

ВЕСТНИК

НАЦИОНАЛЬНОГО ТЕХНИЧЕСКОГО УНИВЕРСИТЕТА "ХПИ"

Сборник научных трудов

Тематический выпуск

"Информатика и моделирование", № 16

17'2011

Издание основано Национальным техническим университетом "Харьковский политехнический институт" в 2001 году

Государственное издание

Свидетельство Госкомитета по
информационной политике Украины
КВ № 5256 от 2 июля 2001 года

КООРДИНАЦИОННЫЙ СОВЕТ:

Председатель

Л.Л. Товажнянский, д-р техн. наук, проф.

Секретарь координационного совета

К.А. Горбунов, канд. техн. наук, доц.

А.П. Марченко, д-р техн. наук, проф.

Е.И. Сокол, д-р техн. наук, проф.

Е.Е. Александров, д-р техн. наук, проф.;

Л.М. Бесов, д-р ист. наук, проф.;

А.В. Бойко, д-р техн. наук, проф.;

Ф.Ф. Гладкий, д-р техн. наук, проф.;

М.Д. Годлевский, д-р техн. наук, проф.;

А.И. Грабченко, д-р техн. наук, проф.;

В.Г. Данько, д-р техн. наук, проф.;

В.Д. Дмитриенко, д-р техн. наук, проф.;

И.Ф. Домнин, д-р техн. наук, проф.;

В.В. Епифанов, канд. техн. наук, проф.;

Ю.И. Зайцев, канд. техн. наук, проф.;

П.А. Качанов, д-р техн. наук, проф.;

В.Б. Клепиков, д-р техн. наук, проф.;

С.И. Кондрашов, д-р техн. наук, проф.;

В.М. Кошельник, д-р техн. наук, проф.;

В.И. Кравченко, д-р техн. наук, проф.;

Г.В. Лисачук, д-р техн. наук, проф.;

В.С. Лупиков, д-р техн. наук, проф.;

О.К. Морачковский, д-р техн. наук, проф.;

П.Г. Перерва, д-р экон. наук, проф.;

Е.В. Рогожкин, д-р техн. наук, проф.;

М.И. Рыщенко, д-р техн. наук, проф.;

В.Б. Самородов, д-р техн. наук, проф.;

Г.М. Сучков, д-р техн. наук, проф.;

Ю.В. Тимофеев, д-р техн. наук, проф.;

М.А. Ткачук, д-р техн. наук, проф.

РЕДАКЦИОННАЯ КОЛЛЕГИЯ:

Ответственный редактор:

В.Д. Дмитриенко, д-р техн. наук, проф.

Ответственный секретарь:

С.Ю. Леонов, канд. техн. наук, доц.

А.Г. Гурин, д-р техн. наук, проф.;

Л.В. Дербунович, д-р техн. наук, проф.;

Е.Г. Жилияков, д-р техн. наук, проф.;

П.А. Качанов, д-р техн. наук, проф.;

Н.И. Корсунов, д-р техн. наук, проф.;

Г.М. Сучков, д-р техн. наук, проф.;

И.И. Обод, д-р техн. наук, проф.;

А.И. Овчаренко, д-р техн. наук, проф.;

А.А. Серков, д-р техн. наук, проф.

Адрес редколлегии: 61002, Харьков,
ул. Фрунзе, 21, НТУ "ХПИ".

Каф. ВТП, тел. (057)-707-61-65

Харьков 2011

Вісник Національного технічного університету "Харківський політехнічний інститут". Збірник наукових праць. Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ", 2011. – № 17. – 197 с.

В збірнику представлені теоретичні та практичні результати наукових досліджень та розробок, що виконані викладачами вищої школи, аспірантами, науковими співробітниками різних організацій та установ.

Для викладачів, наукових співробітників, спеціалістів.

В сборнике представлены теоретические и практические результаты исследований и разработок, выполненных преподавателями высшей школы, аспирантами, научными сотрудниками различных организаций и предприятий.

Для преподавателей, научных сотрудников, специалистов.

Вісник Національного технічного університету "ХПІ" внесено до "Переліку № 9 наукових фахових видань України, в яких можуть публікуватися результати дисертаційних робіт на здобуття наукових ступенів доктора і кандидата наук", затвердженого постановою президії ВАК України від 14 листопада 2001 року, № 2 – 05/9. (Бюлетень ВАК України № 6, 2001 р., технічні науки, збірники наукових праць, № 2) та "Переліку наукових фахових видань України, в яких можуть публікуватися результати дисертаційних робіт на здобуття наукових ступенів доктора і кандидата наук", затвердженого постановою президії ВАК України від 26 травня 2010 р. № 1 – 05/4. (Бюлетень ВАК України № 6, 2010 р., стор. 3, № 20).

**Рекомендовано до друку Вченою радою НТУ "ХПІ"
Протокол № 5 від 27 травня 2011 р.**

© Національний технічний університет "ХПІ"

ISSN 2079-0031



50-летию

*КАФЕДРЫ ВЫЧИСЛИТЕЛЬНОЙ ТЕХНИКИ
И ПРОГРАММИРОВАНИЯ*

*НАЦИОНАЛЬНОГО ТЕХНИЧЕСКОГО
УНИВЕРСИТЕТА*

"ХАРЬКОВСКИЙ ПОЛИТЕХНИЧЕСКИЙ ИНСТИТУТ"

ПОСВЯЩАЕТСЯ

А.М. АХМЕТШИН, д.ф.-м.н., проф. ДНУ, Днепропетровск,
А.А. СТЕПАНЕНКО, к.т.н., ст. преп. ЗНТУ, Запорожье

ИНТЕРФЕРЕНЦИОННЫЙ МЕТОД АНАЛИЗА РАДИОЛОГИЧЕСКИХ ИЗОБРАЖЕНИЙ В ПРОСТРАНСТВЕ ПРИЗНАКОВ ПРЕОБРАЗОВАНИЯ РАДОНА

Рассмотрены информационные возможности нового метода качественного анализа слабоконтрастных радиологических медицинских изображений в комплексном пространстве изображений. Метод позволяет сегментировать визуально неопределенные потенциально информативные участки путем использования фазо-пространственных характеристик. Ил.: 3. Библиогр.: 8 назв.

