.

Поверхня $z = \Theta_4 f(x, y)$ є наближеною математичною моделлю поверхні Σ , яка на кожному із кадрів

$$\Pi_{k,l} = [\alpha_k, \beta_k] \times [\gamma_l, \delta_l], k = \overline{1, m}, l = \overline{1, n}$$

точно відновлює поверхню, а на невідомих смугах $\overline{D}_{1,k,k+1}, k = \overline{1, m-1}$ та $\overline{D}_{2,l,l+1}, l = \overline{1, n-1}$ зображує поверхню за допомогою операторів $\Theta_3 f(x, y)$ та $T_{1,2,k,l}[\Theta_3 f(x, y)]$.

5. Обчислювальний експеримент

Для проведення обчислювальних експериментів було взято тестове зображення, з якого для кожного експерименту штучно були видалені смуги для їх подальшого відновлення викладеним вище методом інтерстріпації.

Результати цих експериментів наведені на рис. 1–4.

Висновки

В роботі наведені та проаналізовані такі класи алгоритмів: текстурні, шаблонні, основані на рівняннях з частинними похідними, гібридні та швидкі напіваавтоматичні. Окремо було розглянуто метод інтерстріпації функції двох змінних. На основі методу інтерстріпації у даній статті було запропоновано модифікований метод інтерстріпації для відновлення зображення поверхні за неповною інформацією про неї. Було проведено обчислювальний експеримент для випадків коли



а



б

Рис. 1. Результати проведення експерименту для інтерстріпації на серії вертикальних смуг (а – пошкоджене зображення, б – відновлене зображення)

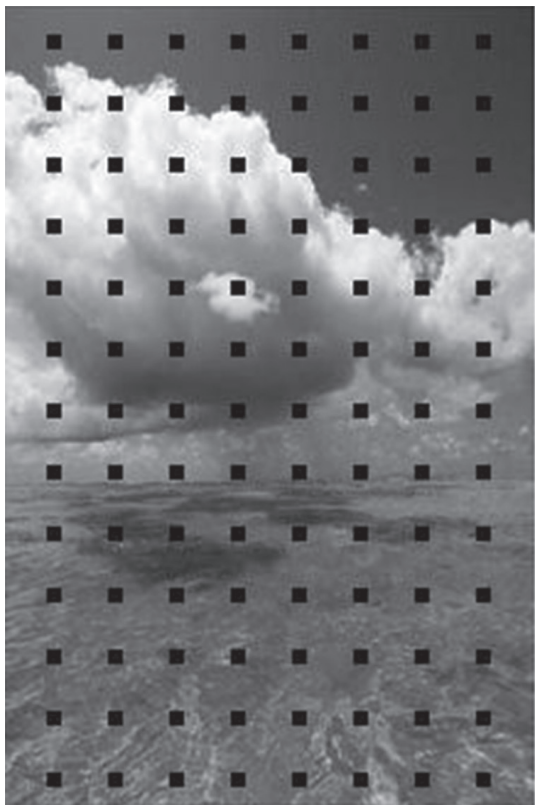


а



б

Рис. 2. Результати проведення експерименту для інтерстріпації на серії горизонтальних смуг
(*а* – пошкоджене зображення, *б* – відновлене зображення)



а



б

Рис. 3. Результати проведення експерименту для інтерстріпації на серії взаємоперпендикулярних смуг
(*а* – пошкоджене зображення, *б* – відновлене зображення)

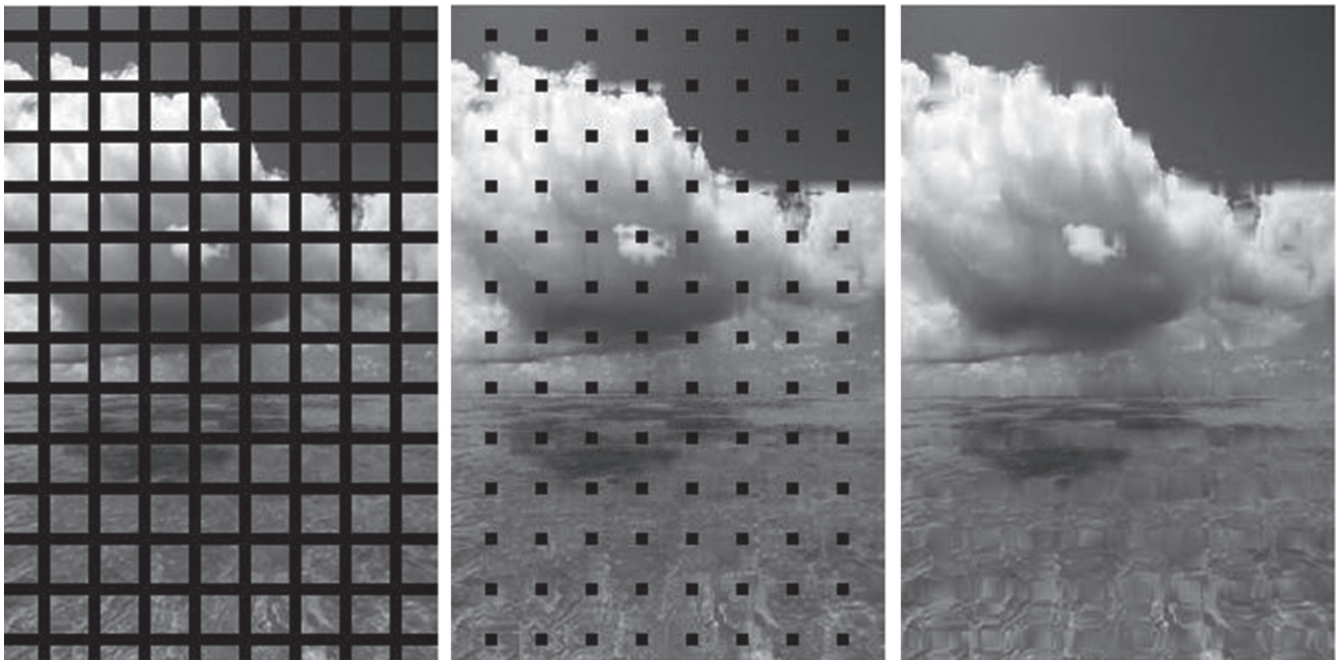


Рис. 4. Результати проведення експерименту для інтерстріпації на серії кадрів
(*а* – пошкоджене зображення, *б* – проміжний крок, *в* – відновлене зображення)

зображення відоме лише на системі горизонтальних, вертикальних, взаємоперпендикулярних смуг та на системі кадрів.

Наведений метод інтерстріпації (як і стандартний метод інтерстріпації) не використовує додаткову інформацію про структуру поверхні між смугами, тобто на невідомих ділянках зображення. В подальшому автори планують дослідити метод текстурної інтерстріпації, в якому буде враховуватися додаткова інформація про деякі властивості зображення між смугами.

Аналіз і розробка методів відновлення зображень є актуальною задачею для різноманітних прикладних областей та потребує подальших досліджень.

Список літератури:

1. Joshua J., Darsan G. Digital inpainting techniques – a survey // International journal of latest research in engineering and technology. – 2016. – vol. 2. – pp. 34–36. 2. Heeger D. J., Bergen J. R. Pyramid-based texture analysis/synthesis // In Proc. of ACM Conf. Comp. Graphics (SIGGRAPH). – 1995. – vol. 29. – pp. 229–233. 3. Yamauchi H., Haber J., Seidel H. Image restoration using multiresolution texture synthesis and image inpainting // In Computer Graphics International. – 2003. – pp. 120–125. 4. Criminisi A., Perez P., Toyama K. Region

filling and object removal by exemplar-based inpainting // IEEE Transactions on Image Processing. – 2004. – vol. 13. – pp. 1200–1212. 5. Drori I., Cohen-Or D., Yeshurun H. Fragment – based image completion // In Proceedings of ACM Conf. Comp. Graphics (SIGGRAPH). – 2003. – pp. 303–312. 6. Hays J., Efros A. Scene completion using millions of Graphics // Computer Graphics Proceedings (SIGGRAPH). – 2007. 7. Bertalmio M., Sapiro G., Caselles V., Ballester C. Image inpainting // Proc. of the 27th Annual Conf. on Computer Graphics and Interactive Techniques. – 2000. – pp. 417–424. 8. Tschumperl D., Deriche R. Vector-valued image regularization with PDE's: A common framework for different applications // IEEE Transactions on Pattern Analysis and Machine Intelligence. – 2005. – vol. 27. – pp. 506–517. 9. Sun J., Yuan L., Jian J., Shum H.-Y. Image completion with structure propagation // In Proc. of ACM Conf. Comp. Graphics. – 2005. – pp. 1–8. 10. Oliveira M., Bowen B., McKenna R., Chang Y.-S. Fast digital image inpainting // In Proc. of Intl. Conf. on Visualization, Imaging and Image Processing. – 2001. – pp. 261–266. 11. Матвеева С. Ю. Математичне моделювання поверхні тіла методами інтерстріпації функцій за даними радіолокації : Дис. канд. фіз.-мат. наук / Матвеева Світлана Юріївна. – Харків, 2014. – 126 с. 12. Литвин О. М. Інтерлінація функцій. – Харків: Основа, 1992. – 234 с. 13. Литвин О. М. Інтерлінація функцій та деякі її застосування. – Харків: Основа, 2002. – 544 с.

Надійшла до редакції 09.11.2016.

УДК 681.3.07

**В.В. Ляшенко¹, О. А. Кобылин², С.Н. Томич³**¹ ХНУРЭ, м. Харьков, Украина, lyashenko.vyacheslav@mail.ru² ХНУРЭ, м. Харьков, Украина, oleg.kobylin@gmail.com³ ХНУРЭ, м. Харьков, Украина, stas.tomich@gmail.com

ОСНОВНЫЕ ЭТАПЫ ОБРАБОТКИ ИЗОБРАЖЕНИЙ ЦИТОЛОГИЧЕСКИХ ПРЕПАРАТОВ С ИСПОЛЬЗОВАНИЕМ ИДЕОЛОГИИ ВЕЙВЛЕТОВ

В данной работе рассмотрена целесообразность применения идеологии вейвлетов в медицине, а именно для обработки изображений цитологических препаратов. Показаны особенности обработки изображений цитологических препаратов с использованием идеологии вейвлетов. Приведены примеры обработки различных изображений цитологических препаратов.

ВЕЙВЛЕТ АНАЛИЗ, КОНТРАСТИРОВАНИЕ, ОБРАБОТКА ИЗОБРАЖЕНИЙ, ЦИТОЛОГИЧЕСКИЕ ПРЕПАРАТЫ

Введение

Фиксация и анализ изображений различных объектов является одним из направлений познания окружающего нас мира. Такой подход дает возможность исследовать не только изменения, происходящие в реальном мире, но и обнаружить и изучить возможные закономерности таких изменений. Более того, всесторонний и комплексный анализ изображений реальных объектов позволяет исследовать процессы, которые невозможно увидеть при помощи простого человеческого зрения.

Одними из таких направлений применения общей идеологии обработки изображений (анализа и интерпретации данных) является медицина. В качестве образов реальных объектов выступают изображения различных органов, тканей, отдельных частей скелета человека, которые получены при помощи специальных методов их визуализации: магнитно-резонансной томографии [1], позитронно-эмиссионной томографии [2], ультразвукового анализа [3], световой и электронной микроскопии [4, 5].

Среди множества образов реальных объектов, позволяющих исследовать организм человека, особо можно выделить изображения цитологических препаратов. Это связано с тем, что:

с одной стороны, цитологические препараты представляют собой объекты микромира, которые позволяют проводить более углубленные исследования человеческого организма, изучать динамику его функционирования и осуществлять диагностику возможных заболеваний на ранних стадиях их развития;

с другой стороны, это специальные изображения, которые отличаются особенностью визуализации объектов микромира, что обуславливает необходимость использования разнообразных методов обработки изображений с целью получения информации об исследуемых объектах, процессах, явлениях.

В общем случае, идеология обработки изображений цитологических препаратов преследует своей целью идентификацию определенных частей изображения (клетки, ядра клетки) для последующего изучения их изменений (изменение формы клетки, изменение площади клетки) либо подсчета определенных количественных характеристик (количества клеток, количества ядер клеток, площади клеток). Для решения поставленных задач могут быть использованы как методы выделения контуров, так и методы сегментации. В тоже время решение таких задач предполагает использование предварительных методов обработки изображений цитологических препаратов, которые направлены на улучшение качества восприятия и последующей обработки исходного изображения [6, 7]. Это связано с тем, что процедура получения изображений цитологических препаратов использует методику окрашивания рассматриваемых клинических образцов, которая, в конечном счете, является одним из ключевых источников возникновения ошибок при обработке таких изображений, вследствие возникающих различий относительной интенсивности окрашивания отдельных частей цитологического препарата. Однако следует отметить, что простым изменением яркости, контрастности либо одной фильтрацией невозможно качественно решить возникающие задачи в обработке изображений цитологических препаратов. В тоже время практика автоматической обработки изображений цитологических препаратов, как правило, основывается на монохромных изображениях, что вносит свои коррективы в восприятие и анализ соответствующих образов на цитологических препаратах [8, 9]. При этом следует отметить, что существуют различные методы и подходы к обработке и анализу изображений цитологических препаратов, каждый из которых имеет как положительные, так и отрицательные моменты. Поэтому разработка и

реализация новых процедур анализа изображений цитологических препаратов является перспективным направлением исследования, которое будет способствовать нахождению наилучших решений в реализации целей обработки медицинских изображений.

Таким образом, представляется целесообразным рассмотреть процедуру обработки изображений цитологических препаратов с использованием идеологии вейвлетов, которая хорошо себя зарекомендовала в различных областях обработки изображений.

1. Основы идеологии вейвлетов в процедуре обработки изображений

Выбор вейвлет анализа для обработки изображений цитологических препаратов обоснован тем, что вейвлет обработка позволяет учесть особенности рассматриваемых изображений за счет разложения исходных данных на множество аппроксимирующих и детализирующих коэффициентов, в частности при выделении контура, что можно рассматривать в качестве основы исследования таких изображений. Более того, вейвлет анализ изображений позволяет получить дополнительную информацию об исследуемых объектах, что позволяет, в частности, повысить качество анализа изображений цитологических препаратов.

Вейвлет анализ основан на вейвлет преобразовании. Вейвлет-преобразование — это разложение сигнала (в частности некоторого изображения) по системе вейвлетов. Вейвлеты получаются путем сдвига и масштабированием одной функции — порождающего вейвлета [10]. Вейвлетом в таком случае является функция, быстро убывающая на бесконечности, среднее значение которой равно нулю. Если сигнал имеет разрыв, то высокие амплитуды будут только у тех вейвлетов, максимумы которых окажутся вблизи точки разрыва. Это позволяет выделять контур на исследуемом изображении. В тоже время в общем виде под разрывами понимается резкий скачкообразный переход в течение какого-либо процесса. Количественно это можно оценить величиной первой производной такого процесса. Там, где имеют место скачки, первая производная очень велика. Если скачек имеет форму разрыва, то первая производная стремится к бесконечности. Однако, реальные процессы, измеренные физически реальными приборами, не могут иметь идеальных разрывов. В действительности, измеренные фрактальные переходы характеризуются конечным значением производной. Чем резче разрыв, тем больше значение производной. Плавные переходы будут иметь небольшие значения производной. Благодаря этому можно

определить наличие особенностей анализируемого изображения, а также и точку в котором возможные особенности проявляется. Подчеркнуть такие особенности помогает многоуровневое разложение исходного изображения на множество аппроксимирующих и детализирующих коэффициентов.

В основе формализации непрерывного вейвлет-преобразования (НВП) лежит использование двух непрерывных и интегрируемых по всей оси t функций [10, 11]:

— вейвлет — функции $\phi(t)$ с нулевым значением интеграла

$$\int_{-\infty}^{\infty} \phi(t) dt = 0, \quad (1)$$

определяющей детали сигнала и порождающей детализирующие коэффициенты;

— масштабирующая функции $\phi(t)$ с единичным значением интеграла

$$\int_{-\infty}^{\infty} \phi(t) dt = 1, \quad (2)$$

определяющей грубое приближение сигнала и порождающей коэффициенты аппроксимации.

Однако функция НВП применима только для одномерных сигналов, а изображение является двумерным сигналом. Поэтому для возможности применения НВП с целью выделения границ изображения предлагается рассматривать следующую процедуру анализа и выделения контура [12]:

— выполним вычисление горизонтальных разрывов исходного изображения F , представленного в виде матрицы, заданной своими отсчетами $f_{ij} \in \{0, 1, \dots, P\}$, $i = 1, 2, \dots, N$, $j = 1, 2, \dots, M$ на квадратной решетке $N \times M$. Для этого используем следующее выражение для получения так называемой матрицы вейвлет-спектрограммы W (исходя из последовательной обработке каждой строки исходного изображения F):

$$W[f_{ij}] = \frac{1}{\sqrt{a}} \int_{-\infty}^{+\infty} f_{ij} \phi\left(\frac{t-b}{a}\right) dt, \quad (3)$$

где $\phi\left(\frac{t-b}{a}\right)$ — материнский вейвлет, удовлетворяющий условию (1); a , b — масштаб и центр временной локализации, которые определяют масштаб и смещение функции $\phi(t)$ в соответствии с условиями масштабирования (2); $[f_{ij}]$ — обозначает номер обрабатываемой строки исходного изображения F для получения множества значений ее вейвлет-спектрограммы W .

Параметры a , b выбираются таким образом, чтобы соответствующие линейные размеры матрицы вейвлет-спектрограммы W коррелировали с линейными размерами исходного изображения F , но при этом были учтены возможные параметры вейвлет преобразования.

Далее на основе анализа полученной спектрограммы (W , для каждой строки исходного изображения F) выбираем определенную ее строку NN , исходя из условия

$$NN = \max\left(\frac{1}{N} \sum_{i=1}^N w_{ij}\right), \quad (4)$$

где w_{ij} — элемент вейвлет-спектрограммы анализируемой строки исходного изображения F .

Такой выбор обусловлен тем, что мы выбираем ту часть спектра строки исходного изображения, которая соответствует наибольшему разрыву исходного сигнала между его отсчетами (смотри замечания сделанные выше).

Выбранная таким образом строка будет соответствовать строке в матрице F_g , которая характеризует матрицу горизонтальных разрывов исходного изображения F .

Обработка всех строк исходного изображения F позволяет в итоге получить матрицу горизонтальных разрывов F_g , благодаря следующей последовательности преобразований:

$$F \xrightarrow{\text{НВП строк } W} \text{выбор строки } F_g.$$

— аналогичным образом производится вычисление вертикальных разрывов исходного изображения F для каждого его столбца. Для этого используется формула (3) и аналогичная формуле (4) формула для выбора определенной строки из полученных вейвлет-спектрограмм для каждого столбца исходного изображения F :

$$MM = \max\left(\frac{1}{M} \sum_{i=1}^M w_{ij}\right). \quad (5)$$

Обработка всех столбцов исходного изображения F позволяет в итоге получить матрицу вертикальных разрывов F_v , благодаря следующей последовательности преобразований:

$$F \xrightarrow{\text{НВП столбца } W} \text{выбор столбца } F_v.$$

— производится сложение матриц вертикальных и горизонтальных разрывов в одну матрицу, которая и отображает контур исходного изображения на основе метода НВП. Для визуальной наглядности матрицы горизонтальных, вертикальных разрывов, а также обобщенная матрица, отображающая контур исходного изображения могут быть инвертированы.

При этом следует отметить, что построение вейвлет-спектрограммы W во многом определяется размерами исходного изображения и используемым параметром масштаба a при проведении вейвлет преобразования исследуемого изображения.

2. Контрастирование как элемент анализа изображений цитологических препаратов

Как правило, в качестве предварительной процедуры обработки микроскопических изображений в медицине, к которым относятся и изображения цитологических препаратов, выделяют контрастирования [13, 14].

Контрастность — одна из основных характеристик изображения, напрямую связанная с яркостью пикселей, которые являются источниками информации об объектах на изображении. Поэтому изменение контраста изображения позволяет повысить как четкость восприятия изображения, так и точность (эффективность) его последующей обработки. В частности, при увеличении контрастности изображения светлые участки становятся еще светлее, а темные темнее. В результате происходит перераспределение пикселей за счет среднего тонового диапазона. При уменьшении контрастности изображения, наоборот происходит расширение среднего тонового диапазона. Темные пиксели становятся более светлыми, а светлые более темными и частично переходят в средние тона. Таким образом, изменение контраста изображения, и прежде всего увеличение контраста позволяет сделать отдельные детали изображения более различимыми. Это очень важно для микроскопических изображений в медицины. Для контрастирования изображений цитологических препаратов могут быть использованы различные методы, среди которых можно выделить [15]:

- выравнивание гистограммы значений яркости элементов изображения (эквализация),
- нелинейное растяжение динамического диапазона значений яркостей изображения,
- использование различных масок фильтрации,
- нечеткое маскирование.

В дальнейшем мы используем метод выравнивания гистограммы значений яркости элементов изображения, как один из самых простых методов контрастирования изображений. Нашей целью, в первую очередь, является показать влияние контрастирования на эффективность применения идеологии вейвлетов в обработке изображений цитологических препаратов.

3. Тестовые изображения

Для проведения экспериментов мы используем различные изображения цитологических препаратов, которые находятся в открытом доступе интернет.

На рис. 1 представлено изображение цитологического препарата — молочной железы, где показано скопление клеток эпителиальных групп разного типа (с сайта <http://screening.iarc.fr/>). В частности,

желтым цветом выделены опухолевые клетки кожи, расположенные в эпидермисе.

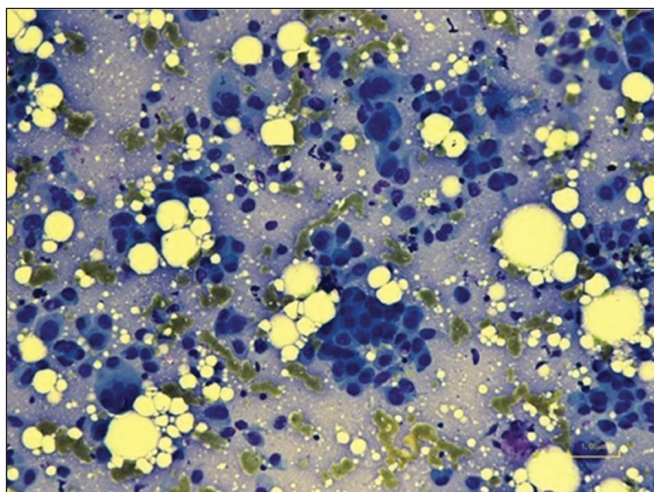


Рис. 1. Цитологический препарат. Молочная железа

На рис. 2 представлено изображение цитологического препарата для диагностики рака щитовидной железы, где синим цветом показаны потенциальные клетки рака (с сайта <http://www.pathologyoutlines.com>).

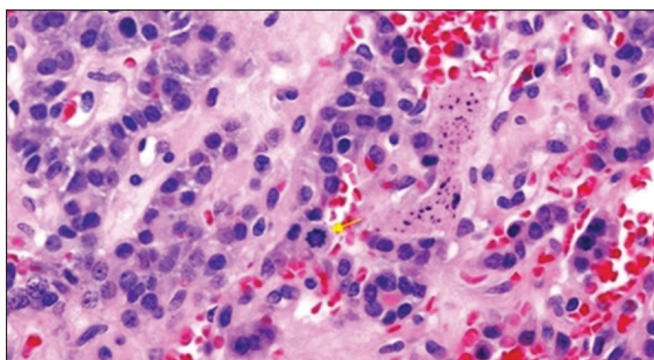


Рис. 2. Цитологический препарат. Щитовидная железа

Представленные изображения цитологических препаратов являются различными по структуре и сложности восприятия, что позволяет оценить возможности применения методологии вейвлет анализа в качестве инструмента их обработки.

4. Эксперименты и обсуждение

Для реализации идеологии вейвлетов с целью обработки и анализа изображений цитологических препаратов мы преобразовываем исходные изображения в полутоновые, а затем контрастируем полутоновые изображения. На рис. 3 представлены результаты преобразования исходного изображения рис. 1 в полутоновое изображение. Результаты контрастирования такого изображения в виде соответствующих гистограмм для исходного полутонового изображения и контрастированного полутонового изображения представлены на рис. 4 (рис. 4, а, рис. 4, б соответственно).

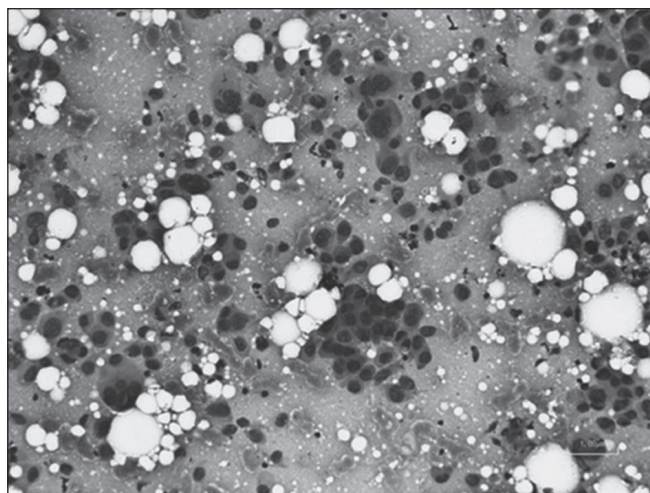
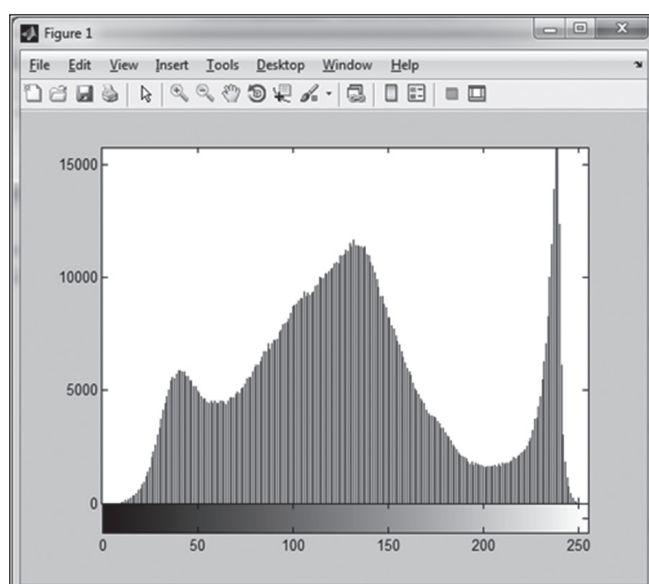
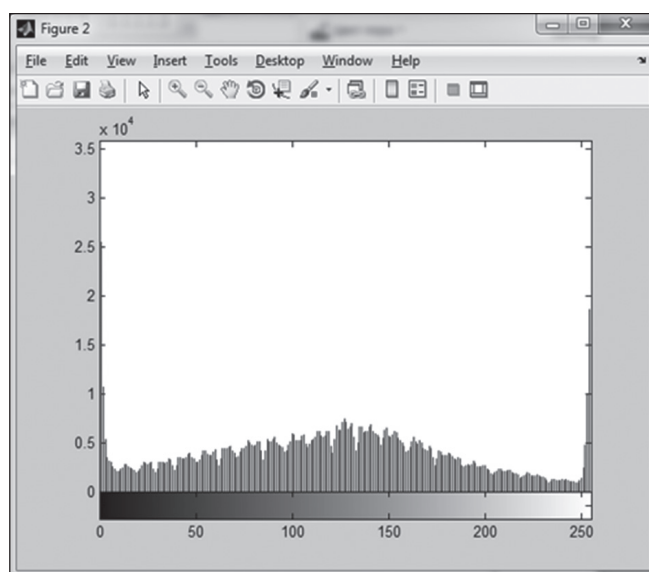


Рис. 3. Полутоновое изображение рис. 1



а



б

Рис. 4. Гистограммы исходного и контрастированного полутонового изображения для рис. 1

На рис. 5 представлены результаты преобразования исходного изображения рис. 2 в полутоновое изображение. Результаты контрастирования такого изображения в виде соответствующих гистограмм для исходного полутонового изображения и контрастированного полутонового изображения представлены на рис. 6 (рис. 6, *а*, рис. 6, *б* соответственно).

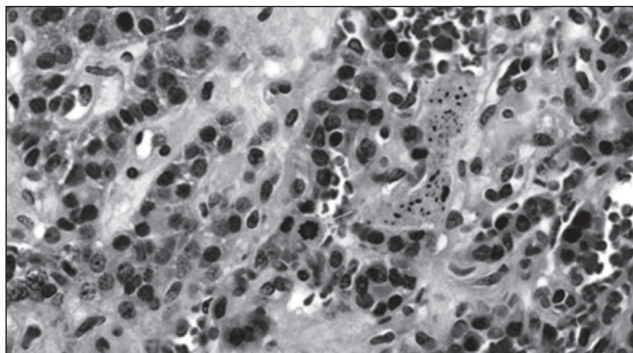
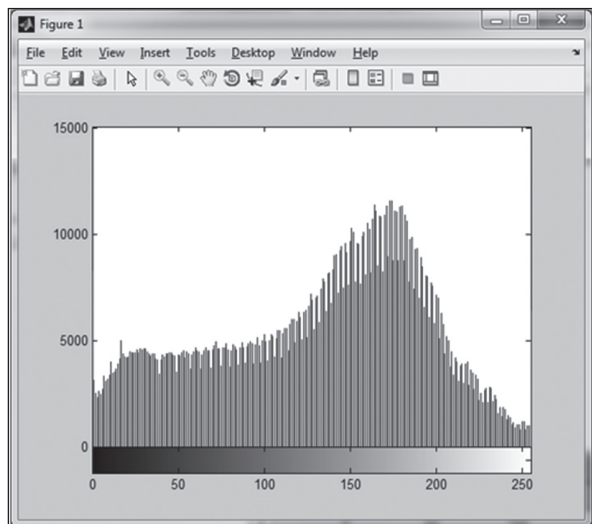
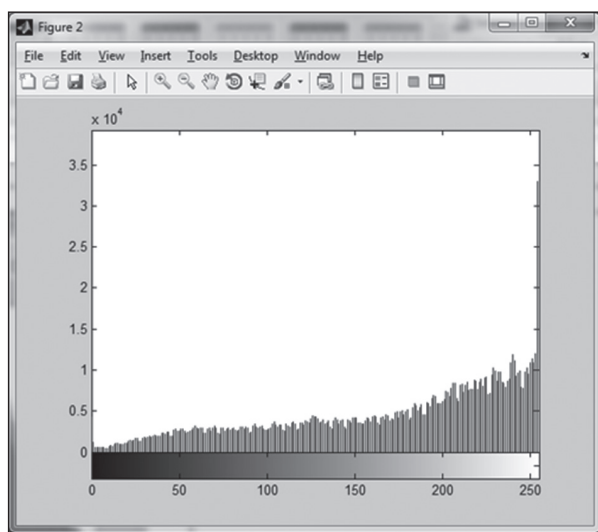


Рис. 5. Полутоновое изображение рис. 2



а

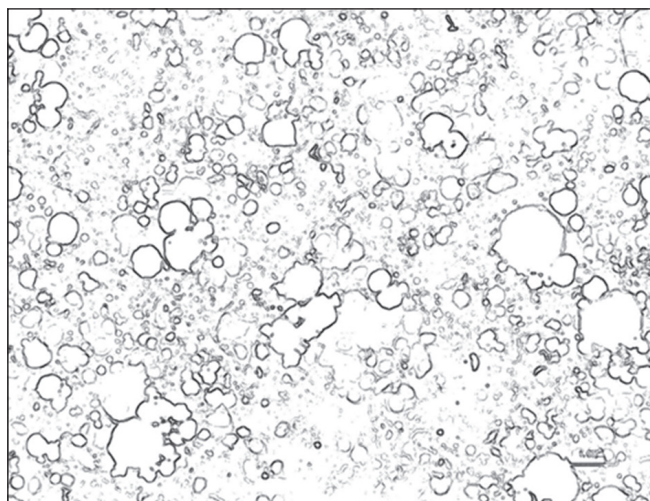


б

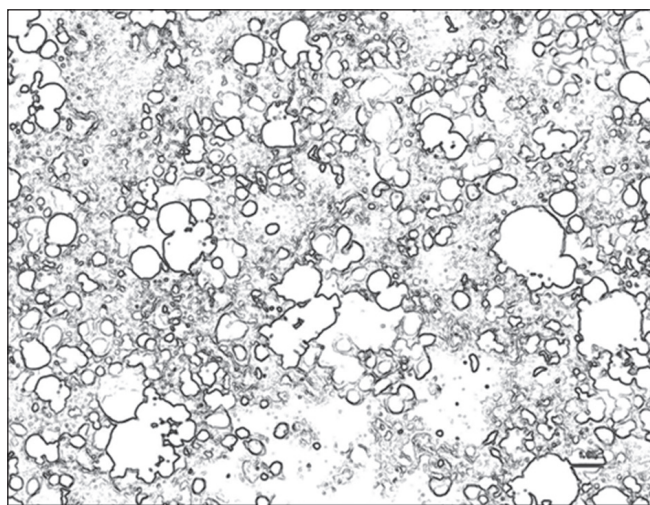
Рис. 6. Гистограммы исходного и контрастированного полутонового изображения для рис. 2

Далее мы применяем идеологию вейвлетов к исходному и контрастированному полутоновым изображениям для анализа изображений цитологических препаратов.

На рис. 7 представлены результаты вейвлет преобразования для изображения рис. 1 (*а* – обработка исходного полутонового изображения, *б* – обработка контрастированного полутонового изображения). Для вейвлет преобразования мы используем вейвлет db1.



а



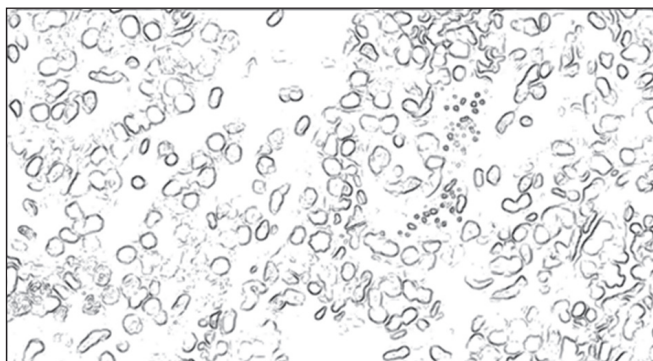
б

Рис. 7. Результаты вейвлет преобразования для изображения 1

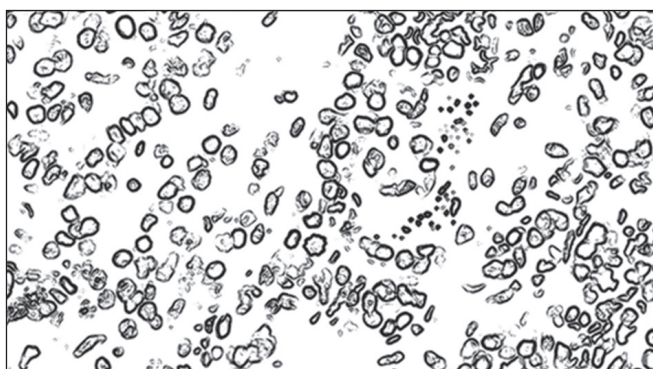
На рис. 8 представлены результаты вейвлет преобразования для изображения рис. 2 (*а* – обработка исходного полутонового изображения, *б* – обработка контрастированного полутонового изображения).

Как видно из рис. 7 и рис. 8 процедура предварительного контрастирования исходного изображения позволяет получить более информативные изображения после применения идеологии вейвлетов к изображениям цитологических препаратов.

Прежде всего, используемое вейвлет преобразование позволяет подчеркнуть присущие особенности, представленных на изображениях цитологических препаратов различных объектов, подчеркнуть границы выделения таких объектов. Таким образом, в дальнейшем можно говорить об организации процедуры кластеризации объектов на изображениях цитологических препаратов с использованием идеологии вейвлетов.



а



б

Рис. 8. Результаты вейвлет преобразования для изображения 2

Следует также отметить о появлении множества фоновых точек в результате предварительного контрастирования исходного изображения и последующего применения к таким изображениям процедуры вейвлет преобразования. Для возможной минимизации излишнего количества фоновых точек в результате предварительного контрастирования исходного изображения цитологических препаратов представляется целесообразным выбор метода контрастирования в зависимости от общей сложности и насыщенности объектами исходного изображения цитологических препаратов. В частности для изображения на рис. 9, используемый метод контрастирования — выравнивание гистограммы значений яркости элементов изображения, дает хорошие результаты (рис. 10, а — обработка исходного полутонового изображения, б — обработка контрастированного полутонового изображения).

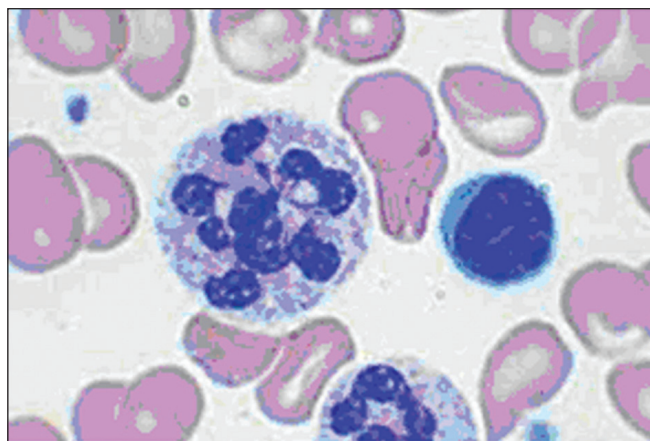
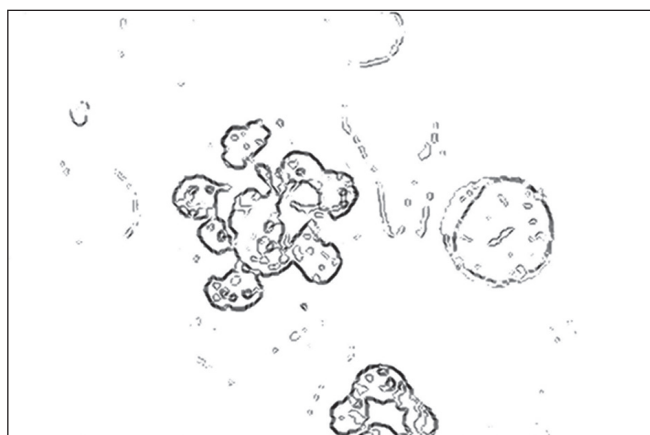
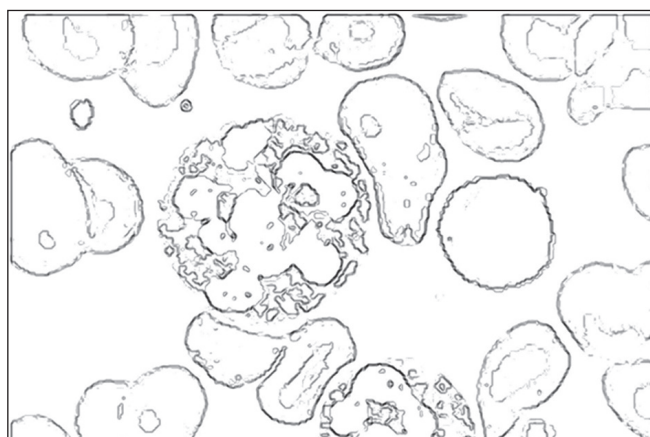


Рис. 9. Цитологический препарат. Клетки крови



а



б

Рис. 10. Результаты вейвлет преобразования для изображения 9

Другим направлением по минимизации множества фоновых точек в результате предварительного контрастирования исходного изображения можно рассматривать процедуру фильтрации, которая направлена на сглаживание незначительных перепадов и удаление небольших объектов на изображении цитологических препаратов. Основой для проведения такой фильтрации является как выбор фильтра определенного типа, так и выбор размера окна для реализации процедуры фильтрации

(определение масштаба используемого фильтра). Определение масштаба используемого фильтра должно коррелировать с размерами информативных объектов на изображении цитологических препаратов. Также важно учитывать и длину волны вейвлета, который используется для преобразования изображений цитологических препаратов. При этом такая длина волны используемого вейвлета и масштаб фильтра для сглаживания незначительных перепадов и удаление небольших объектов на изображении цитологических препаратов в идеале должны коррелировать между собой.

Для повышения эффективности использования идеологии вейвлетов с целью обработки изображений цитологических препаратов также можно использовать полученные результаты для реализации процедуры нечеткого маскирования. Этот вывод основан на сравнении полученных результатов для исходного и контрастированного изображения в соответствии с рис. 7 и рис. 8 и возможностью вычисления существующего несоответствия между результатами вейвлет обработки исходного и контрастированного изображения, где такое несоответствие может быть использовано в качестве элемента усиления в процедуре нечеткого маскирования.

Тем не менее, приведенные результаты показывают возможность и целесообразность применения идеологии вейвлетов для обработки изображений цитологических препаратов. При этом среди основных этапов такой обработки следует выделить: контрастирование, подавление локальных шумов и непосредственное применение вейвлет преобразования.

Выводы

Рассмотрена общая процедура обработки изображений цитологических препаратов на основе использования идеологии вейвлетов. Показаны положительные и отрицательные стороны применения контрастирования изображений цитологических препаратов с целью их дальнейшей обработки при помощи вейвлет преобразования. Проанализированы варианты и пути минимизации возникновения множества фоновых точек в результате предварительного контрастирования исходного изображения для дальнейшей его обработки на основе идеологии вейвлетов. Приведены примеры обработки различных изображений цитологических препаратов с использованием вейвлет преобразования.

Список литературы:

1. Schluter, S. et al. Image processing of multiphase images obtained via X-ray microtomography: a review [Text] S. Schluter et al. //Water Resources Research. — 2014. — Т. 50. — №. 4. — С. 3615–3639.
2. Gaemperli, O. et al. Imaging intraplaque inflammation in carotid atherosclerosis with 11C-PK11195 positron emission tomography/computed tomography [Text] O. Gaemperli et al. //European heart journal. — 2012. — Т. 33. — №. 15. — С. 1902–1910.
3. Sikdar, S. et al. Novel method for predicting dexterous individual finger movements by imaging muscle activity using a wearable ultrasonic system [Text] S. Sikdar et al. //IEEE Transactions on Neural Systems and Rehabilitation Engineering. — 2014. — Т. 22. — №. 1. — С. 69–76.
4. Eklund, A. et al. Medical image processing on the GPU—Past, present and future [Text] A. Eklund et al. //Medical image analysis. — 2013. — Т. 17. — №. 8. — С. 1073–1094.
5. Ciresan, D. et al. Deep neural networks segment neuronal membranes in electron microscopy images [Text] D. Ciresan et al. //Advances in neural information processing systems. — 2012. — С. 2843–2851.
6. Mahendran, G. Automatic segmentation and classification of pap smear cells [Text] G. Mahendran, R. Babu, D. Sivakumar //International Journal of Management, IT and Engineering. — 2014. — Т. 4. — №. 5. — С. 100–108.
7. Singh, S. Identification of components of fibroadenoma in cytology preparations using texture analysis: a morphometric study [Text] S. Singh, R. Gupta //Cytopathology. — 2012. — Т. 23. — №. 3. — С. 187–191.
8. Lyashenko, V. V. Using the methodology of wavelet analysis for processing images of cytology preparations [Text] V. V. Lyashenko, A. M. A. abd allah Babker, O. A. Kobylin //National Journal of Medical Research. — 2016. — Т. 6. — №. 1. — С. 98–102.
9. Lyashenko, V. V. The methodology of wavelet analysis as a tool for cytology preparations image processing [Text] V. V. Lyashenko, A. M. A. A. Babker, O. A. Kobylin //Cukurova Medical Journal. — 2016. — Т. 41. — №. 3. — С. 453–463.
10. Kingsbury, N. Image processing with complex wavelets [Text] N. Kingsbury //Philosophical Transactions of the Royal Society of London A: Mathematical, Physical and Engineering Sciences. — 1999. — Т. 357. — №. 1760. — С. 2543–2560.
11. Heil, C. E. Continuous and discrete wavelet transforms [Text] C. E. Heil, D. F. Walnut //SIAM review. — 1989. — Т. 31. — №. 4. — С. 628–666.
12. Kobylin, O. Comparison of standard image edge detection techniques and of method based on wavelet transform [Text] O. Kobylin, V. Lyashenko //International Journal of Advanced Research. — 2014. — Т. 2. — №. 8. — С. 572–580.
13. Dey, N. et al. Digital analysis of microscopic images in medicine [Text] N. Dey et al. //Journal of Advanced Microscopy Research. — 2015. — Т. 10. — №. 1. — С. 1–13.
14. Lyashenko, V. Contrast Modification as a Tool to Study the Structure of Blood Components [Text] V. Lyashenko, R. Matarneh, O. Kobylin //Journal of Environmental Science, Computer Science and Engineering & Technology. — 2016. — Т. 5. — №. 3. — С. 150–160.
15. Semmlow J. L., Griffel B. Biosignal and medical image processing. [Text] / J. L. Semmlow, B. Griffel — CRC press, 2014.