Ключевые слова: слабоконтрастное изображение, комплексное пространство изображений, сегментировать, информативные участки, фазо-пространственные характеристики.

Постановка проблемы. Одна из актуальных проблем вычислительного интеллекта связана с задачами анализа слабоконтрастных изображений (медицина, дистанционное зондирование, геофизика, неразрушающий контроль). Суть проблемы заключается в том, чтобы сделать визуально неразличимые участки (детали) видимыми. Традиционные методы цифровой обработки изображений (эквализация гистограмм, градиентное отображение и т.д.) зачастую не приводят к успеху, если априори неизвестный объект интереса расположен на существенно неоднородном яркостном фоне. Одно из современных направлений, связанных с анализом слабоконтрастных изображений, базируется на проведении аналогий с наиболее чувствительными физическими методами оптических и радиофизических измерений (интерферометрия, голография, эллипсометрия) [1 – 3].

Основная проблема подобного направления – разработка адекватных физико–математических моделей анализируемых изображений, допускающих использование математического формализма методов физических когерентно–оптических измерений.

Анализ литературы. В литературе, посвященной анализу современных методов цифровой обработки изображений [4, 5], основное внимание уделяется методам повышения качества изображений с точки зрения их психофизиологического восприятия. В частности, в указанных книгах нет ни малейшего намека, например, на оптические голографические методы неразрушающего контроля [6], обеспечивающие обнаружение участков конструкций потенциально

подверженных критическим деформациям. Такая ситуация, по нашему мнению, обусловлена тем обстоятельством, что в голографической интерферометрии анализируется не само изображение тестируемой конструкции, а его интерферограмма, что, естественно, выдвигает особые требования к ее расшифровке (интерпретации).

В работах [7, 8] были рассмотрены информационные возможности виртуального цифрового интерференционного метода повышения визуального качества слабоконтрастных изображений. Идея метода базируется на возможности проведения формальной аналогии с методом фазоконтрастной оптической микроскопии при переходе в пространство комплексных яркостей. Подобный подход действительно повышает контрастность изображений, но не подходит для решения задач сегментации участков слабоконтрастных изображений с "размытыми" границами (потенциально патологические участки на медицинских радиологических изображениях).

Цель статьи – демонстрация информационных возможностей нового метода сегментации визуально неразличимых участков слабоконтрастных медицинских радиологических изображений на основе сочетания двух этапов: преобразование Радона исходного изображения и виртуального синтеза соответствующей ему интерферограммы с интерпретацией результатов по ее фазо–пространственной характеристике.

Основная часть. На рис. 1 представлено изображение томограммы головного мозга до и после введения рентгеноконтрастного вещества и соответствующие им результаты повышения их контраста на основе применения стандартного метода эквализации их гистограмм.

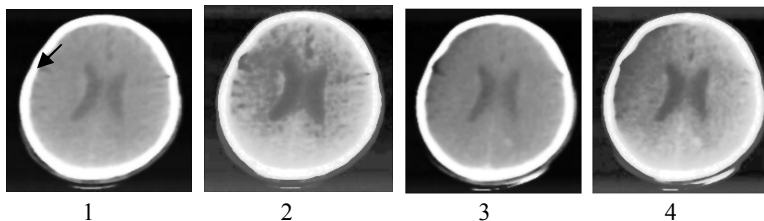


Рис.1. Томограмма головного мозга до (1) и после (3) введения рентгеноконтрастного вещества и соответствующие им изображения (2, 4) результатов применения метода эквализации их гистограмм

На рис.1 видна гематома (указана стрелкой), но совершенно невидима область ее "скрытого" влияния. Введение

рентгеноконтрастного вещества (рис.1.3) не улучшает ситуации. Повышение контраста методом эквализации гистограммы несколько улучшает ситуацию (рис.1.4), однако с точки зрения компьютерного зрения, результат является убедительным.

Текущая тенденция анализа, например, медицинских радиологических изображений базируется на использовании нескольких различных методов цифровой обработки в целях повышения достоверности обнаружения потенциальных патологических участков. Однако такой подход, с точки зрения синергетического принципа, не всегда дает ожидаемые результаты (т.е. 1+1 может быть и меньше двух). В этой связи зачастую может быть целесообразен подход, базирующийся на следующем принципе: переход в новый информационный базис с последующим анализом нового синтезированного изображения на основе методов вычислительного интеллекта уже в этом базисе. В качестве варианта перехода к новому информационному базису в данной работе предлагается использовать преобразование Радона исходного анализируемого изображения $I(x, y)$, базирующееся на вычислении проекций изображения вдоль определенных направлений (углов). Проекция функции $I(x, y)$ на ось x' представляет собой линейный интеграл

$$R_{\theta}(x') = \int_{-\infty}^{\infty} I(x' \cos \theta - y' \sin \theta, x' \sin \theta + y' \cos \theta) dy', \quad (1)$$

где оси x' и y' задаются поворотом на угол θ против часовой стрелки

$$\begin{bmatrix} x' \\ y' \end{bmatrix} = \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix} * \begin{bmatrix} x \\ y \end{bmatrix}. \quad (2)$$

Исходное полутоновое изображение рассматривается как двумерная функция. Таким образом, использование преобразования Радона обеспечивает переход к новому информационному базису, где ось "x" соответствует числу углов проецирования исходного изображения, а ось "y" – соответствует яркостным значениям проекций $R_{\theta}(x')$. С физической точки зрения, использование преобразования Радона позволяет "накопить" значения визуально неразличимых участков, в целях облегчения (упрощения) процедуры их последующей идентификации. В каком-то смысле, эта операция близка к процедуре "усреднения", используемой в области цифровой обработки сигналов для выделения сигналов неизвестной формы на фоне аддитивных измерительных шумов в условиях, когда отношение сигнал/шум намного

меньше единицы.