Поступила в редколлегию 10.11.2016

УДК 004.891.3



О.В. Чалая

ХНУРЭ, г. Харьков, Украина, oksana.chala@nure.ua

МЕТОД ОБОБЩЕНИЯ ПРЕДСТАВЛЕНИЯ ЗНАНИЕ-ЕМКОГО БИЗНЕС-ПРОЦЕССА

Предложена структурная модель знание-емкого процесса на отдельном уровне детализации. Модель обеспечивает возможность построения обобщенного/детализованного представления такого процесса на основе априорно заданных характеристик элементов контекста, что дает возможность согласовать иерархию элементов контекста и иерархию представления процесса. Предложен метод обобщения представления знание-емкого бизнес-процесса на основе анализа его логов. Метод включает в себя этапы выделения представленных в логе подмножеств атрибутов объектов для каждого уровня детализации и проведения обобщения действий процесса с использованием выделенных атрибутов. При проведении обобщения объединяются цепочки последовательных действий, блоки выбора/объединения результатов выбора, а также скрываются частные, редко используемые реализации процесса, что дает возможность интегрировать функциональный и процессный подходы к управлению в рамках одного предприятия. Метод позволяет повысить эффективность интеллектуального анализа процессов путем выделения существенных для анализа и принятия решений подмножеств действий, что создает условия для решения проблемы «спагетти-подобных» workflow –моделей.

ЗНАНИЕ-ЕМКИЙ БИЗНЕС-ПРОЦЕСС, КОНТЕКСТ, ОБЪЕКТ, НЕЯВНЫЕ ЗНАНИЯ, АРТЕФАКТ, WORKFLOW

Введение

Реализация процессного подхода к управлению предполагает построение модели бизнес-процесса (БП), а также последующее ее конфигурирование в процессной информационно-управляющей системе. Традиционно модель представляется в виде workflow - графа, описывающего последовательности действий бизнес-процесса, а также взаимосвязи между этими последовательностями [1,2].

Однако для знание-емких бизнес-процессов (ЗБП) построение такой модели связано со значительными трудностями в силу того, что ход выполнения процесса может быть изменен исполнителями на основе знаний. Для принятия решений по изменению хода процесса могут использоваться как явные, так и неявные знания. Явные знания обычно представлены в документальной форме. Неявные знания отражают опыт исполнителя и очень плохо переводятся в вербальную форму. На основе неявных знаний могут быть выстроены частные варианты реализации ЗБП. Однако причины, по которым были выполнены частные последовательности действий, обычно неизвестны в силу неявного характера знаний. Модель ЗБП, содержащая все известные сценарии выполнения процесса (включая частные случаи) становится многовариантной, характеризуется значительным количеством ветвлений, что затрудняет анализ и практическое использование. Основная причина указанных трудностей состоит в том, что в модели на одном уровне детализации приводятся как ключевые последовательности действий, так и редко используемые частные случаи.

Поэтому проблема многоуровневой детализации знание-емкого бизнес-процесса несомненно является актуальной.

В настоящее время при моделировании бизнес-процессов используются два основных подхода: процедурный (workflow)[1-3] и подход к управлению последовательностью действий процесса на основе данных [4-7].

Первый подход моделирует процесс через действия и управляющие структуры посредством изменения состояния его исполняемого экземпляра. Изменение его состояния происходит только после выполнения действия БП. Недостаток данного подхода состоит в том, что все действия приводятся на одном уровне детализации, без учета их важности для достижений целей организации, а также уровня исполнителей этих действий в организационной иерархии.

Попытка построения детализованного/обобщенного представления бизнес-процесса была предпринята в работах [8,9]. Однако при детализации рассматривались статистические критерии (вероятность достижения заданного состояния, частота появления заданного подмножества состояний и т.п.). Семантика детализации, влияние детализации на принятие решений, а также учет условий среды, в которых выполнялись действия процесса, в этих работах не рассматривались.

Второй подход к моделированию бизнес-процессов интенсивно используется фирмой IBM. Он основан на выделении артефактов – объектов, с которыми взаимодействует бизнес-процесс. В соответствии с данным подходом создается информационная модель артефакта и модель его жизненного цикла в рамках рассматриваемого бизнес-процесса [10]. Первая модель включает в себя набор всех используемых бизнес-процессом атрибутов, а вторая содержит набор

последовательностей действий по обработке артефакта, реализующий все варианты его обработки артефакта [11].

Основанный на артефактах подход позволяет представить бизнес-процесс с различной степенью детализации. Для этого необходимо скрыть подмножество артефактов и связанных с ними действий, которые не влияют на принятие решений на заданном уровне иерархии. Однако данный подход ориентирован на моделирование синхронизации состояний артефактов без явного задания последовательности действий по их использованию. Указанный недостаток не позволяет обобщить/детализовать общую схему действий бизнес-процесса.

Таким образом, вопросы детализации/обобщения представления знание-емкого бизнес-процесса с учетом контекста его выполнения требуют дальнейшего рассмотрения.

1. Постановка задачи

Целью статьи является разработка метода обобщения/детализации представления знание-емкого бизнес-процесса, который бы обеспечивал возможность его анализа и корректировки с заданной степенью грануляции, обеспечивающей эффективное решение задач процессного управления.

2. Метод обобщения представления бизнес-процесса

При реализации управления бизнес-процессами предприятие должно иметь несколько детализованных/обобщенных моделей каждого таких процессов. Очевидно, что руководство организации должно использовать максимально обобщенное описание ЗБП, обеспечивающее принятие общих решений по управлению предприятием. В то же время непосредственные исполнители процесса должны иметь полную схему действий процесса на участке их ответственности, обеспечивающую решение их текущих задач.

Поэтому уровень детализации ЗБП в общем случае зависит от следующих факторов:

- уровня исполнителей в иерархии организации, в которой выполняется процесс;
- уровня решаемых этими исполнителями задач.

При обобщении модели бизнес-процесса часть информации о его действиях скрывается (либо удаляется). При этом необходимо сохранить существенные для каждого уровня информацию и знания, которые необходимы для принятия решений исполнителями.

Предлагаемый метод ориентирован на использование в задачах интеллектуального анализа процессов [12]. Интеллектуальный анализ процессов

предназначен для построения их workflow — моделей на основе анализа логов. Лог содержит набор последовательностей событий, каждая из которых отражает однократное выполнение бизнес-процесса.

Отметим, что при решении задач интеллектуального анализа ЗБП методы process mining формируют спагетти подобные модели [12], использование которых связано со значительными трудностями.

События лога характеризуются меткой времени, а также набором значений атрибутов артефактов, с которым взаимодействует бизнес-процесс. Например, с событием могут быть связаны следующие атрибуты:

- наименование обрабатываемого процессом объекта;
- страна, в которой выполняется бизнес-процесс;
- наименование организации, в которой выполняется бизнес-процесс;
- наименование подразделения организации, в котором выполнялись соответствующие действия бизнес-процесса;
- имя исполнителя;
- роль исполнителя;
- наименование действия процесса, информация о котором сохранена в событии лога;
- текущее состояние действия процесса (в частности, ожидание ресурсов, выполнение, ожидание запроса пользователя).

Из приведенного примера видно, что лог может содержать информацию о иерархической структуре организации в формате:

```

страна->
организация->
подразделение->
исполнитель.

```

Поиск такой информации может быть выполнен аналитиком, либо с использованием методов data mining. Однако в целом можно утверждать, что необходимая для детализации иерархия содержится в логе в форме значений атрибутов событий.

Пусть структурная модель знание-емкого бизнес-процесса задана в виде набора элементов workflow, контекста, а также связывающих их зависимостей, определяющих возможность выполнения действий процесса в зависимости от контекста:

$$BP = (Ct, Kn, Wf), \quad (1)$$

где Ct — контекст бизнес-процесса; Kn — составляющая знаний ЗБП; Wf — workflow — описание ЗБП.

Контекст, в свою очередь, состоит из набора артефактов, между которыми существуют контекстные зависимости:

$$Ct = (Af, R_{Ct}), Af = \{af_i\}, R_{Ct} : af_i \rightarrow af_k \mid \quad (2)$$

$$af_i, af_k \in Af,$$

где Ct — контекст бизнес-процесса; Af — множество артефактов, составляющих контекст; R_{Ct} — контекстные зависимости между артефактами бизнес-процесса, af_i, af_k — артефакты контекста бизнес-процесса.

Тогда структурная модель знание-емкого процесса на текущем уровне детализации имеет вид:

$$BP^* = (Ct^* = \{(a_{si}, v_{si}^j) \mid \quad (3)$$

$$\forall (a, v) \in Ct^* \exists e \in \Pi : e(a) = v\},$$

$$Kn^* \subseteq Kn \mid \forall r \in Kn^* \nexists (a, v) \in Ct^*,$$

$$Wf^* \subseteq Wf \mid \forall wf \in Wf^* \nexists r \in Kn^*$$

где Ct^* — выделенное для текущего уровня подмножество свойств артефактов контекста, которое определяется через набор имеющихся в логе атрибутов событий и значений этих атрибутов; e — событие лог процесса; Π — лог процесса a_{si} — s атрибут i — артефакта; v_{si}^j — значение атрибута a_{si} ; Kn^* — подмножество знаний для текущего уровня детализации; Wf^* — подмножество возможных последовательностей работ для данного уровня детализации.

Из выражения (3) видно, что для построения требуемого уровня детализации представления знание-емкого бизнес-процесса могут быть отобраны только те атрибуты a_{si} и их значения v_{si}^j , которые записаны в логе.

Уровень знаний Kn^* может содержать только такие зависимости r , которые используют выделенные атрибуты артефактов контекста Ct^* .

В свою очередь, любая возможная последовательность работ $wf \in Wf^*$ на заданном уровне детализации задается только теми правилами, которые используют выделенное подмножество атрибутов артефактов контекста Ct^* .

Предлагаемый метод включает в себя следующие этапы.

Этап 1. Выбор основного и вспомогательного критериев для каждого уровня детализации модели бизнес-процесса.

Пусть $attribute_value_{ij}$ — атомарное высказывание, о j — значении i — атрибута. Данное высказывание является истинным в том случае, если в логе существует событие, с которым связан i — атрибут, причем в записи события этот атрибут принимает j — значение. Тогда основной критерий Kr для текущего уровня детализации определяется через конъюнкцию/дизъюнкцию атомарных высказываний $attribute_value_{ij}$ з следующим образом:

$$Kr = \bigvee_j (\bigwedge_i attribute_value_{ij}). \quad (4)$$

Данное определение основного критерия дает возможность задать альтернативные варианты значений для выбранного подмножества атрибутов артефактов бизнес-процесса.

С точки зрения практического применения это означает, что помимо основного варианта поведения процесса в модель на требуемом уровне детализации могут быть включены и альтернативные, уникальные последовательности действий. Это дает возможность выполнить сравнительный анализ типового и единичного решения и усовершенствовать модель на основе результатов анализа.

Вспомогательный критерий базируется на дополнительной извлекаемой из лога информации. Поскольку лог содержит записи о каждой реализации процесса, то в результате его анализа можно получить информацию о частоте возникновения событий, отражающих действий процесса, а также временном интервале, необходимом для достижения конечного состояния процесса.

Дополнительное сокрытие действий для заданного уровня детализации выполняется в том случае, если частота появления событий в логе ниже заданного порога, либо время достижения конечного состояния превышает заданный порог. Такое использование вспомогательного критерия позволяет выделить типовые варианты процесса.

Обратный подход заключается в выделении решений, для которых частота появления в логе ниже заданного порога. В этом случае выделяются частные решения в конкретных ситуациях, которые получены в результате изменения бизнес-процесса исполнителями.

Полный критерий представляет собой конъюнкцию основного и вспомогательного критериев.

Этап 2. Выделение и обобщение линейных последовательностей действий процесса.

Поскольку в модели (3) на заданном уровне детализации существенными являются действия, для фиксирующих событий которых выполняется условие $\forall (a, v) \in Ct^* \exists e \in \Pi : e(a) = v$, то критерием сокрытия линейной последовательности действий является наличие такой пары событий в трассе лога, которая обладает следующими характеристиками:

- события характеризуются атрибутами и их значениями из априорно отобранного подмножества контекста Ct^* ;
- все события между событиями из указанной пары не содержат атрибутов из множества Ct^* .

$$\exists \{e_n\} \in \pi, \Pi = \{\pi\}, i = \overline{1, I}, e_n = \{(a, v)\}_n : \quad (5)$$

$$(\forall i \neq 1, I \{ (a, v)\}_n \cap Ct^* = \emptyset) \wedge$$

$$\forall i = 1, I \{ (a, v)\}_n \subset Ct^*$$

где π — трасса лога, которая содержит запись одно-

кратного выполнения бизнес-процесса, e_n — событие трассы π лог; $\{(a, v)\}_n$ — множество атрибутов и их значений, которые характеризуют n -событие лог.

Этап 3. Обобщение последовательностей событий на нескольких трассах лог

В отличие от этапа 2, в данном случае выполняется сопоставление последовательностей событий на нескольких трассах лог. Сопоставляются пары одинаковых событий (e_1, e_n) , принадлежащие различным трассам лог. Если последовательность событий $(e_2, \dots, e_{i-1}, e_{i+1}, \dots, e_{n-1})$ между e_1 и e_n совпадает для всех трасс лог, то мы имеем линейную последовательность действий, которая скрывается аналогично этапу 2.

В том случае, если последовательности аналогичных событий не совпадают, в лог между событиями e_1 и e_n зафиксирован блок ветвлений. Обработка этого блока выполняется на этапе 4.

Этап 4. Выделение и обобщение блоков ветвлений в виде последовательности действий.

На данном этапе необходимо обработать две ситуации:

- во всех вариантах ветвления между событиями e_1 и e_n отсутствуют события, которые характеризуются атрибутами Ci^* ;
- существуют трассы, на которых имеются промежуточные события с атрибутами из Ci^* .

В первой ситуации выполняется сокрытие событий, которые расположены между событиями e_1 и e_n .

Во второй ситуации рекурсивно выполняется обобщение последовательностей для промежуточных событий согласно этапам 2 и 3, либо обобщение блоков ветвления согласно этапу 4.

Результатом данного этапа является множество пар событий вида (e_1, e_n) для всех трасс процесса, промежуточные события между которыми необходимо скрыть в модели.

Этап 5. Применение вспомогательного частотного критерия. Данный критерий может применяться как для сокрытия, так и для выделения уникальных, разовых последовательностей событий.

В первом случае выполняется обобщение, выделяется наиболее типовое выполнение процесса. Практическое назначение данного уровня модели — поддержка работы в нормальном режиме.

Во втором случае выполняется выделение уникальных решений, основанных на применении исполнителями своих неявных знаний. Практическое назначение данной детализации модели состоит в том, чтобы выделить, проанализировать в дальнейшем включить в модель новые проверенные на практике решения, тем самым повысив ее адекватность.

Этап 6. Построение модели бизнес-процесса с заданным уровнем детализации.

Исходными данными для текущего этапа являются:

- модель процесса, полученная на основе интеллектуального анализа логов и содержащая знаниевую и workflow — составляющие;
- набор пар событий, полученный в результате этапа 5 и отражающий существенные для данного уровня детализации действия процесса.

Вершины workflow — графа, полученные в результате анализа логов, соответствуют событиям процесса. Поэтому на данном этапе из графа потока работ удаляются все промежуточные вершины, между парами событий $E = \{(e_1, e_n)\}$, а также правила из Kn , которые не связаны с событиями из E .

Результатом данного этапа является workflow — граф Wf^* , который содержит только используемые на данном уровне рассмотрения события, соответствующие действиям процесса, также набор правил, задающих связь между подмножеством атрибутов контекста Ci^* и графом Wf^* . Отметим, что в общем случае граф Wf^* представляет собой совокупность последовательностей работ для различных значений атрибутов a_{si} в соответствии с критерием (4).

Выводы

Предложена структурная модель знание-емкого процесса на отдельном уровне детализации. Модель включает в себя: подмножество выделенных атрибутов артефактов контекста и их значений при условии наличия последних в лог процесса; подмножество правил из уровня знаний ЗБП, ограниченное зависимостями, использующими выделенные атрибуты артефактов контекста; подмножество последовательностей работ, использующее выделенные на заданном уровне детализации правила. Модель обеспечивает возможность построения обобщенного/детализованного представления ЗБП на основе априорно заданных характеристик элементов контекста, что дает возможность согласовать иерархию элементов контекста и иерархию представления процесса.

Предложен метод обобщения представления знание-емкого бизнес-процесса на основе анализа его логов. Метод включает в себя этапы выделения представленных в лог подмножеств атрибутов объектов для каждого уровня детализации и проведения обобщения действий процесса с использованием выделенных атрибутов. При проведении обобщения объединяются цепочки последовательных действий, блоки ветвления/объединения результатов ветвления, а также скрываются частные реализации процесса. Критерием обобщения является наличие в описании событий в лог процесса

заданного набора значений атрибутов. Действия процесса, которым в логике соответствуют события без заданного подмножества атрибутов, скрываются из рассмотрения на каждом уровне.

Представленный метод, в отличие от существующих, учитывает связь действий процесса с контекстом. Поскольку контекст включает в себя артефакты, соответствующие организационной структуре предприятия, то метод позволяет выделить уровни детализации/обобщения в соответствии с иерархией подразделений в организации. Это дает возможность интегрировать функциональный и процессный подходы к управлению в рамках одного предприятия.

Метод позволяет повысить эффективность интеллектуального анализа процессов путем выделения существенных для анализа и принятия решений подмножеств действий, что создает условия для решения проблемы «спагетти-подобных» workflow —моделей.

Список литературы:

1. *Weske M.* Business Process Management: Concepts, Languages, Architectures. Second Edition/ M. Weske. — Springer, 2012. — 403 p.
2. *OMG:* Business Process Model and Notation (BPMN), Version 2.0 (2011). — режим доступа: <http://www.omg.org/spec/BPMN/2.0/PDF>
3. *Vom Brocke, J.* Handbook on Business Process Management 1. Introduction, Methods, and Information Systems / J. vom Brocke, M. Rosemann. — Springer-Verlag Berlin Heidelberg, 2015. — 709 p.
4. *Cohn D.* Business artifacts: A data-centric approach to modeling business operations and processes/ Cohn D., Hull R. //IEEE Data Eng. Bull. — 2009. — №32. — pp. 3–9.
5. *Bhattacharya K.* Artifact-centered operational modeling: Lessons from customer engagements/ K. Bhattacharya, N. S. Caswell, S. Kumaran, A. Nigam, F. Y. Wu // IBM Systems Journal. — 2007. — №46 (4). — pp. 703–721.
6. *Nigam A.* Business artifacts: An approach to operational specification/ A. Nigam, N. S. Caswell/ IBM Systems Journal. — 2003. — № 42(3). — pp. 428–445.
7. *Hull R.* Business Artifacts with Guard-Stage-Milestone Lifecycles: Managing Artifact Interactions with Conditions and Events// DEBS, — 2011. — pp. 51–62.
8. *Polyvyanyy A.* Process model abstraction: a slider approach/ A. Polyvyanyy, S. Smirnov, M. Weske //Enterprise Distributed Object Computing Conference. EDOC '08. 12th International IEEE, Munchen, 15–19 Sept. 2008. — режим доступа: <https://www.researchgate.net/publication/4377458>.
9. *Polyvyanyy A.* Reducing complexity of large EPCs/ A. Polyvyanyy, S. Smirnov, M. Weske // Modellierung betrieblicher Informationssysteme — Modellierung zwischen SOA und Compliance, Saarbrücken, Nov 2008. — режим доступа: <https://www.researchgate.net/publication/221149446>.
10. *Muller D.* Data-driven modeling and coordination of large process structures / D. Muller, M. Reichert, J. Herbst // Springer. — LNCS. — Volume 4803. — p.p. 131–149.
11. *Fahland D.* Many-to-many: Some observations on interactions in artifact choreographies/ De Leoni, M., Van Dongen, B. F., van der Aalst, W. M. P. // 3rd Central-European Workshop on Services and their Composition (ZEUS), — 2011. — режим доступа: <http://ceur-ws.org/Vol-705/paper1.pdf>.
12. *Van der Aalst, W. M. P.* Process Mining: Discovery, Conformance and Enhancement of Business Processes / W. M. P. Van der Aalst. — Springer Berlin Heidelberg, 2011. — 352 p.

Поступила в редколлегию 11.11.2016

УДК 681.518:004.93.1'



І.В. Шелехов¹, Д.В. Прилепа²

¹СумДУ, м. Суми, Україна, igor-i@ukr.net

²СумДУ, м. Суми, Україна, prilepa.dmitrij@meta.ua

ВИЗНАЧЕННЯ ЛІНІЇ СИМЕТРІЇ НА ЗОБРАЖЕННІ ОБЛИЧЧЯ ЛЮДИНИ ПРИ ДІАГНОСТУВАННІ ЕМОЦІЙНО-ПСИХІЧНОГО СТАНУ

Запропоновано інформаційно-екстремальний метод визначення центральної лінії симетрії за зображеннями обличчя людини при діагностуванні її емоційно-психічного стану. При цьому показано, що при психодіагностуванні центральна лінія симетрії вибирається за умови мінімальної різниці максимумів інформаційного критерію, визначених в процесі навчання для право- і лівопівкульних портретів людини. За результатами фізичного моделювання доведено, що пошук центральної лінії симетрії зображення обличчя людини дозволяє підвищити достовірність розпізнавання психоемоційного стану людини.

ІНФОРМАЦІЙНО-ЕКСТРЕМАЛЬНА ІНТЕЛЕКТУАЛЬНА ТЕХНОЛОГІЯ, КРИТЕРІЙ ФУНКЦІОНАЛЬНОЇ ЕФЕКТИВНОСТІ, МАШИННЕ НАВЧАННЯ, ЛІНІЯ СИМЕТРІЇ, СИСТЕМА ДІАГНОСТУВАННЯ ЕМОЦІЙНО-ПСИХІЧНОГО СТАНУ

Вступ

Попереднє оброблення вхідних даних відіграє важливу роль в ефективному застосуванні методів комп'ютерної психодіагностики за зображенням обличчя [1]. До задач попереднього оброблення відносять нормалізацію зображення та стабілізацію його яскравості [2], локалізацію області обличчя на фото [3], детекцію областей очей та інших структурних елементів обличчя [4], відновлення поверхні обличчя в 3D-формі [5] тощо. Для більшості цих задач розроблено окремі алгоритми, що характеризуються малоефективним рівнем інтеграції з основними алгоритмами психодіагностування, але вважаються найбільш доречними для підвищення оперативності комп'ютерної психодіагностики. Можливість включення процедур попереднього оброблення вхідних даних у вигляді контурів категорійних моделей навчання психодіагностичної системи підтримки прийняття рішень (СППР) надає так звану інформаційно-екстремальну інтелектуальну технологію (ІЕІ-технологія) аналізу даних, що ґрунтується на максимізації інформаційної спроможності системи в процесі її машинного навчання [6]. При цьому визначення оптимальних функціональних параметрів системи як для етапу попереднього оброблення, так і безпосереднього навчання розпізнаванню емоційно-психічного стану людини за зображенням її обличчя, здійснюється шляхом пошуку глобального максимуму інформаційного критерію функціональної ефективності (КФЕ) в робочій (допустимій) області визначення його функції.

В статті розглядається інформаційно-екстремальний алгоритм машинного навчання системи діагностування емоційно-психічного стану людини за зображенням її обличчя з додатковою процедурою пошуку лінії симетрії обличчя.

1. Постановка задачі

Розглянемо формалізовану постановку задачі інформаційного синтезу системи діагностування емоційно-психічного стану людини за зображенням її обличчя, основною складовою якої є здатна навчатися СППР. Нехай дано алфавіт $\{X_m^o \mid m = \overline{1, M}\}$ класів розпізнавання, які характеризують різні емоційно-психічні стани пацієнта, і навчальну матрицю яскравості пікселів рецепторного поля зображення обличчя пацієнта $\|y_{m,i}^{(j)}\|, i = \overline{1, N}, j = \overline{1, n}$, де N, n – кількість ознак розпізнавання і реалізацій зображень відповідно. При цьому рядок матриці $\{y_{m,i}^{(j)} \mid i = \overline{1, N}\}$ визначає j -у реалізацію, а стовпчик $\{y_{m,i}^{(j)} \mid j = \overline{1, n}\}$ – навчальну вибірку значень i -ї діагностичної ознаки. Відомий структурований вектор параметрів навчання СППР

$$g = \langle x_m, d_m, \delta, h \rangle, \quad (1)$$

де x_m – еталонний (усереднений) вектор-реалізація зображення, вершина якого визначає геометричний центр контейнера класу X_m^o ; d_m – радіус контейнера класу X_m^o , що відновлюється в радіальному базисі простору ознак розпізнавання; δ – параметр симетричного поля контрольних допусків на ознаки розпізнавання; h – відстань лінії симетрії від лівої межі рецепторного поля зображення. Необхідно на етапі навчання СППР оптимізувати контейнери класів розпізнавання, відновленні на h -му кроці навчання, де параметр h вибирається з практичних міркувань наближеним до середини рецепторного поля, за умови, що усереднене значення інформаційного КФЕ навчання системи приймає максимальне значення в робочій (допустимій) області визначення його функції

$$\bar{E} = \frac{1}{M} \sum_{m=1}^M \max_{G_E \cap \{h\}} E_m^{(h)} \quad (2)$$

де $E_m^{(h)}$ – інформаційний КФЕ, обчислений в

процесі навчання при поточному значенні радіуса гіперсферичного контейнера класу X_m^o ; G_E – робоча (допустима) область визначення функції КФЕ; $\{h\}$ – множина кроків навчання СППР.

2. Категорійна модель СППР

Оскільки процес психодіагностування є слабо формалізованим, то математичну категорійну модель здатної навчатися системи психодіагностування розглянемо у вигляді діаграми відображення відповідними операторами множин, що застосовуються в процесі інформаційно-екстремального навчання [7, 8]. При цьому вхідний математичний опис СППР подамо у вигляді структури множин

$$\Delta = \langle G, T, Z, \Omega, Y, X; P, \Phi_1, \Phi_2 \rangle,$$

де G – множина вхідних факторів; T – множина моментів часу зняття інформації; Ω – простір ознак розпізнавання; Z – простір можливих емоційно-психічних станів людини; Y – множина сигналів після первинної обробки інформації; P – оператор попередньої обробки зображення, що виконує формування ліво- та правопівкульного портрету з використанням лінії симетрії з відповідної термножини H ; $\Phi_1 : Z \times \Omega \rightarrow Y$ – оператор оброблення зображення (формування вибіркової множини Y на вході СППР); оператор Φ_2 перетворює вхідну евклідову навчальну матрицю Y у бінарну матрицю X .

Категорійну математичну модель інформаційно-екстремального машинного навчання СППР для розпізнавання емоційно-психічного стану людини показано на рис. 1

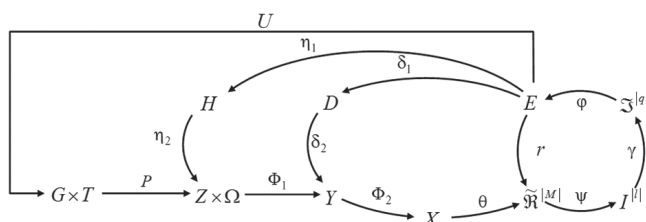


Рис. 1. Категорійна модель навчання СППР

На рис. 1 оператор θ будує розбиття $\mathfrak{R}^{|M|}$ простору ознак розпізнавання; оператор $\Psi : \mathfrak{R}^{|M|} \rightarrow I^{|I|}$ шляхом перевірки основної статистичної гіпотези про належність реалізацій образу формує множину $I^{|I|}$, де I – кількість статистичних гіпотез; оператор γ формує множину точнісних характеристик $\mathfrak{Z}^{|q|}$, де $q = I^2$; а оператор ϕ обчислює термножину E , елементами якої є значення інформаційного КФЕ. Термножина D складається із допустимих значень системи контрольних допусків на ознаки розпізнавання. При цьому оператори δ_1 і δ_2 відповідно виконують оцінку та регламентують процес оптимізації системи контрольних допусків. Оператор r замикає контур оптимізації геометричних параметрів розбиття $\mathfrak{R}^{|M|}$. Додатково

категорійна модель включає контур відновлення лінії симетрії, який реалізується операторами η_1 і η_2 . Крім того, оператор U регламентує процес навчання.

Таким чином, категорійна модель показана на рис. 1 визначає структуру алгоритму машинного навчання системи психодіагностування за зображенням обличчя, який полягає в наближенні інформаційного КФЕ до його максимального граничного значення.

3. Алгоритм навчання

Згідно з категорійною моделлю рис. 1 інформаційно-екстремальне навчання СППР у загальному випадку здійснюється за багатоциклічною ітераційною процедурою пошуку глобального максимуму інформаційного КФЕ в робочій (допустимій) області визначення його функції, яка у випадку оптимізації параметра h положення лінії симетрії на зображенні обличчя має вигляд

$$h^* = \arg \min_{G_h} \left\{ \otimes \max_{G_\delta} \left\{ \max_{G_E \cap \{h\}} \bar{E}^{(h)} \right\} \right\} \quad (3)$$

де $\bar{E}^{(h)}$ – усереднене значення КФЕ навчання системи психодіагностування, обчислене на h -му кроці навчання; G_h – область допустимих значень параметра h положення лінії симетрії обличчя; G_δ – область допустимих значень системи контрольних допусків на ознаки розпізнавання; $\otimes$ – символ операції повторення.

Як КФЕ машинного навчання СППР може використовуватися нормована ентропійна міра Шеннона [9], яка для рівномірних двоальтернативних рішень має вигляд

$$E_m^{(h)}(d) = 1 + \frac{1}{2} \left(\frac{\alpha^{(h)}(d)}{\alpha^{(h)}(d) + D_2^{(h)}(d)} \log_2 \frac{\alpha^{(h)}(d)}{\alpha^{(h)}(d) + D_2^{(h)}(d)} + \frac{\beta^{(h)}(d)}{D_1^{(h)}(d) + \beta^{(h)}(d)} \log_2 \frac{\beta^{(h)}(d)}{D_1^{(h)}(d) + \beta^{(h)}(d)} + \frac{D_1^{(h)}(d)}{D_1^{(h)}(d) + \beta^{(h)}(d)} \log_2 \frac{D_1^{(h)}(d)}{D_1^{(h)}(d) + \beta^{(h)}(d)} + \frac{D_2^{(h)}(d)}{\alpha^{(h)}(d) + D_2^{(h)}(d)} \log_2 \frac{D_2^{(h)}(d)}{\alpha^{(h)}(d) + D_2^{(h)}(d)} \right), \quad (4)$$

де $D_{1,m}^{(h)}(d)$, $D_{2,m}^{(h)}(d)$ – перша та друга достовірності відповідно, обчислені на кожному кроці навчання при поточному радіусі d контейнера класу розпізнавання X_m^o ; $\alpha^{(h)}(d)$, $\beta^{(h)}(d)$ – відповідно помилка першого та другого роду.

В рамках ІЕІ-технології відновлення роздільних гіперповерхонь класів розпізнавання відбувається в радіальному базисі бінарного простору ознак. У цьому випадку замкнену роздільну гіперповерхню,

для якої деяким способом визначено геометричний центр будемо називати контейнером відповідного класу розпізнавання. При цьому тип побудованих на етапі машинного навчання вирішальних правил визначається геометричною формою контейнерів класів розпізнавання.

Аналіз алгоритму (3) показує, що він представляє собою трьохциклічну процедуру пошуку глобального максимуму інформаційного критерію оптимізації (4). При цьому зовнішній цикл оптимізує параметр h переміщення лінії симетрії зображень обличчя людини. Наступний вкладений цикл оптимізує параметр δ поля контрольних допусків на ознаки розпізнавання. Функціями внутрішнього циклу є:

- обчислення на кожному кроці навчання інформаційного критерію (4);
- пошук глобального максимуму інформаційного критерію в робочій області визначення його функції;
- визначення оптимальних геометричних параметрів контейнерів класів розпізнавання.

Таким чином, оптимізація контейнерів класів розпізнавання полягає в ітераційній процедурі пошуку глобального максимуму критерію (4) в робочій (допустимій) області визначення його функції.

4. Результати моделювання алгоритму навчання

Вхідний математичний опис сформовано за графічними даними, наведеними в праці [10]. З метою оцінки функціональної ефективності навчання СППР реалізовано пошук лінії симетрії для фазових портретів, яка визначалася спочатку оператором візуально, а потім в процесі інформаційно-екстремального навчання підбиралася автоматично за процедурою (3). З метою редукції вирішальних правил розглядалися гіперсферичні контейнери класів розпізнавання. Дослідження виконувалося на прикладі розпізнавання одного класу, який характеризував нестабільний стан пацієнта. При цьому фотографії людини розділялися на праву та ліву половини, кожна з яких відображалася дзеркально по лінії симетрії та поєднувалася зі своєю копією.

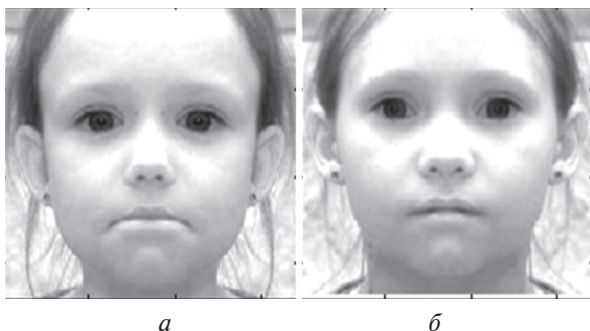


Рис. 2. Графічне відображення зображень відносно лінії симетрії, що визначалася візуально: a – лівопівкульний портрет, b – правопівкульний портрет

На рис. 2 показано зображення лівопівкульного і правопівкульного портретів, сформованих за умови, що лінія симетрії задавалася оператором візуально, тобто вручну. За результатами оброблення показаних на рис. 2 зображень було сформовано вхідну навчальну матрицю яскравості пікселів рецепторного поля розміром 150×150 пікселів. При цьому зчитування яскравості пікселів рецепторного поля здійснювалося в декартовій системі координат.

Результати оптимізації геометричних параметрів контейнерів класів розпізнавання, що відновлюються в радіальному базисі простору ознак за інформаційно-екстремальним алгоритмом навчання (3), за умови, що при формуванні портретів (рис. 2) лінія симетрії визначалася оператором візуально, показано на рис. 3.

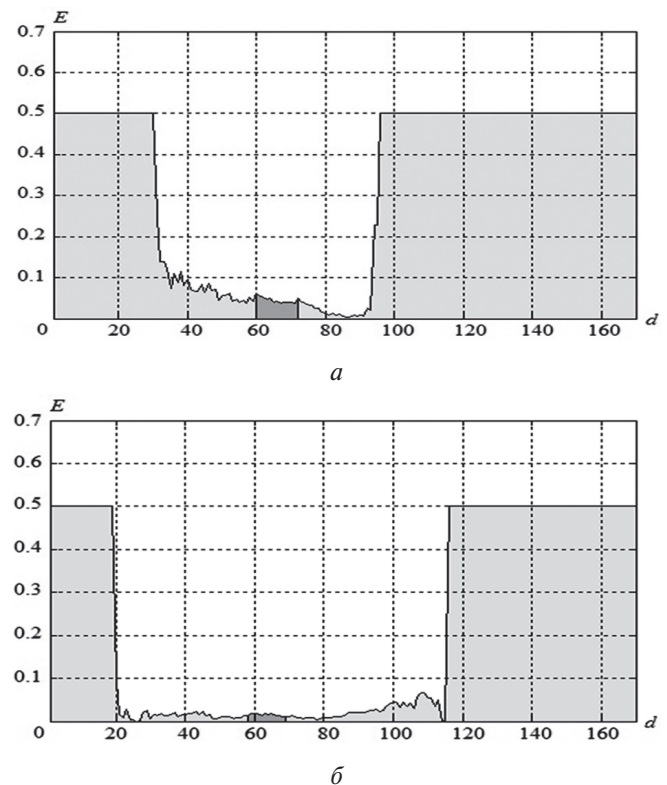


Рис. 3. Графіки залежності КФЕ (4) від радіусів контейнерів класів розпізнавання при ручному пошуку лінії симетрії: a – клас X_1^o ; b – клас X_2^o

На рис. 3 темними ділянками позначена робочі (допустимі) області визначення функції критерію (4), в яких значення першої та другої достовірностей перевершують відповідно значення помилок першого і другого роду. При цьому портрети особи з нестабільним станом характеризують клас X_1^o , сформований із лівопівкульних портретів (рис. 3а), і клас X_2^o , сформований із правопівкульних (рис. 3б).

Автоматичне визначення лінії симетрії в процесі інформаційно-екстремального навчання системи діагностуванні емоційно-психічного стану

людини здійснювалося шляхом пошуку глобального мінімуму усередненого критерію оптимізації параметра зміщення лінії симетрії, який характеризував мінімальну відмінність лівопівкульного і правопівкульного портретів. При цьому як стартова приймалася лінія симетрії, що в попередньому випадку визначалася візуально.

На рис. 4 показано динаміку зміни усередненого КФЕ (4) під час автоматичного пошуку лінії симетрії.

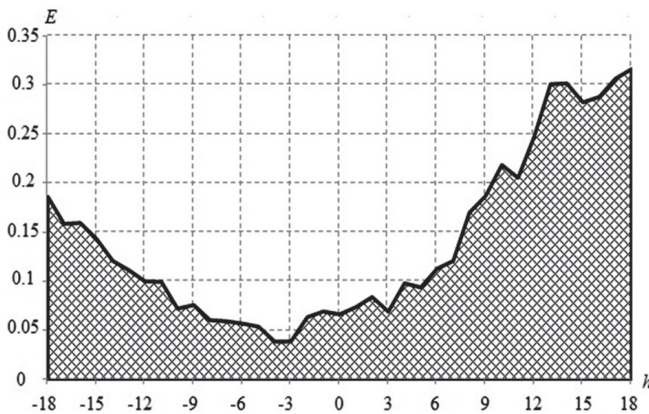


Рис. 4. Графік зміни усередненого КФЕ

Аналіз результатів (рис. 4), показує що мінімальне значення усередненого КФЕ дорівнює 0.04 при стартовому значенні 0.07. При цьому лінія симетрії розташована на три пікселі лівіше ніж та, яку було визначено візуально.

Крім того, варто зазначити, що пошук оптимального параметра h в процесі машинного навчання здійснювався в робочій області визначення інформаційного критерію, оскільки на кожному кроці навчання оптимізувалися як контрольні допуски на ознаки розпізнавання, так і радіуси контейнерів класів розпізнавання.

З метою побудови оптимальних вирішальних правил було реалізовано інформаційно-екстремальне машинне навчання системи психодіагностування із оптимізацією системи контрольних допусків на ознаки розпізнавання при автоматично визначенні за алгоритмом (3) лінії симетрії лівопівкульного і правопівкульного портретів. Попередньо було здійснено реконструкцію показаних на рис. 2 портретів із застосуванням дзеркального відображення відносно автоматично знайденої лінії симетрії відповідно лівопівкульного і правопівкульного портретів. Реконструйовані зображення показано на рис. 5.

На рис. 6 показано графіки залежності інформаційного критерію (4) від радіусів контейнерів класів розпізнавання, що відновлюються в радіальному базисі простору ознак за ітераційним алгоритмом машинного навчання (3) із автоматичним пошуком лінії симетрії.

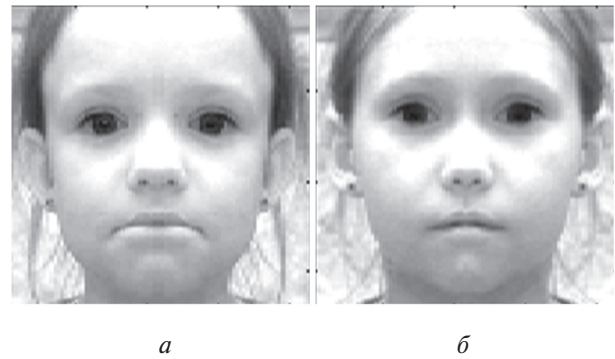


Рис. 5. Графічне відображення зображень відносно лінії симетрії, визначеної автоматично в процесі навчання системи: а — лівопівкульний портрет; б — правопівкульний портрет

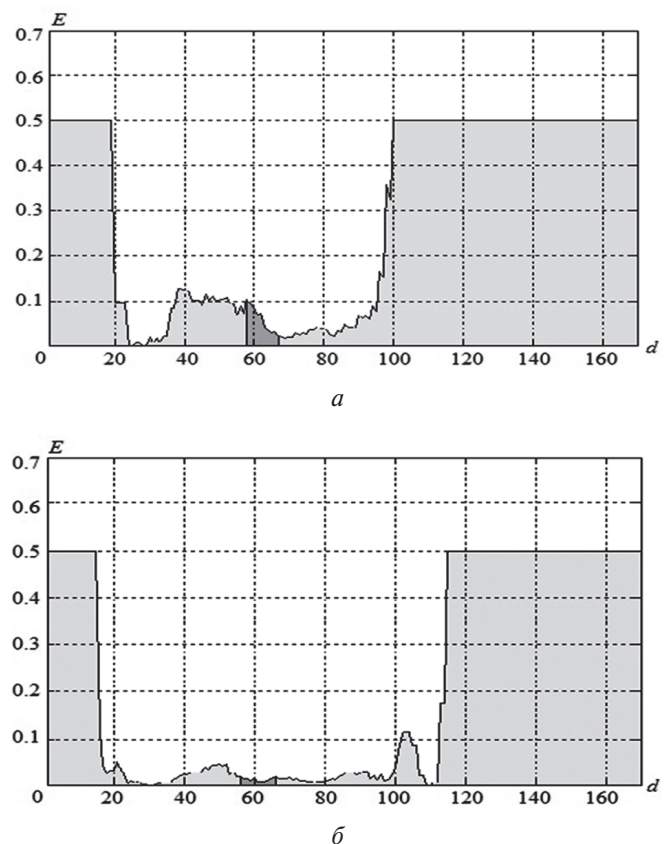


Рис. 6. Графіки залежності КФЕ (4) від радіусів контейнерів класів розпізнавання при автоматичному пошуку лінії симетрії: а — клас X_1^o ; б — клас X_2^o

На рис. 6 оптимальні значення радіусів контейнерів класів розпізнавання визначалися за навчальними матрицями, сформованою за реконструйованими портретами, показаними на рис. 5. При цьому лівопівкульні портрети особи з нестабільним станом характеризувалися класом X_1^o (рис. 6, а), а правопівкульні портрети — класом X_2^o (рис. 6, б). Порівняльний аналіз максимальних значень інформаційного критерію (4), обчислених в робочій області визначення його функції (рис. 3 і рис. 6), дозволяє стверджувати, що автоматичне визначення лінії симетрії лівопівкульних і правопівкульних

портретів показує зменшення їх відмінності, що підвищує достовірність діагностування емоційно-психічного стану людини через усунення впливу антропологічної асиметрії обличчя людини.