Структура алгоритма базируется на использовании следующих этапов.

1. Преобразование Радона исходного изображения $I(x, y) \Rightarrow R_0(x')$.

2. Модуляционное преобразование синтезированного изображения $R_0(x') \Rightarrow H(x, y)$ на основе использования выражения (3)

$$H(x, y) = \exp(j\pi / [\Psi[I(x, y)] + \gamma]) , \quad (3)$$

где Ψ – оператор предварительной трансформации исходного изображения, γ – стабилизирующий параметр. Такое преобразование позволяет провести формальную аналогию с интерференционным методом фазоконтрастной микроскопии, при введении виртуального когерентного опорного поля (волны), что открывает возможность визуализации фазовых "неоднородностей" анализируемого изображения. К "фазовым неоднородностям" относятся участки изображения, где перепад яркостей находится в пределах 2%, что не позволяет осуществить процедуру из непосредственной визуальной идентификации.

3. Введение виртуального когерентного опорного поля $B(x, y)$.

4. Вычисление модуля и аргумента векторной суммы $S(\theta, r) = R_0(\theta) + B(x, y)$, рассматриваемых в качестве новых информативных синтезированных изображений (рис. 2).

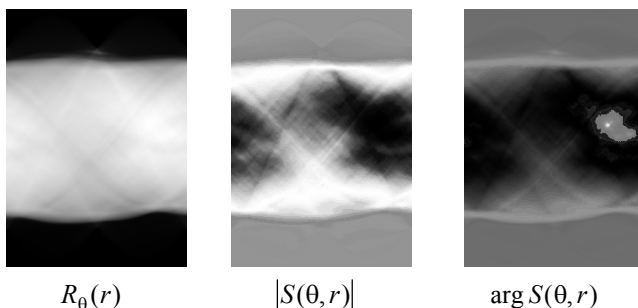


Рис. 2. Результат применения нового метода для томограммы на рис. 1.1

Из рассмотрения рис. 2 видно, что на исходной характеристике $R_0(\theta)$ участок скрытого влияния визуально не идентифицируется, тогда как на фазо–пространственной интерференционной характеристике

преобразования Радона $\arg S(\theta, r)$ она выделяется однозначно.

Из рассмотрения рис. 2 можно заключить, что не было необходимости введения рентгеноконтрастного вещества и использования повторной томографии (двойная доза облучения), поскольку информативный участок томограммы однозначно идентифицируется на характеристике $\arg S(\theta, r)$ (рис. 2).

Конечно, восприятие результатов сегментации томограммы в пространстве признаков преобразования Радона требует определенной психологической "настройки" в силу необычности самого подхода. В этой связи возникает естественный вопрос: возможно ли вернуться из пространства преобразования Радона в обычное евклидово пространство

$\arg S(\theta, r) \xRightarrow{R^{-1}} \hat{I}(x, y)$, где R^{-1} – оператор обратного преобразования Радона, $\hat{I}(x, y)$ – трансформированное изображение в евклидовом пространстве обычных яркостей, т.е.

$$\hat{I}(x, y) = \int_0^{\pi} \int_{-\infty}^{\infty} \arg S(\theta, r) \exp(-j2\pi r(x \cos(\theta) + y \sin(\theta))) |r| dr. \quad (4)$$

К сожалению, обратное преобразование Радона относится к области некорректных задач математической физики, поэтому любые математические преобразования над результатами прямого преобразования Радона (например, синтез $|S(\theta, r)|$ и $\arg S(\theta, r)$) приводят

к тому, что обратное преобразование (т.е. $\hat{I}(x, y)$) будет иметь искаженную форму (рис. 3).

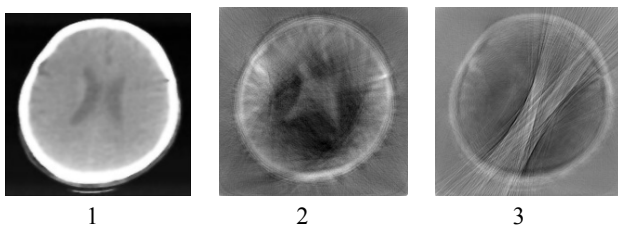


Рис. 3. Результаты обратного преобразования Радона от данных на рис. 2: 1 – $R_0(r)$; 2 – $|S(\theta, r)|$; 3 – $\arg S(\theta, r)$

Выводы. В результате проделанной работы можно заключить:

1. Предложен новый виртуальный физический метод качественного

анализа слабоконтрастных медицинских радиологических изображений.

2. Метод позволяет выделять визуально неразличимые потенциально информативные участки радиологических изображений по фазо–пространственным характеристикам цифровой интерференционной характеристики преобразования Радона.

3. Обратное преобразование Радона от его трансформированных характеристик является неустойчивой операцией.

Список литературы: 1. Афанасьев В.А. Оптические измерения / В.А. Афанасьев – М.: Высшая школа, 1981. – 228 с. 2. Стюард И.Г. Введение в Фурье-оптику / И.Г. Стюард. – М.: Мир, 1985. – 182 с. 3. Аззам Р. Эллипсометрия и поляризованный свет / Р. Аззам, Н. Башира. – М.: Мир, 1981. – 583 с. 4. Pratt W.K. Digital Image Processing / W.K. Pratt. – New York: John Wiley and Sons Inc., 2001. – 723 p. 5. Гонсалес Р. Цифровая обработка изображений / Р. Гонсалес, Р. Вудс. – М.: Техносфера, 2006. – 1070 с. 6. Голографические неразрушающие исследования / Под ред. Р.К. Эрф. – М.: Машиностроение, 1979. – 446 с. 7. Ахметшина Л.Г. Информационные возможности модуляционного преобразования при сегментации мультиспектральных изображений / Л.Г. Ахметшина // Системні технології. – 2004. – № 6. – С. 122-127. 8. Ахметшин А.М. Интерференционные методы повышения качества и чувствительности анализа низкоконтрастных изображений на основе комплексной фазовой модуляции яркостей / А.М. Ахметшин, Л.Г. Ахметшина, И.М. Мацюк // Искусственный интеллект. – 2007. – № 3. – С. 194–204.