Висновки

В роботі запропоновано модифікацію відомої технології відео-комп'ютерної психодіагностики та корекції шляхом її інтеграції з інформаційно-екстремальною інтелектуальною технологією, що дозволяє підвищити достовірність діагностування за довільних початкових умов формування зображень обличчя людини, що діагностується.

Розроблено та програмно реалізовано алгоритм функціонування інтелектуальної системи оцінки емоційно-психічного стану пацієнта в режимі навчання системи діагностування за інформаційно-екстремальною інтелектуальною технологією дозволяє визначити автоматично лінію симетрії дзеркального відображення лівопівкульних і правопівкульних портретів, що дозволяє зменшити вплив антропологічної асиметрії обличчя на достовірність психодіагностування.

Список літератури:

1. *Leeland, K.B.* Face Recognition: New Research [Text] / K.B. Leeland — Nova Publishers, 2008. — 157 P.
2. *Huang, K.*, Symmetry-based photo-editing. Pattern Recognition / K. Huang, W. Hong, Y. Ma// - vol.38(6), 2005. P. 825 - 834.
3. *Young-Jun Song*, Face recognition robust to left/right shadows; facial symmetry / Young-Jun Song, Young-Gil Kim, Un-Dong Chang, Heak Bong Kwon// — Pattern Recognition, 2006. — vol.39(8), P. 1542–1545.
4. *Kukharev, G.*, Techniki Biometryczne: Część / G. Kukharev, A. Kuźmiński // — Szczeric (Poland): Pracownia Poligraficzna WI PS, 2003. — 310 P.
5. *Kevin, W.* A survey of approaches and challenges in 3D and multi-modal 3D + 2D face recognition / Kevin W. Bowyer, Kyong Chang, Patrick Flynn// — Computer Vision and Image Understanding, 2006. — vol.101(1), P. 1–15.
6. *Довбиш, А.С.* Основи проектування інтелектуальних систем: навчальний посібник [Текст] / А.С. Довбиш — Суми: Вид-во СумДУ, 2009. — 171 с.
7. *Шелехов, І.В.* Оптимізація параметрів навчання комп'ютеризованої системи діагностування емоційно-психічного стану людини [Текст] / І.В. Шелехов, Д.В. Прилепа // - *Радіоелектронні і комп'ютерні системи*, 2014. — №1(65). — С. 161-167.
8. *Довбиш, А.С.* Інформаційно-екстремальний алгоритм навчання системи діагностування емоційно-психічного стану людини [Текст] / А.С. Довбиш, І.В. Шелехов, Д.В. Прилепа // *Радіоелектроніка, інформатика, управління*, 2014. — №2(31). — С. 156 — 163.
9. *Довбиш, А.С.* Інтелектуальні інформаційні технології в електронному навчанні [Текст] / А.С. Довбиш, А.В. Васильєв, В.О. Любчак // — Суми: Видавництво СумДУ, 2014. — 172 с.
10. *Ануашвили, А.Н.* Объективная психология на основе волновой модели мозга [Текст] / А.Н. Ануашвили - Москва: Экон-Информ, 2008. — 292 с.

Поступила в редколегію 15.11.2016

УДК: 004.93.1'

**А.В. Васильєв¹, А.С. Довбиш², Є.С. Кулік³, З.В. Козлов⁴**¹ СумДУ, м. Суми, Україна, rector@sumdu.edu.ua² СумДУ, м. Суми, Україна, kras@id.sumdu.edu.ua³ СумДУ, м. Суми, Україна, jenyakulik92@gmail.com⁴ СумДУ, м. Суми, Україна, zakhar.kozlov@yandex.ua

ІНФОРМАЦІЙНО-АНАЛІТИЧНА СИСТЕМА АДАПТАЦІЇ НАВЧАЛЬНОГО КОНТЕНТУ ВИПУСКОВОЇ КАФЕДРИ ДО ВИМОГ РИНКУ ПРАЦІ

Розглядається інформаційний синтез здатної навчатися інформаційно-аналітичної системи оцінки відповідності навчального контенту випускової кафедри вищого навчального закладу вимогам ринку праці. Запропоновано в рамках інформаційно-екстремальної інтелектуальної технології, яка ґрунтується на максимізації інформаційної спроможності системи в процесі її навчання, алгоритм оптимізації геометричних параметрів контейнерів класів розпізнавання, що відновлюються в радіальному базисі простору ознак. Формування вхідного математичного опису інформаційно-аналітичної системи здійснювалося за результатами опитування роботодавців та випускників кафедри з досвідом роботи за базовою спеціальністю.

ІНФОРМАЦІЙНО-ЕКСТРЕМАЛЬНА ІНТЕЛЕКТУАЛЬНА ТЕХНОЛОГІЯ, НАВЧАЛЬНИЙ КОНТЕНТ, МАШИННЕ НАВЧАННЯ, ОПТИМІЗАЦІЯ, КРИТЕРІЙ ФУНКЦІОНАЛЬНОЇ ЕФЕКТИВНОСТІ

Вступ

Існуючі європейські рамкові освітні стандарти, які спрямовані на вирішення проблеми оцінки якості освіти, суттєву увагу приділяють відповідності навчального контенту вимогам ринку праці. При цьому аналіз сучасних інформаційних систем оцінки якості навчального контенту здійснюється в основному за загальними кількісними показниками працевлаштування випускників вищих навчальних закладів та рейтингом підприємств-роботодавців [1]. Оскільки кількісні показники є непрямими критеріями якості навчального контенту, то для отримання об'єктивних його оцінок в основному застосовується метод анкетування роботодавців. Недоліком такого підходу є недосконалість сучасних соціологічних методів аналізу результатів опитування респондентів, що вимагає суттєвих матеріальних і часових витрат з боку організаторів анкетування. Оскільки сучасні системи аналізу результатів соціопитування в основному базуються на застосуванні методів багатовимірного статистичного аналізу [2], то отримані результати є усередненими, що ускладнює їх детальну інтерпретацію. В праці [3] запропонована експертна система оцінки ефективності та контролю якості інформаційно-аналітичних систем оцінки знань студентів на основі нечіткої логіки, основними недоліками якої є чутливість до багатовимірності і неоднозначності рішень. Крім того, недоліками існуючих систем оцінки якості освіти є їх негнучкість і відсутність машинного аналізу зворотного зв'язку між випусковою кафедрою вищого навчального закладу, роботодавцями та студентами різних форм навчання.

Аналіз сучасного стану і тенденцій розвитку інформаційно-аналітичних систем оцінки якості навчального контенту випускової кафедри свідчить про необхідність переходу від експертних систем до систем підтримки прийняття рішень (СППР), здатних автоматично формувати базу знань, аналізувати дані та видавати рекомендації користувачам.

Одним із перспективних шляхів вирішення цієї задачі є застосування ідей і методів інформаційно-екстремальної інтелектуальної технології (ІЕІ-технології) аналізу даних, яка ґрунтується на максимізації інформаційної спроможності системи в процесі її машинного навчання [6, 7]. В праці [8] в рамках ІЕІ-технології розглядалася задача інформаційного синтезу інформаційно-аналітичної системи оцінки відповідності навчального контенту вимогам ринку праці. При цьому в процесі машинного навчання були побудовані полімодальні вирішальні правила, які враховували геометричні параметри оптимальних гіперсферичних контейнерів класів розпізнавання з окремими центрами розсіювання їх векторів-реалізацій (далі просто реалізації). Така структура вирішальних правил не дозволила отримати високу функціональну ефективність машинного навчання через суттєвий перетин класів розпізнавання. Одним із шляхів підвищення функціональної ефективності інформаційно-аналітичної системи оцінки якості навчального контенту є реалізація її машинного навчання з відновленням в радіальному базисі простору ознак вкладених контейнерів класів розпізнавання, які характеризуються єдиним центром розсіювання їх реалізацій, що має місце при прийнятті рішень за оціночними шкалами.

В статті розглядається інформаційно-екстремальний алгоритм машинного навчання інформаційно-аналітичної системи адаптації навчального контенту випускової кафедри до вимог ринку праці із вкладеними контейнерами класів розпізнавання.

1. Постановка задачі

Розглянемо формалізовану постановку задачі інформаційно-екстремального синтезу здатної навчатися інформаційно-аналітичної системи адаптації навчального контенту випускової кафедри до вимог ринку праці.

Нехай дано алфавіт класів розпізнавання $\{X_m^o | m = \overline{1, M}\}$, де M — кількість класів, які характеризують рівні якості навчального контенту, і сформовану за результатами оцінок респондентами змістовних модулів навчальних дисциплін за стобальною шкалою тривимірну навчальну матрицю $\|y_{m,i}^{(j)} | i = \overline{1, N}, j = \overline{1, n}\|$, де N — кількість ознак розпізнавання, яка дорівнює кількості змістовних модулів навчальних дисциплін, що оцінюються респондентами; n — кількість структурованих векторів-реалізацій (далі просто реалізації), координатами яких є ознаки розпізнавання. Крім того, для вкладених контейнерів класів розпізнавання з єдиним центром розсіювання їх реалізацій відомий вектор параметрів навчання інформаційно-аналітичної системи, які прямо впливають на її функціональну ефективність,

$$g_m = \langle x_1, R_m^{\text{BH}}, R_m^{\text{ЗОВ}}, \delta_K \rangle, \quad (1)$$

де x_1 — вектор-реалізація базового (внутрішнього) класу розпізнавання X_1^o , вершина якого визначає єдиний для всіх вкладених класів центр розсіювання їх реалізацій; $R_m^{\text{BH}}, R_m^{\text{ЗОВ}}$ — внутрішній і зовнішній радіуси вкладеного контейнера класу X_m^o ; δ_K — параметр симетричного поля контрольних допусків на ознаки розпізнавання. При цьому на параметр навчання δ_K накладається обмеження:

$$\delta_K < \frac{\delta_H}{2},$$

де δ_H — нормоване поле допусків на ознак розпізнавання, яке визначає область значень відповідного поля контрольних допусків.

Необхідно на етапі навчання у рамках ІЕІ-технології оптимізувати параметри навчання вектора (1) шляхом пошуку глобального максимуму усередненого за алфавітом класів розпізнавання інформаційного критерію функціональної ефективності (КФЕ) в робочій області визначення його функції:

$$\bar{E}^* = \frac{1}{M} \sum_{m=1}^M \max_{G_E \cap \{k\}} E_m^{(k)}, \quad (2)$$

де $E_m^{(k)}$ — обчислений на k -му кроці ітераційної

процедури КФЕ навчання системи розпізнавати реалізації класу X_m^o ; G_E — робоча область визначення функції КФЕ; $\{k\}$ — впорядкована множина кроків машинного навчання.

При цьому побудоване в радіальному базисі бінарного простору ознак оптимальне (тут і далі в інформаційному розумінні) вкладене розбиття класів розпізнавання відповідає умовам:

$$\begin{aligned} & (\forall X_m^o \in \tilde{\mathfrak{R}}^{|M|}) [X_m^o \neq \emptyset]; \\ & (\exists X_m^o \in \tilde{\mathfrak{R}}^{|M|}) (\exists X_c^o \in \tilde{\mathfrak{R}}^{|M|}) [X_m^o \neq X_c^o \rightarrow X_m^o \cap X_c^o \neq \emptyset]; \\ & (\forall X_{m-1}^o \in \tilde{\mathfrak{R}}^{|M|}) (\forall X_m^o \in \tilde{\mathfrak{R}}^{|M|}) (\forall X_{m-1}^o \in \tilde{\mathfrak{R}}^{|M|}) [(R_{m-1}^{\text{BH}} < \\ & < R_m^{\text{BH}} < R_{m+1}^{\text{BH}}) \& (R_{m-1}^{\text{ЗОВ}} < R_m^{\text{ЗОВ}} < R_{m+1}^{\text{ЗОВ}})]; \\ & \bigcup_{X_m^o \in \tilde{\mathfrak{R}}^{|M|}} X_m^o \subseteq \Omega_B; m \neq c, m, c = \overline{1, M}, \end{aligned}$$

де X_c^o — сусідній клас розпізнавання, що межує з класом X_m^o .

Крім того, на практиці виконуються умови

$$R_1^{\text{BH}} = 0 \text{ і } R_M^{\text{ЗОВ}} = N.$$

Таким чином, задача інформаційно-екстремального машинного навчання полягає у відновленні в радіальному просторі ознак на кожному кроці навчання контейнерів класів розпізнавання шляхом цілеспрямованого наближення значення інформаційного критерію (2) до його максимального граничного.

На етапі екзамєну, тобто безпосереднього прийняття рішень в робочому режимі, необхідно прийняти високодостовірне рішення про належність реалізації образу, що розпізнається, до одного із класів розпізнавання із заданого алфавіту.

2. Математична модель

Розглянемо категорійну модель оптимізації геометричних параметрів вкладених контейнерів класів розпізнавання в рамках інформаційно-екстремального машинного навчання інформаційно-аналітичної системи оцінки навчального контенту випускової кафедри щодо його відповідності вимогам ринку праці.

Категорійна модель представляє узагальнення орієнтованого графу, в якому відображення множин здійснюється операторами, які виконують відповідні функції машинного навчання. Вхідний математичний опис категорійної моделі подамо у вигляді структури

$$I = \langle G, T, \Omega, Z, Y, X; \Phi_1, \Phi_2 \rangle,$$

де G — простір вхідних сигналів (факторів), T — множина моментів часу одержання інформації від респондентів; Ω — простір ознак розпізнавання; Z — простір станів якості навчального контенту, Y — вибіркова множина, яка утворює вхідну

багатовимірну навчальну матрицю; X – бінарна навчальна матриця; Φ_1 – оператор формування вхідної навчальної матриці Y ; Φ_2 – оператор перетворення матриці Y в бінарну матрицю X шляхом квантування за рівнем елементів вхідної матриці навчання.

На рис. 1 показано категорійну модель інформаційно-екстремального навчання СППР з оптимізацією вкладених гіперсферичних контейнерів класів розпізнавання.

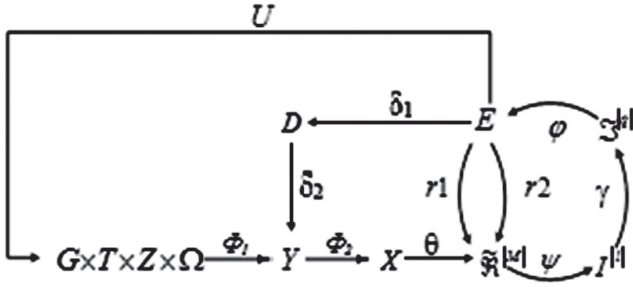


Рис. 1. Категорійна модель машинного навчання із вкладеними контейнерами класів розпізнавання

На рис. 1 оператор $\theta: X \rightarrow \tilde{\mathcal{R}}^{|M|}$ будує в загальному випадку нечітке розбиття $\tilde{\mathcal{R}}^{|M|}$ бінарного простору ознак на класи розпізнавання, а оператор класифікації Ψ перевіряє основну статистичну гіпотезу про належність вхідної реалізації класу X_m^o і таким чином формує множину гіпотез $I^{|l|}$, де l – кількість статистичних гіпотез. Оператор γ шляхом оцінки прийнятих гіпотез формує множину точнісних характеристик $\mathcal{Z}^{|q|}$, де $q = l^2$, а оператор ϕ обчислює множину значень інформаційного КФЕ, який є функціоналом від точнісних характеристик. Контур, який замикається операторами r_1 і r_2 , оптимізує внутрішні та зовнішні радіуси вкладених контейнерів класів розпізнавання шляхом пошуку глобального максимуму КФЕ в робочій (допустимій) області визначення його функції. Крім того, на рис. 1 показано контур оптимізації системи контрольних допусків, яка задається терм-множиною D .

Таким чином, побудова вирішальних правил в процесі інформаційно-екстремального машинного навчання здійснюється шляхом оптимізації геометричних параметрів вкладених гіперсферичних контейнерів класів розпізнавання, а допустимі перетворення робочої бінарної навчальної матриці – шляхом оптимізації системи контрольних допусків на ознаки розпізнавання.

3. Інформаційно-екстремальний алгоритм машинного навчання

Згідно з категорійною моделлю, показаною на рис. 1, інформаційно-екстремальний алгоритм машинного навчання розглядається у вигляді двохциклічної процедури оптимізації параметра δ_K поля контрольних допусків на ознаки розпізнавання

$$\delta_K^* = \arg \max_{G_\delta} \left\{ \frac{1}{M} \sum_{m=1}^M \max_{G_E \cap \{k\}} E_m^{(k)} \right\}, \quad (3)$$

де δ_K^* – оптимальний параметр поля контрольних допусків на ознаки розпізнавання; G_δ – допустима область значень параметра δ_K поля контрольних допусків.

Вхідними даними є тривимірна вхідна навчальна матриця і нормоване поле допусків δ_H , яке на практиці дорівнює 20 градаціям стобальної оціночної шкали.

Розглянемо основні етапи реалізації алгоритму інформаційно-екстремального навчання СППР з оптимізацією вкладених гіперсферичних контейнерів класів розпізнавання:

- 1) ініціалізація лічильника класів розпізнавання: $m := 0$;
- 2) $m := m + 1$;
- 3) ініціалізація лічильника кроків зміни параметра δ_K поля контрольних допусків на ознаки розпізнавання: $l := 0$;
- 4) $l := l + 1$;
- 5) формування бінарної навчальної матриці $\|x_{m,i}^{(j)}\|$, елементи якої визначаються за правилом:

$$x_{m,i}^{(j)} = \begin{cases} 1, & \text{if } y_{m,i}^{(j)} \in \delta_K[l]; \\ 0, & \text{if } y_{m,i}^{(j)} \notin \delta_K[l]; \end{cases}$$

- 6) визначення координат вектора $x_1[l]$, вершина якого задає єдиний центр розсіювання реалізацій класів розпізнавання, за правилом

$$x_{1,i}[l] = \begin{cases} 1, & \text{if } \frac{1}{n} \sum y_{m,i}^{(j)} \geq \rho; \\ 0, & \text{if } \text{else}, \end{cases}$$

де ρ – рівень селекції (квантування) координат вектора x_1 , який за замовчуванням дорівнює 0.5;

- 7) якщо $l \leq l_{\max}$, то виконується пункт 4, інакше – пункт 18;

- 8) ініціалізація лічильника кроків зміни внутрішнього радіуса $R_m^{(\text{BH})}[l]$;

$$9) R_m^{(\text{BH})}[l] := R_m^{(\text{BH})}[l] + 1;$$

- 10) якщо $m = 1$, то $R_m^{*(\text{BH})}[l] := 0$, інакше виконується пункт 11;

- 11) якщо $l \leq l_{\max}$, то виконується пункт 4, інакше – пункт 12;

- 12) обчислюється значення інформаційного критерію $E_m^{(\text{BH})}[l]$;

- 13) якщо $R_m^{(\text{BH})}[l] \leq N$, то виконується пункт 8, інакше – пункт 12;

- 14) для робочої області визначення функції інформаційного критерію $E_m^{(\text{BH})}[l]$ визначаються його максимальне значення $E_{m=1}^{(\text{BH})*}[l]$ і оптимальне значення радіуса $R_m^{(\text{BH})*}[l]$;

- 15) ініціалізація лічильника кроків зміни зовнішнього радіуса $R_m^{(30B)}[I] := 0$;
- 16) $R_m^{(30B)}[I] := R_m^{(30B)}[I] + 1$;
- 17) обчислюється значення інформаційного критерію $E_m^{(30B)}[I]$;
- 18) якщо $R_m^{(30B)}[I] \leq N$, то виконується пункт 15, інакше – пункт 18;
- 19) для робочої області визначення функції інформаційного критерію $E_m^{(30B)}[I]$ визначаються його максимальне значення $E_m^{(30B)*}[I]$ і оптимальне значення радіусу $R_m^{(30B)*}[I]$;
- 20) якщо $m < M$, то виконується пункт , інакше – пункт 20;
- 21) якщо $m = M$, то $R_m^{(30B)*}[I] := N$, інакше виконується пункт 22;
- 22) обчислюється максимальне середнє значення інформаційного КФЕ

$$\bar{E}^*[I] = \frac{1}{M-1} \sum_{m=1}^{M-1} \frac{E_m^{(30B)*}[I] + E_{m+1}^{(BH)*}[I]}{2} \quad (4)$$

і визначається оптимальне значення параметра δ_K^* ;
23) ЗУПИН.

Як КФЕ навчання СППР розглядалася модифікована інформаційна міра Кульбака у вигляді [5]

$$E_m^{(k)} = \log_2 \left(\frac{2 - (\alpha_m^{(k)}(d) + \beta_m^{(k)}(d))}{\alpha_m^{(k)}(d) + \beta_m^{(k)}(d)} \right) * [1 - (\alpha_m^{(k)}(d) + \beta_m^{(k)}(d))], \quad (5)$$

де $\alpha_m^{(k)}(d)$ – помилка першого роду прийняття рішення на k -му кроці навчання; $\beta_m^{(k)}(d)$ – помилка другого роду; d – дистанційна міра, яка визначає радіуси гіперсферичних контейнерів, побудованих в радіальному базисі простору Хеммінга.

Зважаючи на можливість перетину вкладених контейнерів, вирішальні правила можна побудувати у вигляді предикатного виразу

$$\begin{aligned} & [\forall x^j \in \tilde{\mathcal{R}}^{|M|}] \{x^j \in X_m^o, \\ & \text{якщо } \left[\frac{R_m^{(30B)*} - d(x_1 \oplus x^j)}{R_{m+1}^{(BH)*} - d(x_1 \oplus x^j)} > 1 \right] \\ & \text{або } [R_{M-1}^{(30B)*} < d(x_1 \oplus x^j) < N], m = \overline{1, M}, \end{aligned} \quad (6)$$

де $d(x_1 \oplus x^j)$ – кодова відстань між вектором x_1 і реалізацією x^j , що розпізнається.

Результати моделювання

Реалізація запропонованого алгоритму машинного навчання здійснювалася при створенні інформаційно-аналітичної системи оцінки відповідності навчального контенту спеціальності «Комп'ютерні науки та інформаційні технології» в Сумському державному університеті. З метою формування вхідної навчальної матриці було запропоновано 120 респондентам оцінити 50 змістовних

модулів з 10 навчальних дисциплін бакалаврського рівня, пов'язаних із професійною підготовкою фахівця в галузі інформаційних технологій. Як респонденти виступали провідні фахівці десяти ІТ-компаній, серед яких переважну кількість склали випускники кафедри комп'ютерних наук. За результатами відповідей респондентів було автоматично сформовано навчальну матрицю із трьох класів. Клас X_1^o відповідав навчальному контенту з оцінкою «добре», клас X_2^o – оцінці «задовільно», а X_3^o – «незадовільно».

Машинне навчання інформаційно-аналітичної системи здійснювалося за алгоритмом (3), який реалізовував спочатку паралельну оптимізацію системи контрольних допусків на ознаки розпізнавання за інформаційним критерієм (5).

На рис. 2 показано графік залежності усередненого КФЕ (4) від параметра δ_K поля контрольних допусків на ознаки розпізнавання. На графіку світлою ділянкою позначено робочу (допустиму) область визначення функції інформаційного критерію (5), в якій помилки $\alpha_m^{(k)}(d)$ і $\beta_m^{(k)}(d)$ менше 0,5.

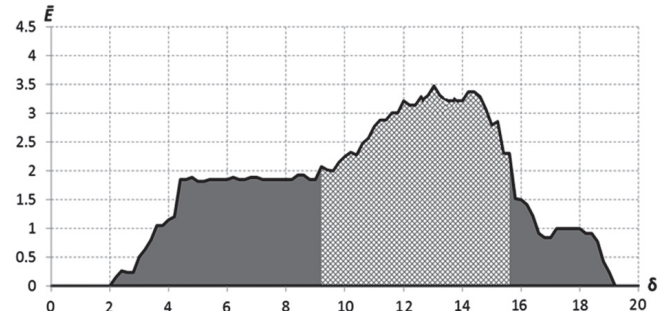


Рис. 2. Графік залежності КФЕ від параметра поля

контрольних допусків на ознаки розпізнавання

Аналіз рис. 2 показує, що оптимальне значення параметра поля контрольних допусків на ознаки розпізнавання дорівнює $\delta_K^* = \pm 13$ градацій стобальної оціночної шкали при максимальному значенні усередненого КФЕ $\bar{E}_{\max} = 3.47$, яке обчислювалося в робочій області визначення його функції.

Порівняльний аналіз результатів машинного навчання, отриманих при застосуванні вкладеної структури гіперсферичних контейнерів класів розпізнавання, показав, що у цьому випадку значення інформаційного КФЕ втричі більше ніж при полімодальній структурі ($\bar{E}_{\max} = 1.09$), яка досліджувалася в праці [8].

З метою побудови вирішальних правил (6) в процесі машинного навчання здійснювалася оптимізація радіусів вкладених контейнерів класів розпізнавання (рис. 3).

Аналіз результатів оптимізації показує, що оптимальними радіусами вкладених контейнерів класів розпізнавання є: для класу X_1^o зовнішній

радіус дорівнює $R_1^{(зов)*} = 21$ (тут і далі в кодових одиницях); для класу X_2^o внутрішній радіус дорівнює $R_2^{(вн)*} = 15$ і зовнішній — $R_2^{(зов)*} = 40$; для класу X_3^o внутрішній радіус дорівнює $R_3^{(вн)*} = 37$.

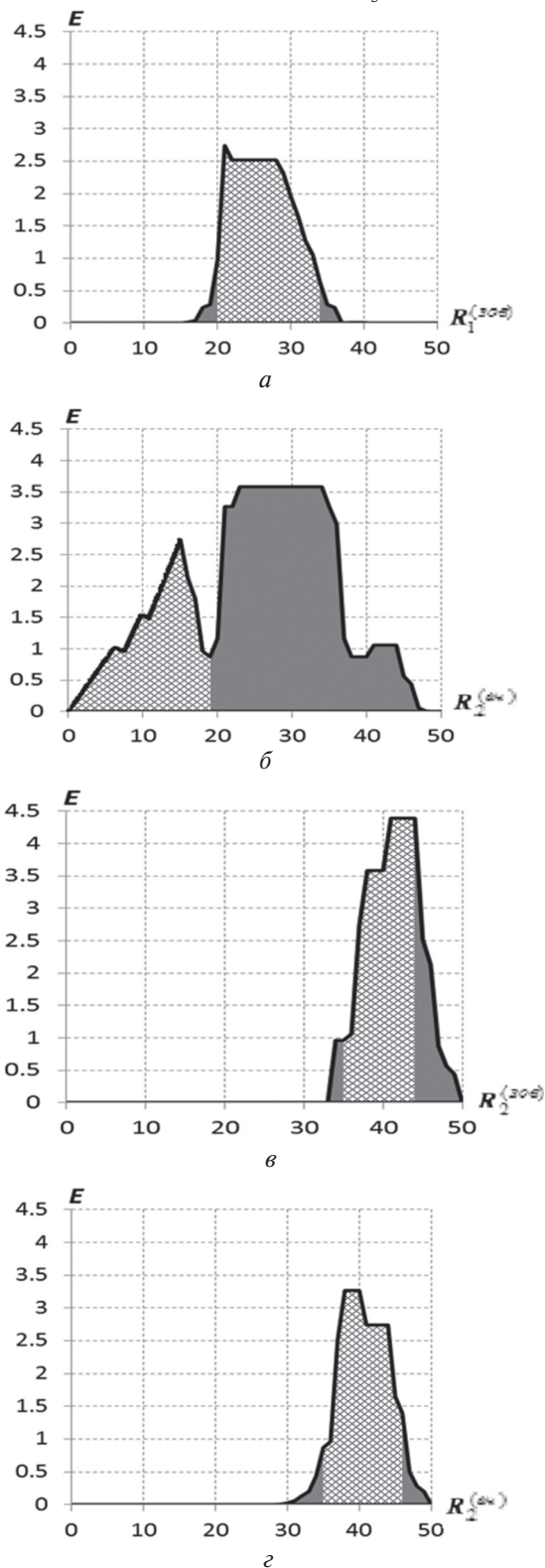


Рис. 3. Графіки залежності КФЕ (5) від радіусів контейнерів класів розпізнавання: а — зовнішній радіус $R_1^{(зов)}$ класу X_1^o ; б — внутрішній радіус $R_2^{(вн)}$ класу X_2^o ; в — зовнішній радіус $R_2^{(зов)}$ класу X_2^o ; г — внутрішній радіус $R_2^{(вн)}$ класу X_3^o

Оптимальним геометричним параметрам контейнерів класів розпізнавання відповідають такі значення КФЕ і точнісних характеристик: для класу X_1^o — $E_1^{(зов)*} = 2,74$ (перша достовірність $D_1^{(зов)*} = 0,93$, помилка другого роду $\beta^{(зов)*} = 0,05$); для класу X_2^o — $E_2^{(вн)*} = 2,65$ ($D_1^{(вн)*} = 0,90$, $\beta^{(вн)*} = 0,03$), $E_2^{(зов)*} = 4,85$ ($D_1^{(зов)*} = 0,98$, $\beta^{(зов)*} = 0,01$) і для класу X_3^o — $E_3^{(вн)*} = 3,75$ ($D_1^{(вн)*} = 0,95$, $\beta^{(вн)*} = 0$).

Для підвищення функціональної ефективності машинного навчання було застосовано алгоритм послідовної оптимізації системи контрольних допусків на ознаки розпізнавання. При цьому одержані на етапі паралельної оптимізації квазі-оптимальні контрольні допуски на ознаки розпізнавання розглядалися як стартові для їх послідовної оптимізації.

На рис. 4 показано графік зміни усередненого критерію Кульбака, отриманий в процесі послідовної оптимізації системи контрольних допусків на ознаки розпізнавання.

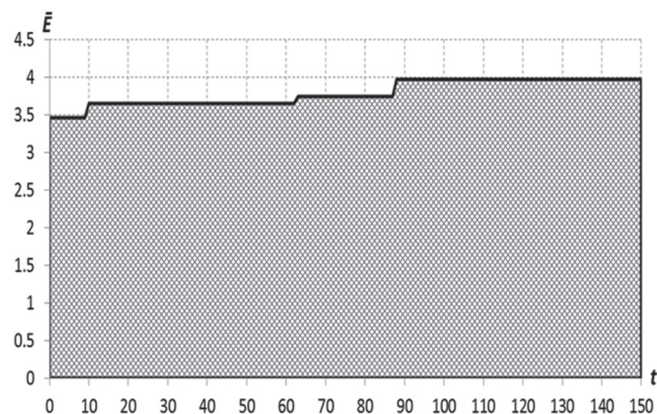


Рис. 4. Графік зміни КФЕ в процесі оптимізації контрольних допусків за послідовним алгоритмом

Як показує аналіз рис. 4, максимальне значення усередненого КФЕ було одержано вже на другому прогоні ітераційної процедури оптимізації після 93 ітерацій. При цьому кількість ітерацій на одному прогоні визначалася кількістю ознак розпізнавання, тобто дорівнювала 50. Максимальне значення усередненого КФЕ при цьому підвищилось у порівнянні з паралельною оптимізацією контрольних допусків на ознаки розпізнавання (рис. 2) і дорівнює $\bar{E}_{\max} = 3,98$. Крім того, це значення суттєво перевершує отримане в праці [7] при полімодальному класифікаторі ($\bar{E}_{\max} = 1,40$).

Висновки

Запропоновано інформаційно-екстремальний алгоритм машинного навчання інформаційно-аналітичної системи адаптації навчального контенту випускової кафедри до вимог ринку праці. За

результатами фізичного моделювання доведено, що застосування вкладених контейнерів класів розпізнавання, які в процесі навчання відновлюються в радіальному базисі простору ознак розпізнавання, дозволило суттєво підвищити функціональну ефективність машинного навчання системи у порівнянні із полімодальною структурою вирішальних правил. Крім того, застосування вкладеної структури контейнерів класів розпізнавання підвищує оперативність алгоритму машинного навчання системи, оскільки не вимагає реалізації процедури визначення для класу, що відновлюється, найближчого сусіда.

Аналіз одержаних в праці результатів показав, що побудоване в процесі навчання розбиття простору ознак є нечітким через перетин контейнерів класів розпізнавання, які обумовлюють при класифікації наявність похибок першого та другого роду. Для побудови безпомилкових за навчальною матрицею вирішальних правил згідно з принципом відкладених рішень Івахненка О. Г. необхідно здійснювати в процесі навчання оптимізацію інших параметрів навчання або переходити до більш складних типів радіально-базисних вирішальних правил.

Таким чином, аналіз одержаних в праці результатів дозволяє зробити висновок, що в задачах інформаційного синтезу здатних навчатися систем оцінки якості, де існує структурованість алфавіту

класів розпізнавання, застосування вкладених контейнерів дозволяє підвищити їх функціональну ефективність.

Список літератури:

1. Агапова, М. О. Про оптимізацію навчального матеріалу / М. О. Агапова, В. Г. Кремень // Теорія і практика управління соціальними системами. Науково-практичний журнал. — 2012. — № 1. — С. 34–39.
2. Берестнева О. Г. Компьютерные технологии в оценке качества обучения студентов / О. Г. Берестнева, А. С. Глазырин // Известия Томского политехнического университета. — 2003. — № 6. — С. 106–112.
3. Петров Е. Г. Методи і засоби прийняття рішень у соціально-економічних системах: Навчальний посібник / Е. Г. Петров, М. В. Новожилов, І. В. Гребеннік. — К.: Техніка, 2004. — 256 с.
4. Солодовников И. В. Экспертная система оценки эффективности обучения на основе математического аппарата нечеткой логики / И. В. Солодовников, О.В. Рогозин, О.В. Шуруев — Научный журнал «Качество Инновации Образование» №1, 2006.
5. Довбиш А. С. Основи проектування інтелектуальних систем: Навчальний посібник / А.С. Довбиш. — Суми: Видавництво Сумського державного університету, 2009. — 171 с.
6. Довбиш А. С. Інтелектуальні інформаційні технології в електронному навчанні / А.С. Довбиш, А.В. Васильєв, В.О. Любчак. — Суми: Видавництво Сумського державного університету, 2013. — 172 с.
7. Довбиш А. С. Інформаційно-екстремальне навчання системи оцінки якості навчального контенту випускової кафедри / А.С. Довбиш, Є. С. Кулік, З. В. Козлов, А. С. Осадчий // Радіоелектронні і комп'ютерні системи. — 2016. — №3. — С. 71–77.

Надійшла до редакції 18.11.2016

РЕФЕРАТИ

РЕФЕРАТЫ

ABSTRACTS

УДК 658.012.011.56

Лінгвістика та системний підхід. Частина 2 / В.А. Широков // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 3–11.

У статті детально розглянуті принципи системного підходу, що базуються на триаді «структура — субстанція — суб'єкт». Запропоновано список лінгвістичних презумпцій. Визначено поняття субстанції думки і природної мови як інструменту для формалізації процесу мислення. Проведено аналіз визначення терміну штучний інтелект як форми індивідуалізації технічних систем, що володіють мовним статусом. Як результат дослідження пропонуються десять системних теорем.

УДК 658.012.011.56

Лингвистика и системный подход. Часть 2 / В.А. Широков // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 3–11.

В статье детально рассмотрены принципы системного подхода, базирующиеся на триаде «структура — субстанция — субъект». Предложен список лингвистических презумпций. Определено понятие субстанции мысли и естественного языка как инструмента для формализации процесса мышления. Проведен анализ определения термина «искусственный интеллект» как формы индивидуализации технических систем, владеющих языковым статусом. Как результат исследования предлагается десять системных теорем.

UDC 658.012.011.56

Linguistics and systematic approach. Part 2 / V.A. Shirokov // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 3–11.

The principles of the system approach based on definition “structure — substance — subject” were analyzed in detail. The concept of substance thought and natural language as a tool for the formalization of the thinking process was defined. Author propose a discussion the concept of artificial intelligence as a form of individualization of technical systems with language status. Ten theorems are proposed.

УДК 81'322.2'33

Моделювання базових складових процесу розуміння тексту в системі автоматичного реферування / О.Ю. Айвас, О.В. Лазаренко, Д.І. Панченко // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 12–15.

Запропонована процедура побудови прообразів рефератів в процесі смислового аналізу тексту в системі автоматичного реферування з використанням текстових баз і концептуальних інваріанти текстів.

Бібліогр.: 6 найм.

УДК 81'322.2'33

Моделирование базовых составляющих процесса понимания текста в системе автоматического реферирования / Е.Ю. Айвас, О.В. Лазаренко, Д.И. Панченко // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 12–15.

Предложена процедура построения прообразов рефератов в процессе смыслового анализа текста в системе автоматического реферирования с использованием текстовых баз и концептуальных инвариант текстов.

Библиогр.: 6 найм.

УДК 81'322.2'33

Modeling of the basic components of the process of understanding the text in the system of auto summarization / E.Yu. Ivas, O.V. Lazarenko, D.I. Panchenko // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 12–15.

The procedure of the creation of summary prototypes in the process of the text semantic analysis based on the text bases and conceptual invariants of texts in the automatic text summarization system is proposed.

Ref.: 6 items.

УДК 004.853

Метод двохетапної класифікації електронних текстів / Л.Е. Чала, С.Г. Удовенко, Є.В. Кушвід // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 16–23.

У статті розглянуто метод двохетапної класифікації електронних документів. На першому етапі використовується поліноміальна модель байєсівського класификатора для фільтрації псевдоспаму. На другому етапі здійснюється тематична класифікація документів основного масиву з використанням зв'язних лінгвістичних дескрипторів.

Іл. 8. Бібліогр.: 8 найм.

УДК 004.853

Метод двухэтапной классификации электронных текстов / Л.Э. Чала, С.Г. Удовенко, Е.В. Кушвид // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 16–23.

В статье рассмотрен метод двухэтапной классификации электронных документов. На первом этапе используется полиномиальная модель байесовского классификатора для фильтрации псевдоспама. На втором этапе осуществляется тематическая классификация документов основного массива с использованием связанных лингвистических дескрипторов.

Ил. 8. Библиогр.: 8 найм.

UDK 004.853

Method of a twostage electronic text classification / L.E. Chala, S.G. Udovenko, Ye.V. Kushvid // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 16–23.

The method of a twostage classification of electronic documents is considered in article. On the first stage the polynomial model of Bayesian classifier is used for filtration of pseudospam. On the second stage thematic classification of documents of basic document array is realised with the use of coherent linguistic descriptors.

Fig. 8. Ref.: 8 items.

УДК 519.7

Лексикографічна система електронного термінологічного словника / О.С. Пузік // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 24–29.

Наведено підходи до обробки сканованих текстів. Розглянуті варіанти побудови моделі даних для тримовного словника. Наведено та детально описано схему лексикографічної бази даних та архітектури програмної системи. Основною перевагою системи є рівноправність мов та відносно простий перехід до багатомовного словника при подальшому розвитку системи.

Рис.: 5. Бібліогр.: 7 найм.

УДК 519.7

Лексикографическая система электронного терминологического словаря / А.С. Пузик // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 24–29.

Приведены подходы к обработке сканированных текстов. Рассмотрены варианты построения модели данных для трехязычного словаря. Приведена и детально описана схема лексикографической базы данных и архитектуры программной системы. Основным преимуществом системы является равноправность языков и относительно простой переход к многоязычному словарю при дальнейшем развитии системы.

Рис.: 5. Библиогр.: 7 найм.

UDC 519.7

Lexicographic system of electronic terminological dictionary / O.S. Puzik // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 24–29.

Approaches to handle scanned texts are described. Alternatives for trilingual dictionary data model organization are considered. The article describes scheme of lexicographical database of the dictionary and software architecture as well. The main feature of the system is the equality of languages and ease switching to multilingual dictionary on further development of the system.

Fig.: 5. Ref.: 7 items.

УДК 004.942

Застосування глибоких згорткових нейронних мереж для класифікації риноманометричних даних / А.Л. Єрохін, А.С. Нечипоренко, А.С. Бабій, О.П. Турута // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 30–34.

Стаття присвячена розробці методу попередньої обробки риноманометричних даних, що дозволяє здійснювати автоматизоване виявлення і видалення некоректних вимірювань у режимі реального часу. Розглянуто особливості застосування математичного апарату у задачах медичної діагностики. Запропоновано формальну модель процесу діагностики порушень функції носового дихання на базі згорткових нейронних мереж. Пропонується рішення задачі класифікації риноманометричних даних на основі використання глибоких конволюційних нейронних мереж.

Ил. 6, табл. 1. Бібліогр.: 24 найм.

УДК 004.942

Применение глубоких сверточных нейронных сетей для классификации риноманометрических данных / А.Л. Ерохин, А.С. Нечипоренко, А.С. Бабий, А.П. Турута // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 30–34.

Статья посвящена разработке метода предварительной обработки риноманометрических данных, что позволяет осуществлять автоматизированное выявление и удаление некорректных измерений в режиме реального времени. Рассмотрены особенности применения математического аппарата в задачах медицинской диагностики. Предложена формальная модель процесса диагностики нарушений функции носового дыхания на базе сверточных нейронных сетей. Предлагается решение задачи классификации риноманометрических данных на основе использования глубоких конволюционных нейронных сетей.

Ил. 6, табл. 1. Библиогр.: 24 найм.

UDC 004.942

Application of Deep Convolutional Neural Networks for Classification Rinomanometric Data / A.L. Yerokhin, A.S. Nechiporenko, A.S. Babii, O.P. Turuta // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 30–34.

Article is devoted to development of preprocessing method of rinomanometric data. The software system allows detecting and removing of incorrect measurements in real-time. A formal model of diagnosis of nasal breathing function is proposed. Classification procedure of rinomanometric data is based on usage of deep convolution neural networks.

Fig. 6, Tabl. 1, Ref.: 24 items.

УДК 681.513

Методологічні аспекти сегментної обробки Кирліан-об'єктів / А.С. Свіридов, Ю.Ю. Завизиступ, О.П. Міхаль // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 35–41.

Стосовно алгоритмізації обробки Кирліан-об'єктів для медичної діагностики, з'ясовано доцільність обраного підходу. Сформульовано загальну структуру та етапність рішення. Серед варіантів реалізації, запропоновано фрагментування та фрактальний опис. Ключові моменти розглянутих алгоритмічних рішень одмакетовані з демонстрацією їх працездатності.

Табл. 1. Лл. 3. Бібліогр.: 10 найм.

УДК 681.513

Методологические аспекты сегментной обработки Кирлиан-объектов / А.С. Свиридов, Ю.Ю. Завизиступ, О.Ф. Михаль // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 35–41.