УДК 004.93

Интерференційний метод аналізу радіологічних зображень у просторі ознак перетворення Радону / Ахметшин О.М., Степаненко О.О. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 4–9.

Розглянуті інформаційні можливості нового методу якісного аналізу низько контрастних радіологічних медичних зображень у комплексному просторі яскравостей. Метод дозволяє сегментувати візуально невизначені потенційно інформативні ділянки шляхом використання фазо-просторових характеристик. Іл.: 3. Бібліогр.: 8 назв.

Ключові слова: низько контрастне зображення, комплексний простір яскравостей, сегментування, інформативні ділянки, фазо-просторові характеристики.

UDC 004.93

Interferometric method of radiological image analysis in a feature space Radon's transformation / Akhmetshin A.M., Stepanenko A.A. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 4–9.

Information possibilities a new method quality analysis of low contrast radiological medical images in a complex brightness space are considered. The method is possible to make segmentation hidden areas on base using phase–space characteristics. Figs.: 3. Refs.: 8 titles.

Key words: low contrast image, complex brightness space, segmentation, information hidden areas, phase-space characteristic.

Поступила в редакцію 28.01.2011

Я.Г. ВЕЛИКАЯ, Национальный исследовательский университет "БелГУ", Белгород

ПРОБЛЕМА ЭКВИВАЛЕНТНОСТИ В СТРУКТУРИРОВАННЫХ МОДЕЛЯХ ВЫЧИСЛЕНИЙ

В статье предлагается модификация трансформационного метода, позволяющая решить проблему эквивалентности для конечных детерминированных автоматов. Ил.: 2. Библиогр.: 8 назв.

Ключевые слова: трансформационный метод, проблема эквивалентности, конечные детерминированные автоматы.

Постановка проблемы. Для моделей вычислений существует ряд фундаментальных проблем: проблема эквивалентности; проблема построения полной системы эквивалентных преобразований; проблема минимизации. Для некоторых подклассов моделей вычислений данные проблемы решены. Существуют различные подходы для решения описанных проблем. Одним из подходов является подход, основанный на задании структуры модели вычислений в графическом виде. Модели вычислений, структура которых задана графически, будем называть структурированными. В работах Р.И. Подловченко и В.Е. Хачатряна [1] был предложен трансформационный метод для решения проблемы эквивалентности многоленточных автоматов, представленных в графическом виде. Ранее было доказано, что трансформационный метод позволяет доказать разрешимость проблемы эквивалентности для некоторых подклассов моделей вычислений, в частности, многоленточных автоматов с непересекающимися циклами, но не решает её для конечных детерминированных автоматов.

Анализ литературы. Под проблемой эквивалентности обычно понимается нахождение алгоритма, распознающего эквивалентность моделей вычислений. Доказано [2], что проблема эквивалентности в общем случае разрешима для многоленточных автоматов и, в частности, для конечных автоматов. Однако алгоритм разрешения многоленточных автоматов в работе [2] не был предложен. В статье Р. Берда [3] приводится решение проблемы эквивалентности для двухленточных автоматов и пример того, что предложенный в статье подход не решает проблему для трехленточных автоматов. В статье [4] предложен новый подход решения проблемы эквивалентности многоленточных автоматов. Что касается детерминированных конечных автоматов, то для них существует общеизвестный алгоритм разрешения эквивалентности [5].

Целью статьи является модификация трансформационного метода и обоснование того, что обобщенный трансформационный метод позволяет решить проблему эквивалентности для конечных детерминированных автоматов.

Трансформационный метод и его модификация. Как уже было сказано, трансформационный метод работает с многоленточными автоматами, представленными в графическом виде. Модель вычислений будем задавать в виде диаграммы. Последние диаграммы строятся над двумя конечными алфавитами: $P = \{p_1, p_2, \dots, p_n\}$ и $Q = \{0, 1\}$, где n – количество лент в автомате. По определению, диаграмма – это конечный ориентированный граф с размеченными вершинами и дугами. Его структура удовлетворяет следующим требованиям: в нем имеются две выделенные вершины, называемые входом и выходом диаграммы; из выхода нет исходящих дуг, а из всех остальных вершин исходят по две дуги; все вершины, кроме выхода, помечены символами алфавита P , а выходящие из вершин дуги помечены символами алфавита Q , причем дуги, выходящие из одной вершины, помечены различными символами. Любой конечный ориентированный путь u в диаграмме может быть описан историей $L(u) = ((a_1, \varepsilon_1), (a_2, \varepsilon_2), \dots, (a_n, \varepsilon_n))$, где a_i – это метка вершины, из которой выходит i -я дуга, ε_i – это метка i -й дуги пути L , $i = 1, 2, \dots, n$.

p_i -проекцией пути u называется слово, полученное из $L(u)$ удалением всех пар, не содержащих символа p_i , где $i = 1, 2, \dots, n$.

Определим маршрут, как путь, начинающийся во входе диаграммы. Маршрутом через диаграмму, назовем маршрут, заканчивающийся на выходе диаграммы.

Диаграммы D_1 и D_2 назовем эквивалентными, если для любого маршрута L_1 через одну из диаграмм найдется маршрут L_2 через другую диаграмму, такой что p_i -проекции маршрутов L_1 и L_2 совпадают.