Касательно алгоритмизации обработки Кирлиан-объектов для медицинской диагностики, определена целесообразность выбранного подхода. Сформулирована общая структура и этапность решения. Среди вариантов реализации предложены фрагментирование и фрактальное описание. Ключевые моменты рассмотренных алгоритмических решений отмакетированы с демонстрацией их работоспособности.

Табл. 1. Ил. 3. Библиогр.: 10 наим.

UDK 681.513

Methodological aspects of segment processing of the Kirlian-objects / A.S. Sviridov, Yu.Yu. Zavizistup, O.Ph. Mikhal // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 35–41.

The appropriateness is found regarding the algorithmical processing of the Kirlian-objects for the medical diagnosis purposes. The general structure is defined. The fragmentation and fractal description are proposed among the options of implementing. The key points are discussed as to the algorithmic solutions.

Tabl. 1. Fig. 3. Ref.: 10 items.

УДК 681.513

Еволюціонування мультиагентної системи як аналог формування індивідуального інтелекту людини / О.П. Міхаль // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 42–47.

Еволюція мультиагентної системи, як інструменту посилення людського інтелекту, зіставлена з етапністю розвитку індивідуального інтелекту людини. Факт встановлення відповідності цікавий в плані прогнозування напрямків розвитку конкретних мультиагентних розробок.

Табл. 1. Лл. 1. Бібліогр.: 4 найм.

УДК 681.513

Эволюционирование мультиагентной системы как аналог формирования индивидуального интеллекта человека / О.Ф. Михаль // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 42–47.

Эволюция мультиагентной системы как инструмента усиления человеческого интеллекта сопоставлена с этапностью развития индивидуального интеллекта человека. Факт установления соответствия интересен в плане прогнозирования направлений развития конкретных мультиагентных разработок.

Табл. 1. Ил. 1. Библиогр.: 4 наим.

UDK 681.513

Evolution of the multi-agent system as an analogue of the formation of individual human intellect / O.Ph. Mikhal // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 42–47.

Evolution of multi-agent systems, as a tool for strengthening of human intelligence, is comparable with the phasing of the individual human intellect. The fact of the establishment of conformity is interesting in terms of forecasting the development of specific multi-agent applications.

Tabl. 1. Fig. 1. Ref.: 4 items.

УДК 519.62

Мультиалгебраїчні дискретні системи при наявності алгебраїчної структури носія. Повідомлення 1 / А.Г. Каграманян, Г.Г. Четвериков, В.В. Шляхов // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 48–52.

Знайдені вимоги існування мультиалгебраїчної системи, де елементи — класи еквівалентностей, які індукуються бінарними та тернарними відношеннями. Розглянуто вид конкретної алгебраїчної структури у вигляді лінійного простору.

УДК 519.62

Мультиалгебраические дискретные системы при наличии алгебраической структуры носителя. Сообщение 1 / А.Г. Каграманян, Г.Г. Четвериков, В.В. Шляхов // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 48–52.

Найдены условия существования мультиалгебраической системы, элементами которой являются классы эквивалентностей, индуцированных бинарными и тернарными отношениями. Рассмотрен конкретный вид алгебраической структуры в виде линейного пространства.

UDC 519.62

Multialgebraic discret systems in presence of algebraic structure of carrier. Part 1 / A.G. Kagramanyan, G.G. Chetverikov, V.V. Shlyahov // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 48–52.

The study describes conditions of existence of multialgebraic system which elements are equivalence classes that induced by binary and ternary relations. It is described certain type of algebraic structure in the form of linear space.

УДК 519.62

Мультиалгебраїчні дискретні системи при наявності алгебраїчної структури носія. Повідомлення 2 / А.Г. Каграманян, Г.Г. Четвериков, В.В. Шляхов // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 53–58.

Розглянуто вид конкретної алгебраїчної структури у вигляді лінійного простору. При цьому багато питань, що пов'язані з ідентифікацією лінійних операторів в умовах компаратора наводять до ситуації, коли область визначення операторів уявляється як частина множини входних сигналів. Конус лінійного простору — це один з найбільш розповсюджених та цікавих випадків.

УДК 519.62

Мультиалгебраические дискретные системы при наличии алгебраической структуры носителя. Сообщение 2 / А.Г. Каграманян, Г.Г. Четвериков, В.В. Шляхов // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 53–58.

Рассмотрен конкретный вид алгебраической структуры в виде линейного пространства. При этом многие вопросы, связанные с идентификацией линейных операторов в условиях компаратора приводят к ситуации, когда область определения операторов представляет часть множества входных сигналов. Конус линейного пространства один из наиболее распространенных интересных случаев.

UDC 519.62

Multialgebraic discret systems in presence of algebraic structure of carrier. Part 2 / A.G. Kagramanyan, G.G. Chetverikov, V.V. Shlyahov // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 53–58.

A specific type of algebraic structure of a linear space is described. However, many issues related to the identification of linear operators in conditions of comparator lead to situation when operators' domain is a part of variety of input signals. Cone of linear space is one of the most common cases of interest.

УДК 519.876.5

Про нечіткі рекурентні відображення при мультиагентному моделюванні популяційної динаміки / Д.І. Чумаченко, С.В. Яковлев // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 59–63.

У статті розглянуто підходи до побудови нечіткої лінгвістичної моделі для опису систем популяційної динаміки. Розроблена модель дозволяє позбутися від невизначеностей, пов'язаних зі змінними, граничними умовами, початковими станами, значеннями параметрів і т.д. в мультиагентних системах.

Табл. 1. Лл. 0. Бібліогр: 19 найм.

УДК 519.876.5

О нечетких рекуррентных отображениях при мультиагентном моделировании популяционной динамики / Д.И. Чумаченко, С.В. Яковлев // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 59–63.

В статье рассмотрены подходы к построению нечеткой лингвистической модели для описания систем популяционной динамики. Разработанная модель позволяет избавиться от неопределенностей, связанных с переменными, граничными условиями, начальными состояниями, значениями параметров и т.д. в мультиагентных системах.

Табл. 1. Ил. 0. Библиогр.: 19 наим.

UDC 519.876.5

About fuzzy recurrent mapping at multiagent simulation of population dynamics / D.I. Chumachenko, S.Vs. Yakovlev // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 59–63.

The article describes the approaches to the construction of a fuzzy linguistic model to describe the population dynamics systems. The developed model allows to get rid of the uncertainties associated with variables, boundary conditions, initial conditions, parameter values, etc. in multiagent systems.

Tab. 1. Fig. 0. Ref: 19 items.

УДК 004.82

Розробка узагальненої моделі процесу розв'язання задачі з інтервальним представленням часу / І.В. Левикін // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 64–68.

Запропоновано трьохрівневу модель процесу розв'язання задачі як елементу прецеденту. Узагальнена модель процесу розв'язання охоплює рівні упорядкованих послідовностей подій, дискретного та інтервального представлень, а також правила перетворень між рівнями. Запропонована модель забезпечує можливість побудови інтервальної моделі процесу шляхом доповнення інтервалами виконання дій процесу дискретного рівня представлення.

Бібліогр.: 10 найм.

УДК 004.82

Разработка обобщенной модели процесса решения задачи с интервальным представлением времени / И.В. Левыкин // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 64–68.

Предложена трехуровневая модель процесса решения задачи как элемента прецедента. Обобщенная модель процесса решения охватывает уровни упорядоченных последовательностей событий, дискретного и интервального представлений, а также правила преобразований между уровнями. Предложенная модель обеспечивает возможность построения интервальной модели процесса путем дополнения интервалами выполнения действий процесса дискретного уровня представления.

Библиогр.: 10 наим.

UDC 004.82

Development of generalized model of the process of solving the problem with the representation of the time interval / I.V. Levykin // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 64–68.

A three-level model of the process of solving the problem as a precedent element. The generalized model of decision process covers the levels of ordered sequences of events, discrete and interval performances, as well as the transformation rules between the levels proposed model makes it possible to construct a model of the process interval by supplementing perform actions in intervals discrete representation of the level of the process.

Ref.: 10 items.

УДК 681.513

Мультиагентна інтерпретація генерації подій в моделі системи масового обслуговування / О.П. Міхаль, О.Г. Лебедев // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 69–76.

На прикладі системи масового обслуговування продемонстровано мультиагентну інтерпретацію застосування принципів локально-паралельної обробки при побудові статистичної моделі.

Табл. 0. Іл. 4. Бібліогр.: 8 найм.

УДК 681.513

Мультиагентная интерпретация генерации событий в модели системы массового обслуживания / О.Ф. Михаль, О.Г. Лебедев // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 69–76.

На примере системы массового обслуживания продемонстрирована мультиагентная интерпретация применения принципов локально-параллельной обработки при построении статистической модели.

Табл. 0. Ил. 4. Библиогр.: 8 наим.

УДК 681.513

Multiagent interpretation of the generation of events in the model of the mass service system / O.Ph. Mikhal, O.G. Lebedev // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 69–76.

The multiagent interpretation is demonstrated on example of the system of mass service with the using of principle of local-parallel processing in building of the statistical model.

Tabl. 0. Fig. 4. Ref.: 8 items.

УДК 519.6

Метод сплайн-інтерлінації при знаходженні найбільших (найменших) значень функції двох змінних в замкнутій області / О.М. Литвин, О.В. Ярмош, Т.І. Чорна // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 77–82.

В даній роботі поставлено задачу знайти найбільше або найменше значення функції за допомогою метода сплайн-інтерлінації функції двох змінних. Для цього двомірна задача зводиться до послідовності одномірних задач на знаходження найбільшого або найменшого значення на системі паралельних горизонтальних та вертикальних прямих, що перетинають область дослідження. Результати обчислювального експерименту підтвердили його ефективність та обумовлюють можливість використання даного методу при розв'язуванні багатоекстремальних задач.

Табл. 2. Іл. 6. Бібліогр.: 13 найм.

УДК 519.6

Метод сплайн-интерлинации при нахождении наибольших (наименьших) значений функции двух переменных в замкнутой области / О.Н. Литвин, Е.В. Ярмош, Т.И. Черная // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 77–82.

В данной работе поставлена задача найти наибольшее или наименьшее значения функции с помощью метода сплайн-интерлинации функции двух переменных. Для этого двухмерная задача сводится к последовательности одномерных задач на нахождение наибольшего и наименьшего значения на системе параллельных горизонтальных и вертикальных прямых, пересекающих область исследования. Результаты вычислительного эксперимента подтвердили его эффективность и обуславливают возможность использования данного метода при решении многоэкстремальных задач.

Табл. 2. Ил. 6. Библиогр.: 13 наим.

УДК 519.6

Spline interlineation method in finding the largest and least values for function of two variables in the closed domain / O.N. Litvin, E.V. Yarmosh, T.I. Chorna // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 77–82.

In this article it is set the task to find the maximum or minimum value of a function using the method of spline interlineation function of two variables. To resolve it the two-dimensional problem is reduced to a sequence of one-dimensional problems of finding the largest and the least values in the system of parallel horizontal and vertical lines intersecting area of research. The

results of computational experiment have confirmed its effectiveness and give rise to the possibility of using this method in solving problems with many extreme points.

Tabl. 2. Fig. 6. Ref.: 13 items.

УДК 004.932.2:004.93'1

Класифікація зображень на основі кластерного подання структурних описів / В.О. Гороховатський, В.С. Столяров // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 83–87.

У роботі запропоновані методи класифікації в структурному розпізнаванні зображень на основі перетворення їх описів до кластерного вигляду. Наведено результати експериментів по класифікації на моделях структурних описів.

Табл. 03. Іл. 01. Бібліогр.: 07 найм.

УДК 004.932.2:004.93'1

Классификация изображений на основе кластерного представления структурных описаний / В.А. Гороховатский, В.С. Столяров // Бионика интеллекта: науч.-техн. журнал. — 2016. — № 2 (87). — С. 83–87.

В работе предложены методы классификации в структурном распознавании изображений на основе преобразования их описаний к кластерному виду. Приведены результаты экспериментов по классификации на моделях структурных описаний.

Табл. 03. Ил. 01. Библиогр.: 07 наим.

UDC 004.932.2:004.93'1

Classification of images based on cluster definitions of structural representations / V.A. Gorohovatsky, V.S. Stolyarov // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 83–87.

This paper provides the classification methods in structural image recognition based on transformation of their descriptions to the cluster view. The results of experiments on the classification of models of structural descriptions.

Fig. 01. Tab. 03. Ref.: 07 items.

УДК 519.6

Відновлення зображень в зонах відсутності попиксельної інформації з використанням інтерстрипації / О.М. Литвин, О.О. Литвин, Г.Д. Лесной, О.В. Славик // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 88–93.

Задача відновлення пошкоджених зображень виникає в різних галузях людської діяльності. В даній роботі проведено огляд існуючих методів відновлення зображень. На основі стандартного методу інтерстрипації, запропонованого Литвином О.Н. та Матвеевою С.Ю., в даній статті наведено метод інтерстрипації для відновлення пошкоджених зображень в зонах відсутності інформації про нього. Проведено серію обчислювальних експериментів.

Іл. 4. Бібліогр.: 13 найм.

УДК 519.6

Восстановление изображений в зонах отсутствия попиксельной информации с использованием интерстрипации / О.Н. Литвин, О.О. Литвин, Г.Д. Лесной, А.В. Славик // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 88–93.

Задача восстановления поврежденных изображений возникает в разных отраслях человеческой деятельности. В данной работе проведен обзор существующих методов восстановления изображений. На основе стандартного метода интерстрипации, предложенного Литвиным О.Н. и Матвеевой С.Ю., в данной статье приведен метод интерстрипации для восстановления поврежденных изображений в зонах отсутствия информации о нем. Проведена серия вычислительных экспериментов.

Ил. 4. Библиогр.: 13 наим.

UDC 519.6

Image restoration in the zones of absence of pixel data using interstripation / O.M. Lytvyn, O.O. Lytvyn, G.D. Lisny, O.V. Slavik // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 88–93.

The problems of restoring corrupted images arise in different areas of human activity. In this work, was reviewed the existing image inpainting techniques. On the basis of the standard interstripation method, proposed in works of Lytvyn O.M. and Matveeva S.Y., in this article describes a method of interstripation to recover corrupted images in the areas of the absence of information about it. Series of numerical experiments was carried out.

Fig. 4. Ref.: 13 items.

УДК 681.3.07

Основні етапи обробки зображень цитологічних препаратів з використанням ідеології вейвлетів / В.В. Ляшенко, О.А. Кобилін, С.М. Томич // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 94–100.

Розглянуто загальну процедуру обробки зображень цитологічних препаратів на основі використання ідеології вейвлетів. Відображено позитивні та негативні сторони застосування контрастування зображень цитологічних препаратів з метою їх обробки за допомогою вейвлет перетворення.

Іл. 14. Бібліогр.: 15 найм.

УДК 681.3.07

Основные этапы обработки изображений цитологических препаратов с использованием идеологии вейвлетов / В.В. Ляшенко, О.А. Кобылин, С.Н. Томич // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 94–100.

Рассмотрена общая процедура обработки изображений цитологических препаратов на основе использования идеологии вейвлетов. Отображены положительные и отрицательные стороны использования контрастирования изображений цитологических препаратов с целью их обработки с помощью вейвлет преобразования.

Ил. 14. Библиогр.: 15 наим.

UDK 681.3.07

The main processing stages of cytological drugs images using wavelets ideology / V.V. Lyashenko, O.A. Kobylin, S.M. Tomich // *Bionics of intelligence: Sci. mag.* — 2016. — № 2 (87). — P. 94–100.

The general procedure of processing cytologic drugs images on the basis of the wavelets ideology. Positive and negative aspects of usage contrasting images of cytologic drugs for their processing via wavelet transform.

Fig. 14. Ref.: 15 items.

УДК 004.891.3

Метод узагальнення представлення знання-ємного бізнес-процесу / О.В. Чала // *Біоніка інтелекту: наук.-техн. журнал.* — 2016. — № 2 (87). — С. 101–105.

Запропоновано структурну модель знання-ємного процесу на окремому рівні деталізації. Модель забезпечує можливість побудови узагальненого / деталізованого представлення такого процесу на основі апріорно заданих характеристик елементів контексту, що дає можливість узгодити ієрархію елементів контексту і ієрархію представлення процесу. Запропоновано метод узагальнення представлення знання-ємного бізнес-процесу на основі аналізу його логів. Метод включає в себе етапи виділення представлених в лозі підмножин атрибутів об'єктів для кожного рівня деталізації і узагальнення дій процесу з використанням виділених атрибутів. При проведенні узагальнення об'єднуються ланцюжки послідовних дій, блоки вибору / об'єднання результатів вибору, а також ховаються приватні реалізації процесу. Це дає можливість інтегрувати функціональний і процесний підходи до управління в рамках одного підприємства. Метод дозволяє підвищити ефективність інтелектуального аналізу процесів шляхом виділення істотних для аналізу і прийняття рішень підмножин дій, що створює умови для вирішення проблеми «спагетті-подібних» workflow-моделей.

Бібліогр.: 12 найм.

УДК 004.891.3

Метод обобщения представления знание-емкого бизнес-процесса / О.В. Чалая // *Бионика интеллекта: научн.-техн. журнал.* — 2016. — № 2 (87). — С. 101–105.

Предложена структурная модель знание-емкого процесса на отдельном уровне детализации. Модель обеспечивает возможность построения обобщенного / детализированного представления такого процесса на основе априорно заданных характеристик элементов контекста, который дает возможность согласовать иерархию элементов контекста и иерархию представления процесса. Предложен метод обобщения представления знание-емкого бизнес-процесса на основе анализа его логов. Метод включает в себя этапы выделения представленных в логе подмножеств атрибутов объектов для каждого уровня детализации и обобщения действий процесса и использованием выделенных атрибутов. При проведении обобщения объединяются цепочки последовательных действий, блоки выбора / объединения результатов выбора, а также скрываются частные реализации процесса. Это дает возможность интегрировать функциональный и процессный подходы к управлению в рамках одного подмножества. Метод позволяет повысить эффективность интеллектуального анализа процессов путем выделения существенных для анализа и принятия решений подмножеств действий, который создает условия для решения проблемы «спагетти-подобных» workflow-моделей.

Библиогр.: 12 наим.

UDC 004.891.3

Developing a method of forming a multi detailed representation of the knowledge-intensive business process / O.V. Chala // *Bionics of intelligence: Sci. mag.* — 2016. — № 2 (87). — P. 101–105.

A structural model of multi detailed representation of the knowledge-intensive business process is proposed. The model provides the possibility of constructing a generalized / a detailed presentation of this process on the basis of a priori given the characteristics of context elements, which makes it possible to agree on a hierarchy of context elements and the hierarchy of the reporting process. The method of presentation summarizing the knowledge-intensive business process based on an analysis of its logs. The method includes the steps of selection provided in the log subsets of attributes of objects for each level of detail and the process of consolidating actions with selected attributes. During the generalization of the chain together a sequence of actions, select the blocks / association of election results, as well as hiding the private implementation. This makes it possible to integrate functional and process approaches to management within an enterprise. The method allows increasing the efficiency of the mining process through the provision of essential for analysis and decisioning making subsets of action that creates the conditions for solving the problem of spaghetti-like workflow -models

Ref.: 12 items.

УДК 681.518:004.93.1'

Визначення лінії симетрії на зображенні обличчя людини при діагностуванні емоційно-психічного стану / І.В. Шелехов, Д.В. Прилепа // *Біоніка інтелекту: наук.-техн. журнал.* — 2016. — № 2 (87). — С. 106–110.

Розглянуто інформаційно-екстремальний метод визначення лінії симетрії портрету людини при діагностуванні її емоційно-психічного стану. Ідея методу полягає в виборі лінії симетрії право- та лівополушарних портретів шляхом

мінімізації в процесі машинного навчання ентропійного критерія, що є мірою різноманітності. Розроблений та програмно реалізований метод автоматичного пошуку ліній симетрії портрета людини забезпечив підвищення достовірності розпізнавання емоційно-психічного стану людини в порівнянні з візуальним вибором.

Ил. 6. Бібліогр.: 10 найм.

УДК 681.518:004.93.1'

Определение линии симметрии на изображении лица человека при диагностировании эмоционально-психического состояния / И.В. Шелехов, Д.В. Прилепа // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 106–110.

Рассмотрено информационно-экстремальный метод определения линии симметрии портрета человека при диагностировании его эмоционально-психического состояния. Идея метода состоит в выборе линии симметрии право и левополушарных портретов путём минимизации в процессе машинного обучения энтропийного критерия, являющегося мерой разнообразия. Разработанный и программно реализованный метод автоматического поиска линий симметрии портрета человека обеспечил повышение достоверности распознавания эмоционально-психического состояния человека по сравнению с визуальным выбором.

Ил. 6. Библиогр.: 10 наим.

UDK 681.518:004.93.1'

Definition symmetry line in the image of a human face in diagnosing emotional and mental state / I.V. Shelekhov, D.V. Prylepa // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 106–110.

Information-extreme method of determining the line of symmetry of a person portrait in the diagnosis of emotional and mental state are considered. The idea of the method is to select the line of symmetry right- and the left-hemisphere symmetry portraits by minimization the process of machine learning entropy criterion, which is a diversity measure. Development and implementation of method for automatic determining the lines of symmetry of a portrait provided to increase reliability of detection of emotional and mental state of a person in contrast with a visual selection.

Fig. 6. Ref.: 10 items.

УДК: 004.93.1'

Інформаційно-аналітична система адаптації учбового контенту випускаючої кафедри до ринку праці / А.В. Васильєв, А.С. Довбиш, Є.С. Кулик, З.В. Козлов // Біоніка інтелекту: наук.-техн. журнал. — 2016. — № 2 (87). — С. 111–116.

В статті запропоновано інформаційно-екстремальний алгоритм машинного навчання інформаційно-аналітичної системи адаптації учбового контенту випускаючої кафедри до вимог ринку праці. Описано процедуру машинного навчання системи з оптимізацією геометричних параметрів гіперсферичних вкладених контейнерів класів розпізнавання, що характеризують рівні якості учбового контенту. Наведено приклад реалізації запропонованого алгоритму з метою побудови вирішальних правил для класифікації відповідей респондентів, що є провідними спеціалістами в області інформаційних технологій.