Диаграммы D_1 и D_2 назовем строго эквивалентными, если для любого маршрута L_1 через одну из диаграмм найдется маршрут L_2 через другую диаграмму, такой что истории маршрутов L_1 и L_2 совпадают.

Предложенная модель будет интерпретироваться как конечный детерминированный автомат, если при сравнении маршрутов потребовать совпадение только меток дуг.

На рис. 1 приведен пример диаграммы конечного автомата (метка ленты опущена): вход диаграммы обозначен черным кружком, а выход перечеркнутым, дуги с меткой единица снабжены жирной точкой в начале дуги.

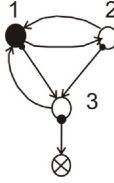


Рис. 1. Пример диаграммы конечного автомата

Трансформационный метод основан на полной системе фрагментных эквивалентных преобразований. Определим фрагмент автомата как часть автомата, определяемая заданным множеством состояний автомата и содержащая вместе с этими состояниями все инцидентные им дуги. Вершины и инцидентные им дуги сохраняют приписанные им в автомате метки. Под фрагментным преобразованием будем понимать замену в автомате одного фрагмента другим. В работе [6] построена полная система эквивалентных преобразований многоэлементных автоматов.

Определим характеристику диаграммы D , называемую покрытием. Это древовидный фрагмент, обозначим его $F(D)$, все вершины и дуги, которого являются образами вершин и дуг автомата D с их метками и список пар эквивалентных вершин, обозначим его S . Корнем древовидного фрагмента является образ входа автомата D . Обозначим через V список всех вершин автомата D , лежащих на маршрутах через автомат, за исключением его выхода. Внося в $F(D)$ какую-либо вершину из списка, будем вычеркивать ее образ из V .

На первом шаге в $F(D)$ вносится корень – образ входа v_0 автомата D , и вершина v_0 удаляется из списка V . Пусть на некотором шаге в $F(D)$ внесена вершина u , являющаяся образом вершины v автомата D и вершина v вычеркнута из V , причём вершина u – не выход и из неё ещё не выходят дуги. Обозначим α_1 и α_2 – дуги, исходящие из вершины v . Пусть α_i , $i = 1, 2$ оканчивается в вершине v_i . Если v_i содержится в V , создаем образ вершины v_i и направляем в нее дугу α_i ; удаляем v_i из списка V . Если v_i не содержится в списке V , но содержалась ранее и не является выходом, то создаем образ вершины v_i , обозначим его v_i' , объявляем его выходом фрагмента $F(D)$ и в нее направляем дугу α_i с ее меткой; пару (v_i, v_i') заносим в список S . Если v_i не содержится в списке V , и не содержалась там, то строящийся фрагмент $F(D)$ не меняется. Наконец, если v_i является выходом D , то он также будет выходом для $F(D)$, и дугу α_i с ее меткой, направляем в этот выход. В общем случае покрытие диаграммы строится неоднозначно. На рис. 2 изображены различные

покрытия диаграммы, изображенной на рис. 1. Диаграмму, для которой можно построить единственное покрытие назовем однозначной.

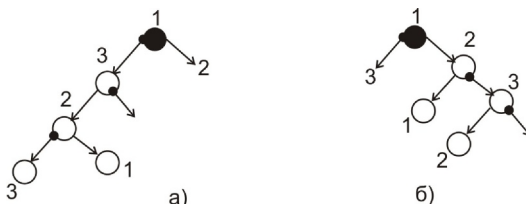


Рис. 2. Различные покрытия диаграммы

Опишем шаги процесса сравнения на эквивалентность диаграмм D_1 , D_2 , основанного на трансформационном методе.

Шаг 1. Построим покрытие $F(D_1)$ диаграммы D_1 и определим список пар эквивалентных вершин $S = \{(s_1, s_1'), \dots, (s_n, s_n')\}$.

Шаг 2. Трансформируем диаграмму D_2 , используя эквивалентные преобразования [6], в диаграмму D_3 , начинающуюся куполом, изоморфным $F(D_1)$. Купол диаграммы – это дерево, состоящее из некоторых вершин и инцидентных им дуг диаграммы, корнем, которого является вход диаграммы.

Если такое преобразование не возможно, то процесс заканчивает свою работу с заключением о том, что диаграммы D_1 и D_2 не являются эквивалентными. В противном случае строится список пар вершин $R = \{(r_1, r_1'), \dots, (r_n, r_n')\}$, каждый элемент которого является изоморфным образом вершин списка S .

Шаг 3. Выполним шаги 1 – 3 для каждой пары диаграмм, входы которых заданы парами из списка $R = \{(r_1, r_1'), \dots, (r_n, r_n')\}$. Это пары диаграмм: $(D_3(r_1), D_3(r_1')), \dots, (D_3(r_n), D_3(r_n'))$.

Описанный процесс прослеживается на дереве потомков $T(D_1, D_2)$.

Дерево потомков строится параллельно с вышеописанными шагами. Меткой корня дерева служит пара сравниваемых диаграмм (D_1, D_2) , а метками вершин – пары сравниваемых подграфов. Непосредственными потомками корня дерева $T(D_1, D_2)$ будут вершины, помеченные парами $(D_3(r_1), D_3(r_1')), \dots, (D_3(r_n), D_3(r_n'))$. У пар диаграмм $(D_3(r_1), D_3(r_1')), \dots, (D_3(r_n), D_3(r_n'))$ будут свои потомки и т.д. Сечение дерева $T(D_1, D_2)$ называется α -сечением, если все вершины сечения помечены парами изоморфных диаграмм. В [7] доказано, что применение трансформационного метода к паре эквивалентных конечных детерминированных автоматов может привести к построению дерева потомков, в котором нет α -сечения.

Предложим следующую модификацию трансформационного метода:

Для каждой пары диаграмм дерева потомков в процессе сравнения диаграмм на эквивалентность необходимо выполнить дополнительный шаг, а именно, первую из диаграмм предварительно преобразовать в однозначную диаграмму. В работе [8] доказано, что любую диаграмму эквивалентными преобразованиями можно трансформировать в однозначную диаграмму.

Обозначим через μ алгоритм сравнения на эквивалентность двух диаграмм, использующий модификацию трансформационного метода. Тогда можно доказать следующие утверждения:

Лемма 1. Если диаграммы D_1 и D_2 – строго эквивалентны, то в дереве потомков $T(D_1, D_2)$ непременно имеется α -сечение.

Лемма 2. α -сечение в дереве потомков $T(D_1, D_2)$, где D_1 и D_2 – строго эквивалентные диаграммы, строится за конечное число шагов, которое задается только количеством вершин в диаграммах D_1 и D_2 .

Теорема. Алгоритм μ является алгоритмом разрешения проблемы эквивалентности диаграмм.

Выводы. Модификация трансформационного метода, предложенная в статье, позволила решить проблему эквивалентности для детерминированных конечных автоматов. Данный метод работает со структурированными моделями вычислений и нацелен в общем случае на разрешение проблемы эквивалентности многоленточных автоматов.

Список литературы: 1. Подловченко Р.И. Об одном подходе к разрешению проблемы эквивалентности / Р.И. Подловченко, В.Е. Хачатрян // Программирование. – 2004. – № 3. – С. 3–20. 2. Harju T. The equivalence of multitape finite automata / T. Harju, J. Karhumäki // Theoret. Computer Sci. – 1991. – № 78. – P. 347–355. 3. Bird R. The equivalence problem for deterministic two-tape automata / R. Bird // J. of Computer and System Science. – 1973. – № 4. – P. 218–236. 4. Letichevsky Alexander A. The equivalence problem of deterministic multitape finite automata: a new proof of solvability using a multidimensional tape / Alexander A. Letichevsky, Arsen S. Shoukourian, Samvel K. Shoukourian // Language and Automata Theory and Applications. Lecture Notes in Computer Science. – 2010. – Vol. 6031/2010. – P. 392–402. 5. Карпов Е.А. Теория автоматов / Е.А. Карпов. – СПб: Питер, 2003. – 208 с. 6. Хачатрян В.Е. Полная система эквивалентных преобразований для многоленточных автоматов / В.Е. Хачатрян // Программирование. – 2003. – №1. – С. 62–77. 7. Хачатрян В.Е. Проблема эквивалентных преобразований для однородных многоленточных автоматов / В.Е. Хачатрян // Программирование. – 2008. – № 3. – С. 77–80. 8. Хачатрян В.Е. Модели вычислений с однозначным покрытием / В.Е. Хачатрян, Я.Г. Великая // Научные ведомости БелГУ. – 2009. – № 7 (62). – С. 116–121.

Статья представлена д.т.н., проф. НИУ "БелГУ" Корсуновым Н.И.

УДК 519.1: 681.3

Проблема еквівалентності в структурованих моделях обчислень / Велика Я.Г.
// Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 10 – 15.

У статті пропонується модифікація трансформційного методу, що дозволяє вирішити проблему еквівалентності для кінцевих детермінованих автоматів. Іл.: 2. Бібліогр.: 8 назв.

Ключові слова: трансформційний метод, проблема еквівалентності, кінцеві детерміновані автомати.

UDC 519.1: 681.3

Equivalence problem in the structured models of calculations/ Velikaya Ya.G.
// Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 10 – 15.

In article updating of the transformational method is offered, allowing to solve a problem of equivalence for final deterministic automata. Figs.: 2. Refs.: 8 titles.

Keywords: the transformational method, a problem of equivalence, finite deterministic automata.

Поступила в редакцію 10.05.2010

Є.О. ГОФМАН, аспірант, ЗНТУ, Запоріжжя,
А.О. ОЛІЙНИК, к.т.н., ЗНТУ, Запоріжжя,
С.О. СУББОТІН, к.т.н., доц. ЗНТУ, Запоріжжя

АГЕНТНИЙ МЕТОД СИНТЕЗУ ДЕРЕВ РІШЕНЬ

Розглянуто завдання побудови моделей у вигляді дерев рішень. Проаналізовано мультиагентний метод з непрямым зв'язком між агентами. Запропоновано мультиагентний метод ідентифікації дерев рішень. Розроблено програмне забезпечення, що реалізує запропонований метод. Бібліогр.: 10 назв.

Ключові слова: агент, дерево рішень, модель, мультиагентний метод.

Постановка проблеми. Розв'язання прикладних завдань у різних областях промисловості, медицини, економіки, веб-технологій пов'язано з необхідністю побудови моделей досліджуваних об'єктів, процесів та систем, для чого запропоновано різні підходи, зокрема методи регресійного аналізу, нейронні мережі, нечітке, нейро-нечітке моделювання й інші [1, 2]. Проте часто необхідно побудувати модель, яка не тільки забезпечую прийнятну точність, а і дозволяє легко зрозуміти, чому приймається відповідне рішення. Досить інтерпретабельними та зрозумілими є моделі, побудовані у вигляді дерев рішень [3 – 7].

Нехай задана навчальна вибірка даних, що складається з N екземплярів, кожний з яких характеризується P атрибутами. При цьому кожний атрибут може відноситися до деякого лінгвістичного терму T . Для кожного i -го екземпляра вказані входження до лінгвістичних термів для кожного атрибута та вказано лінгвістичний терм вихідної змінної. Тоді необхідно побудувати таке дерево рішень, яке дозволяє виконувати віднесення вихідного параметра до лінгвістичного терму із заданою точністю:

$$Q^* \geq Q_{threshold},$$

де Q^* – точність прогнозування за синтезованим деревом рішень;
 $Q_{threshold}$ – прийнятна точність прогнозування.