Ил. 4. Бібліогр.: 7 найм.

УДК: 004.93.1'

Информационно-аналитическая система адаптации учебного контента выпускающей кафедры к требованиям рынка труда / А.В. Васильев, А.С. Довбыш, Е.С. Кулик, З.В. Козлов // Бионика интеллекта: научн.-техн. журнал. — 2016. — № 2 (87). — С. 111–116.

В статье предложен информационно-экстремальный алгоритм машинного обучения информационно-аналитической системы адаптации учебного контента выпускающей кафедры к требованиям рынка труда. Описана процедура машинного обучения системы с оптимизацией геометрических параметров гипersферических вложенных контейнеров классов распознавания, характеризующих уровни качества учебного контента. Приведен пример реализации предложенного алгоритма с целью построения решающих правил для классификации ответов респондентов, являющихся ведущими специалистами в области информационных технологий.

Ил. 4. Библиогр.: 7 наим.

UDK: 004.93.1'

Quality assessment of educational content in graduating departments using information-extreme intellectual technology of data analysis / A.V. Vasylyev, A.S. Dovbysh, E.S. Kulik, Z.V. Kozlov // Bionics of intelligence: Sci. mag. — 2016. — № 2 (87). — P. 111–116.

This article proposes informational and extreme algorithm for machine learning of informational and analytical system for adapting educational content of graduating department to the requirements of the labour market. Article describes machine learning procedure of a system with geometrical parameters optimization of hyperspherical nested containers for recognition classes, which describe levels of educational content quality. The example of implementation of suggested algorithm is offered to build decision rules for classification of respondent's answers, who are leading specialists in IT.

Fig. 4. Ref.: 7 items.

ОБ АВТОРАХ

Айвас Елена Юрьевна	12	аспирантка кафедры теории и практики перевода Харьковского гуманитарного университета «Народная украинская академия»
Бабий Андрей Степанович	30	старший преподаватель, аспирант кафедры программной инженерии Харьковского национального университета радиоэлектроники
Васильев Анатолий Васильевич	111	канд. техн. наук, профессор, ректор Сумского государственного университета
Гороховатский Владимир Алексеевич	83	д-р техн. наук, профессор кафедры информационных технологий и высшей математики Харьковского института ГВУЗ «Университет банковского дела»
Довбыш Анатолий Степанович	111	д-р техн. наук, профессор, заведующий кафедрой компьютерных наук Сумского государственного университета
Ерохин Андрей Леонидович	30	д-р техн. наук, декан факультета компьютерных наук, профессор кафедры программной инженерии Харьковского национального университета радиоэлектроники
Завизиступ Юрий Юрьевич	35	канд. техн. наук, проф. каф. ЭВМ Харьковского национального университета радиоэлектроники
Каграманян Александр Георгиевич	48, 53	канд. техн. наук, доцент кафедры естественных наук Харьковского национального университета им. В.Н. Каразина
Кобылин Олег Анатольевич	94	канд. техн. наук, доцент кафедры информатики Харьковского национального университета радиоэлектроники
Козлов Захар Викторович	111	аспирант кафедры компьютерных наук Сумского государственного университета
Кулик Евгения Сергеевна	111	аспирант кафедры компьютерных наук Сумского государственного университета
Кушвид Евгений Сергеевич	16	студент кафедры искусственного интеллекта Харьковского национального университета радиоэлектроники
Лазаренко Ольга Владимировна	12	канд. техн. наук, доцент, заведующая лабораторией когнитивной лингвистики Харьковского гуманитарного университета «Народная украинская академия»
Лебедев Олег Григорьевич	69	канд. техн. наук, доцент кафедры ЭВМ Харьковского национального университета радиоэлектроники
Левыкин Игорь Викторович	64	канд. техн. наук, профессор кафедры медиасистем и технологий Харьковского национального университета радиоэлектроники
Лесной Георгий Дмитриевич	88	д-р геолог. наук, доцент кафедры геофизики Киевского университета имени Тараса Шевченко, заместитель директора ООО «Тутковский интегрированные решения»
Литвин Олег Николаевич	77, 88	д-р физ.-мат. наук, профессор, заведующий кафедрой высшей и прикладной математики Украинской инженерно-педагогической академии
Литвин Олег Олегович	88	канд. физ.-мат. наук, доцент, декан Технологического факультета Украинской инженерно-педагогической академии
Ляшенко Вячеслав Викторович	94	научный сотрудник кафедры информатики Харьковского национального университета радиоэлектроники

Михаль Олег Филиппович	35, 42, 69	д-р техн. наук, профессор кафедры ЭВМ Харьковского национального университета радиоэлектроники
Нечипоренко Алина Сергеевна	30	канд. техн. наук, доцент кафедры биомедицинской инженерии, докторант кафедры программной инженерии Харьковского национального университета радиоэлектроники
Панченко Дмитрий Игоревич	12	канд. филол. наук, доцент кафедры теории и практики перевода Харьковского гуманитарного университета «Народная украинская академия»
Пузик Алексей Сергеевич	24	аспирант кафедры программной инженерии Харьковского национального университета радиоэлектроники
Прилепа Дмитрий Викторович	106	ведущий специалист кафедры компьютерных наук Сумского государственного университета
Свиридов Артем Сергеевич	35	аспирант кафедры ЭВМ Харьковского национального университета радиоэлектроники
Славик Алексей Валериевич	88	аспирант кафедры высшей и прикладной математики Украинской инженерно-педагогической академии
Столяров Виталий Сергеевич	83	студент Харьковского национального университета радиоэлектроники
Томич Станислав Николаевич	94	магистр кафедры программной инженерии Харьковского национального университета радиоэлектроники
Турута Алексей Петрович	30	канд. техн. наук, доцент кафедры программной инженерии Харьковского национального университета радиоэлектроники
Удовенко Сергей Григорьевич	16	д-р техн. наук, профессор, заведующий кафедрой информатики и компьютерной техники Харьковского национального экономического университета им. С. Кузнеца
Чалая Лариса Эрнестовна	16	канд. техн. наук, доцент кафедры искусственного интеллекта Харьковского национального университета радиоэлектроники
Чалая Оксана Викторовна	101	канд. эконом. наук, доцент кафедры информационных управляющих систем Харьковского национального университета радиоэлектроники
Черная Татьяна Ивановна	77	старший преподаватель Украинской инженерно-педагогической академии, г. Харьков
Четвериков Григорий Григорьевич	48, 53	д-р техн. наук, профессор кафедры программной инженерии Харьковского национального университета радиоэлектроники
Чумаченко	59	
Шелехов Игорь Владимирович	106	канд. техн. наук, доцент кафедры компьютерных наук Сумского государственного университета
Широков Владимир Анатольевич	3	академик НАН Украины, директор Украинского языково-информационного фонда НАН Украины, г. Киев
Шляхов Владислав Викторович	48, 53	канд. техн. наук, доцент, ведущий научный сотрудник кафедры информатики Харьковского национального университета радиоэлектроники
Яковлев	59	
Ярмош Елена Виталиевна	77	канд. физ.-мат. наук, доцент Украинской инженерно-педагогической академии, г. Харьков

ПРАВИЛА оформлення рукописів для авторів науково-технічного журналу «БІОНІКА ІНТЕЛЕКТУ»

Науково-технічний журнал «Біоніка інтелекту» приймає до друку написані спеціально для нього оригінальні рукописи, які раніше ніде не друкувались. Структура рукопису повинна бути такою: індекс УДК, заголовок, відомості про авторів, анотація, ключові слова, вступ, основний текст статті, висновки, список використаної літератури.

Відповідно до Постанови ВАК України від 15.01.2003 №7-05/1 (Бюлетень ВАК, №1, 2003, с. 2) стаття повинна мати такі необхідні елементи: постановка проблеми в загальному вигляді та її зв'язок із важливими науковими чи практичними завданнями; аналіз останніх досліджень і публікацій і виділення не вирішених раніше частин загальної проблеми в даній області; формулювання цілей та завдань дослідження; виклад основного матеріалу досліджень з повним обґрунтуванням отриманих наукових результатів; висновки з даного дослідження та перспективи подальших досліджень у даному напрямку.

Статті мають бути виконані в редакторі Microsoft Word. Формат сторінки — A4 (210x297 мм), поля: верхнє — 25 мм, нижнє — 20 мм, ліве, праве — 17 мм. Кількість колонок — 2, з інтервалом між ними 5 мм, основний шрифт Times New Roman, кегль основного тексту — 10 пунктів, міжрядковий інтервал — множник (1,1), абзацний відступ — 6 мм. Обсяг рукопису — від 4 до 12 сторінок (мови: російська, українська, англійська).

УДК друкується з першого рядка, без відступів, вирівнювання по лівому краю.

Назва статті друкується прописними літерами; шрифт прямий, напівжирний, кегль 12. *Назви розділів* нумерують арабськими цифрами, виділяють жирним шрифтом. Відступи для назви статті, ініціалів та прізвищ авторів, відомостей про авторів, назв розділів, вступу та висновків, списку літератури: зверху — 6 пт, знизу — 3 пт.

Анотацію (мовою статті, абзац 4-10 рядків, кегль 9) розміщують на початку статті, в ній має бути розміщена інформація про результати описаних досліджень.

Ключові слова (4-10 слів з тексту статті, які з точки зору інформаційного пошуку несуть змістовне навантаження) наводять мовою рукопису, через кому в називному відмінку, кегль 9.

Рисунки та таблиці (чорно-білі, контрастні) розміщуються у тексті після першого посилання у вигляді окремих об'єктів і нумерують арабськими цифрами наскрізною нумерацією за наявності більше ніж одного об'єкта. Невеликі схеми, що складаються з 3-4 елементів виконують, використовуючи вставку об'єкта Рисунок Microsoft Word. Більш складні виконують у графічних редакторах у вигляді чорно-білих графічних файлів форматів .tiff, .jpg, .wmf, .cdr із розділенням 300 dpi. Рисунки мають міститися у текстовому файлі й обов'язково

подаватися окремим файлом з відповідною назвою (наприклад, Рис.1.cdr).

Усі елементи рисунка, включаючи написи, повинні бути згруповані. Усі написи в рисунках і таблицях мають бути виконані шрифтом Times New Roman, кегль у рисунках — 10, у таблицях — 9.

Рисунок повинен мати центрований підпис (поза малюнком), шрифт 9, відступи зверху і знизу по 6 пт. Ширина рисунка має відповідати ширині колонки (або ширині сторінки).

Формули, символи, змінні, повинні бути набрані в редакторі формул MathType або Microsoft Equation. Формули розміщують посередині рядка й нумерують за наявності посилань на них у рукописі. Шрифт — Times New Roman. Висота змінної — 10 пунктів, великих і малих індексів — 8 пт, основний математичний символ — 12 (10) пт. Змінні, позначені латинськими літерами, набирають курсивом, грецькі літери, скорочення російських слів і цифри — прямим написанням. Змінні, які є в тексті, також набирають у редакторі формул.

Список літератури вміщує опубліковані джерела, на які є посилання в тексті, укладені у квадратні дужки, друкують без абзацного відступу, кегль 9 пт, відступ зверху — 6 пт.

Після списку літератури з відступом зверху 6 пт зазначають дату подання статті до редколегії. Число та місяць задають двозначними числами через крапку. Розмір шрифта — 9 пт, курсив, вирівнювання по правому краю.

Реферати (Times New Roman, кегль — 9 пунктів, 3-4 речення) подають російською та англійською мовами. Реферат не повинен дублювати текст анотації.

Разом із рукописом (на аркушах білого паперу формату A4 щільністю 80-90 г/м², надрукованим на лазерному принтері, у 2-х примірниках) необхідно подати такі документи:

1. Заяву, яку повинні підписати всі автори.
2. Акт експертизи про можливість опублікування матеріалів у відкритому друці.
3. Рецензію, підписану доктором наук.
4. Відомості про авторів.
5. Електронний варіант рукопису, реферату та відомостей про авторів.
6. Оплату за публікацію.

Необхідно також зазначити один з наступних тематичних розділів, якому відповідає рукопис:

1. Теоретичні основи інформатики та кібернетики. Теорія інтелекту
2. Математичне моделювання. Системний аналіз. Прийняття рішень
3. Інтелектуальна обробка інформації. Розпізнавання образів
4. Інформаційні технології та програмно-технічні комплекси
5. Структурна, прикладна та математична лінгвістика
6. Дискусійні повідомлення

INSTRUCTIONS for authors of manuscripts of the scientific journal «BIONICS OF INTELLIGENCE»

The scientific journal "Bionics of intelligence" accepts for publication original manuscripts which have not been published earlier. The manuscript structure should be as follows: Universal Decimal Classification (UDC) title, authors' initials and surname (in alphabetical order), abstract, key words, introduction, main text, conclusions, references.

According to the Editorial board resolution, based on the Presidium Convention of Ukraine's Supreme Attestation Committee of 15.01.2003 №7-05/1 (Bulletin of Supreme Attestation Committee, №1, 2003, p. 2) manuscripts must have the following required elements: introduction (general statement of a problem and its relation to important scientific and practical tasks; analysis of recent research, publications and highlighting of unsolved parts of the general problem in the given field); formulating aims and tasks of research; presentations of the main research material with full substantiation of scientific results obtained; conclusions and perspectives of further research in the given field.

Manuscripts should be submitted in Microsoft Word. Page format - A4 (210x297mm), margins: top - 25mm, bottom - 20mm; left, right - 17mm. Double column format with 5mm spacing, font - Times New Roman, font size - 10 points, line spacing - multiplier (1,1), indentation - 6mm. The manuscript should be from 4 to 12 pages (languages: Russian, Ukrainian, English).

The UDC is published from the first line, without indentation, the alignment is by a left edge. The title is in capital letters; the type is medium bold-faced Roman; type size 12. The names of sections are of extra bold type and numbered in Arabic figures. There are indentions for the names of manuscripts, initials and surnames of authors, information about authors, the names of sections, introduction and conclusions, references: top - 6 pt; bottom - 3 pt.

An abstract (in the language of a manuscript, an indentation is made up of 4-10 lines; type 9) is in the beginning of an article and contains information about the results of described studies.

Key words (4-10 words from the text of an article, which from the point of view of information search bear sense in the language of a manuscript, by way of a comma in nominative case, type 9).

Figures and tables (black-and-white, sharp and of good quality) should be in a text after a first reference in the form of embedded item and numbered separately by Arabic numerals in case of more than one item. All legends of figures and tables, including inscriptions, must be grouped. All inscriptions in figures and tables must be in Times New Roman, font size in figures - 10, in

tables - 9. A table title is to the right above the table (font size - 9). The figure should contain a centered figure legend (outside a figure), font size 9, in the centre, top and bottom indentions - 6pt. The figure width must agree with the column width (or page width).

Equations, symbols, variables should be submitted in Math Type (Equation). Equations are centered and numbered in case of references in the text. The font - Times New Roman. The size of variable - 10 points, superscript and subscript characters - 8 pt, a main math. symbol - 12 (10) pt. Variables, designated by Latin letters, should be italicized; Greek letters, abbreviations of Russian words and figures should be set in Roman type. Variables which are in the text are also submitted in Math Type (Equation).

References, submitted to the state standards, include published sources that are referred to in the main text in square brackets, without an indentations, 9pt., top indentation - 6 pt.

The date of receiving an article by the Editorial board is designated after the references with top indentations - 6 pt. Date and month should be given in numbers by way of a full stop. The font size - 9 pt, italic type, alignment should be done on the right edge.

Abstracts should be submitted in two languages: Ukrainian and Russian (Times New Roman, 9 pt, 3-4 sentences). The text of a resume must not duplicate an abstract.

The following documents must be submitted together with a manuscript:

1. An application of the following form signed by all the authors:

"You are kindly requested to accept the paper (authors' full names and the name of an paper should be indicated) in pages (the number of pages should be indicated) for publication in the scientific journal "Bionics of intelligence". We guarantee the payment.

Information about the authors (surname, first name and patronimic of each authors, place of work, degree, academic status, contact telephone, mailing and electronic addresses should be indicated).

Signatures of authors".

2. The text of a manuscript on A4 format white color sheets of 80-90gr/m2 density typed on a laser printer.

3. A certificate of expertise about a possibility of having the materials published in the press.

4. A review signed by a doctor of sciences.

5. Information about the authors.

6. An electronic variant of a manuscript, an abstract and information about the authors (on a 3.5" diskette or by electronic mail).

7. A receipt of payment for publication.

АЛГЕБРО-ЛОГІЧНІ ЗАСОБИ МОДЕЛЮВАННЯ ПРИРОДНОЇ МОВИ

Проведено аналіз алгебро-логічної структури природної мови. Розглянуто концептуально-методолічний підхід до мови людини, що дозволяє сприймати її як деяку алгебру, а її тексти — як формули цієї алгебри.

МОВА ПРИРОДНА, АЛГЕБРА ПРЕДИКАТИВ, ВІДНОШЕННЯ, АЛГЕБРА ПРЕДИКАТНИХ ОПЕРАЦІЙ

Вступ

Формальним моделям семантико-синтаксичних структур мови відводиться вирішальна роль у сучасній проблематиці комп'ютерної лінгвістики та системах штучного інтелекту (ШтІ). Це пов'язано з необхідністю створення програмно-апаратного комплексу генерації та аналізу речень природної мови (ПМ).

1. Дослідження алгебро-логічної структури природної мови

У роботі використовується апарат алгебри предикатів [1]. Множина U може бути як скінченною, так нескінченною. У першому випадку простір U^m називатимемо скінченним, а в іншому — нескінченним.

$$P(x_1, x_2, \dots, x_n) = \begin{cases} 0, & \text{якщо } (x_1, x_2, \dots, x_n) \notin T \\ 1, & \text{якщо } (x_1, x_2, \dots, x_n) \in T. \end{cases} \quad (1)$$

Згідно з (1) можливий перехід від будь-якого відношення T до відповідного йому предикату P . Предикат P , що знаходимо по (1), називатимемо характеристичною функцією відношення T .

2. Шляхи автоматизації обробки мовної інформації

У даний час в системах штучного інтелекту машинний словник та комплекс програм (тезауруси) використовуються, як правило, для виконання будь-якої однієї функції.

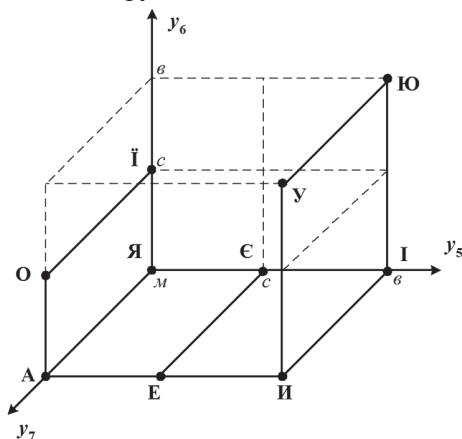


Рис. 1. Тривимірний простір ознак для голосних фонем

Висновки

У статті запропонована узагальнену структуру інтелектуальної системи, яка відповідає новій інформаційній технології рішення задач на ЕОМ, що орієнтовані на досягнення високорівневої технології обробки мовної інформації (отримання нової якості). Істотно новим в роботі є розширення алгебри скінченних предикатів (АСП). Тепер вона охоплює не тільки скінченні предикати, а також — нескінченні. Тепер область її рекомендованого застосування розширена та охоплює довільні відношення, які далі будемо описувати за допомогою ДКАП.

Список літератури:

Надійшла до редколегії 15.02.2012

УДК 519.62

Алгебро-логические средства моделирования естественного языка / Г.Г. Четвериков, И.Д. Вечирская // Бионика интеллекта: науч.-техн. журнал. — 2007. — № 1 (66). — С. 00-00.

В статье рассматриваются перспективные направления развития современных цифровых устройств, сетей и систем. Утверждается, что развитие средств вычислительной техники является основой автоматизации умственной деятельности человека.

Ил. 5. Библиогр.: 7 назв.

UDK 519.7

Algebra-logical tools of modeling natural language / G.G. Chetverikov, I.D. Vechirskaya // Bionics of Intelligence: Sci. Mag. — 2007. — № 1 (66). — С. 00-00.

In article the perspective directions of modern digital devices, networks and systems development are considered. The carried out analysis shows means of computer facilities development is a baseline of automation of the man intellectual activity.

Fig. 5. Ref.: 7 items.

Видавництво здійснює остаточне форматування тексту відповідно до вимог друку.

Адреса редакції:

Україна, 61166, м.Харків, пр. Леніна 14, ХНУРЕ
к.127, тел. 702-14-77, факс 702-10-13,
e-mail: ira_se@list.ru

Наукове видання

БІОНІКА ІНТЕЛЕКТУ
інформація, мова, інтелект

Науково-технічний журнал

№ 2 (87)

2016

Головний редактор — *Ю.П. Шабанов-Кушнарєнко*

Зам. головного редактора — *Г. Г. Четвериков*

Відповідальний редактор — *І. Д. Вечірська*

Комп'ютерна верстка — *О. Б. Ісаєва*

Рекомендовано Вченою Радою
Харківського національного університету радіоелектроніки
(протокол № 7 від 03.07.2015 р.)

Адреса редакції:

Україна, 61166, Харків-166, просп. Леніна, 14,
Харківський національний університет радіоелектроніки, к. 127
тел. 702-14-77, факс 702-10-13,
e-mail: ira_se@list.ru

Підписано до друку 23.12.2016. Формат 60 x 84 ¹/₈. Друк ризографічний.
Папір офсетний. Гарнітура Newton. Умов. друк. арк. 15,1. Обл.-вид. арк. 15,0.
Тираж 100 прим. Зам. № .

Надруковано в навчально-науковому видавничо-поліграфічному центрі ХНУРЕ
61166, Харків-166, просп. Леніна, 14