Аналіз літератури. Відомі різні методи ідентифікації дерев рішень, зокрема ID3, CART, CHAID, QUEST, C5.0 та ін. [3, 4]. Однак більшість із них мають певні недоліки, пов'язані з великою обчислювальною складністю, проблемами формування дерева рішень (ріст дерева, відсікання частини дерева) та ін. [3 – 7].

У зв'язку з цим актуальною є розробка нових методів синтезу дерев рішень, вільних від недоліків існуючих. Одним з нових напрямків

штучного інтелекту є мультиагентні методи з непрямим зв'язком між агентами, що дозволяють вирішувати різні оптимізаційні завдання [2, 8, 9]. Такі методи є особливо ефективними при розв'язанні завдань дискретної оптимізації, тому в цій статті пропонується застосувати мультиагентний підхід з непрямим зв'язком між агентами для розв'язання завдання ідентифікації дерев рішень.

Мета статті – розробка мультиагентного методу з непрямим зв'язком між агентами для синтезу дерев рішень.

Основними завданнями роботи є:

- дослідження основних принципів роботи дерев рішень;
- аналіз мультиагентного методу з непрямим зв'язком між агентами;
- розробка мультиагентного методу ідентифікації дерев рішень;
- розробка програмного забезпечення, що реалізує запропонований мультиагентний метод.

Дерева рішень. Дерева рішень являють собою спадну систему, засновану на підході "розділяй і пануй", основною метою якої є розподіл дерева на взаємно непересічні підмножини [3, 5]. Кожна підмножина являє собою підзадачу класифікації.

Дерево рішень описує процедуру прийняття рішення про приналежність певного екземпляра до того або іншого класу.

Дерево рішень є деревоподібною структурою, що складається із внутрішніх і зовнішніх вузлів, зв'язаних ребрами [6]. Внутрішні вузли – модулі, що приймають рішення, – розраховують значення функції розв'язку, на підставі чого визначають дочірній вузол, який буде відвідано далі. Зовнішні вузли (іноді називаються кінцевими вузлами), навпаки, не мають дочірніх вузлів і описують або мітку класу, або значення, що характеризує вхідні дані. У загальному випадку, дерева рішень використовуються в такий спосіб. Спочатку передаються дані (як правило, це вектор значень вхідних змінних) на кореневий вузол дерева рішень. В залежності від отриманого значення функції рішення, використовуваної у внутрішньому вузлі, відбувається перехід до одному з дочірніх вузлів. Такі переходи тривають доти, поки не буде відвідано кінцевий вузол, що описує або мітку класу, або значення, зв'язане з вхідним вектором значень ознак.

Мультиагентний метод з непрямим зв'язком між агентами. Мультиагентний метод з непрямим зв'язком між агентами є багатоагентним евристичним ітеративним методом випадкового пошуку [2, 8]. Поведінка агентів моделюється як процес переміщення й дослідження простору пошуку. Особливістю моделюемого переміщення є моделювання виділення феромонів, які агенти залишають на шляху в

процесі свого переміщення. Феромони в процесі роботи випаровуються. Таким чином, на найкращому шляху залишається більша кількість феромонів, оскільки додавання феромонів відбувається частіше, ніж випаровування. А оскільки вибір шляху для переміщення агентів ґрунтується на інформації про кількість феромонів, то агенти вибирають кращий шлях.

На етапі ініціалізації задаються параметри методу, що впливають на його роботу. Далі відбувається пересування агентів між вузлами графа, у результаті чого, після закінчення пересування кожного агента формуються розв'язки, з яких вибирається кращий на даній ітерації. Далі виконується перевірка критеріїв зупинення. Якщо перевірка була виконана успішно, відбувається закінчення пошуку, у процесі якого вибирається найкращий розв'язок з усіх, що зустрічалися на пройдених ітераціях. Якщо ж перевірка була неуспішною, то проводиться відновлення граней, яке полягає в імітації випаровування феромонів, та перезапуск агентів.

Ґрунтуючись на принципах мультиагентного методу, його різновидах і областях застосування [2, 8, 9] можна виділити наступні переваги та недоліки.

До переваг методу можна віднести те, що:

- він може використовуватися в динамічних застосуваннях (агенти адаптуються до змін навколишнього середовища);
- у процесі пошуку метод використовує пам'ять усієї колонії, що досягається за рахунок моделювання виділення феромонів;
- збіжність методу до оптимального розв'язку гарантується;
- стохастичність оптимізаційного процесу, тобто випадковість пошуку, за рахунок чого виключається можливість зациклення в локальному оптимумі;
- мультиагентність методу;
- можливість застосування до розв'язання різних завдань оптимізації.

До недоліків методу варто віднести такі:

- теоретичний аналіз ускладнений, оскільки підсумковий розв'язок формується в результаті послідовності випадкових розв'язків; розподіл ймовірностей змінюється при ітераціях; дослідження є більш експериментальними, ніж теоретичними;
- збіжність гарантується, але час збіжності не визначений;
- висока ітеративність методу;
- результат роботи методу досить сильно залежить від початкових параметрів пошуку, які підбираються експериментально.

Таким чином, можна відзначити, що розглянутий мультиагентний метод з непрямим зв'язком може ефективно вирішувати завдання переважно дискретної оптимізації, які можуть бути узгоджені з наступними вимогами:

- відповідне подання задачі – задача повинна бути описана у вигляді графа з набором вузлів і граней між вузлами;
- евристична придатність елементів графа, на основі яких формується розв'язок, – можливість застосування евристичної міри адекватності окремих елементів у графові пошуку;
- складання альтернативних розв'язків, за допомогою чого можна раціонально визначати припустимі розв'язки;
- правило відновлення феромонів – правило, яке визначає ймовірність переміщення агента з одного вузла графа до іншого.

Мультиагентний метод синтезу дерев рішень. Для виконання ідентифікації дерев рішень з використанням мультиагентного підходу з непрямим зв'язком між агентами слід перетворити основні етапи методу відповідно до особливостей розв'язуваного завдання.

Далі приводяться основні зміни в етапах мультиагентного методу.

1. Ініціалізація. На даному етапі в мультиагентному методі створюється граф пошуку, встановлюються параметри роботи методу, а також розраховуються евристичні міри важливості вузлів графа. Для розв'язання завдання побудови дерев рішень граф пошуку буде складатися з вузлів, що представляють окремі лінгвістичні терми, до яких можуть ставитися лінгвістичні змінні. При цьому для кожного лінгвістичного терма вихідної лінгвістичної змінної створюється окремий граф пошуку, для якого розраховуються окремі матриці евристичних значимостей і феромонів. У зв'язку з цим пошук на кожному графові пошуку виконується окремою множиною агентів. Такий підхід викликаний тим, що при розв'язанні завдання ідентифікації дерев рішень важливість лінгвістичних термів залежить від лінгвістичних термів вихідної змінної, а також має значення порядок відвідування вузлів агентами.

2. Пересування агентів. При пересуванні агенти приймають рішення, у який вузол переміститися, таким чином, формуються окремі дерева рішень для кожного лінгвістичного терма вихідної лінгвістичної змінної. Для такого розв'язку пропонується застосовувати правило випадкового вибору, що базується на евристичних мірах важливості та мірі пріоритетності, заснованої на моделюванні виділення феромонів у процесі пересування. Рішення про завершення переміщення окремого агента слід приймати, виходячи з того, наскільки добре побудоване

дерево рішень виділяє відповідні класи екземплярів вихідної навчальної вибірки.

3. Зміна ступеню значимості вузлів. При розв'язанні завдання синтезу дерев рішень як міру пріоритетності необхідно використовувати якість покриття окремого дерева екземплярів відповідного класу. Крім того, пропонується використовувати елітну стратегію, що дозволить забезпечити більш швидку збіжність до підсумкового розв'язку.

4. Відновлення феромонів. Процедура відновлення феромонів не має істотних особливостей для розв'язуваного завдання, тому її можна застосовувати в традиційному виді.

Виходячи з виділених особливостей, які повинен мати розроблюваний мультиагентний метод ідентифікації дерев рішень, було створено метод синтезу дерев рішень з непрямым зв'язком між агентами, який представлений у вигляді наступної послідовності етапів.

Eman 1. Ініціалізація. Задаються статичні параметри роботи методу: коефіцієнти α, β, ρ . Для кожного з можливих лінгвістичних термів вихідних значень створюється свій граф пошуку, що представляє собою лінгвістичні терми, які можуть бути включені в дерево рішень, і, відповідно, своя окрема множина агентів. Також важливою особливістю розроблювального методу є те, що створюються вузли, які характеризують інверсні лінгвістичні терми, тобто це необхідно для випадку, коли в дереві рішень вибирається варіант, що умова у вузлі не спрацювала. Крім того, для кожного графа пошуку розраховуються евристичні значення значимості окремого терма для відповідного лінгвістичного терма вихідної змінної (1):

$$\eta_p^q = \frac{\sum_{o=1}^N \min(\mu_p(o), \mu_q(o))}{\sum_{o=1}^N \mu_q(o)}, \quad \forall p = \overline{1, T}, \quad q = \overline{1, K}, \quad (1)$$

де η_p^q – значення евристичної значимості лінгвістичного терма p для опису класу q ; o – екземпляр вхідної вибірки, що містить N екземплярів; $\mu_p(o)$, $\mu_q(o)$ – значення функції приналежності об'єкта o терму p й класу q , відповідно; T – кількість лінгвістичних термів для вхідних змінних; K – кількість лінгвістичних термів вихідної змінної.

У кожному просторі пошуку кожному вузлу графа пошуку ставиться у відповідність початкове значення кількості феромонів τ_{init} :

$$\tau_p^q(1) = \tau_{init}, \quad \forall p = \overline{1, T}, \quad q = \overline{1, K}, \quad (2)$$

де $\tau_p^q(1)$ – значення кількості феромонів для p -го терма в просторі пошуку для q -го класу на першій ітерації пошуку.

Етап 2. Вибір терму для додавання в правило j -го агента в графі пошуку i -го лінгвістичного терму вихідної змінної.

Для j -го агента на основі випадкового правила вибору розраховується ймовірність включення k -го лінгвістичного терму в правило, що описує i -ий лінгвістичний терм вихідної змінної (3):

$$P_k^{i,j} = \frac{\eta_k^i \cdot \tau_k^i(t)}{\sum_{p \in R^j} \eta_p^i \cdot \tau_p^i(t)}, \quad (3)$$

де $P_k^{i,j}$ – імовірність додавання k -го терма в дерево рішень j -го агента в графові пошуку для i -го класу; R^j – множина термів, які можуть бути додані в дерево рішень j -го агента. Оскільки додавання певного лінгвістичного терму означає, що дерево рішень перейшло на наступний рівень і вже аналізується інша вхідна змінна, то крім обраного терму, виключаються також усі терми, що описують дану вхідну змінну.

Після цього перевіряється умова (4)

$$P_k^{i,j} > rand(1), \quad (4)$$

де $rand(1)$ – випадкове число з інтервалу $[0; 1]$.

Якщо умова виконується, тоді лінгвістичний терм k додається в дерево рішень j -го агента, видаляються терми, пов'язані з відповідною вхідною змінною, з множини можливих термів для даного агента й виконується перехід на наступний етап. В іншому випадку продовжується розгляд термів, що залишилися.

Етап 3. Перевірка завершення переміщення j -го агента.

Якщо множина термів, які j -ий агент може додати у формоване правило, є порожньою, то виконується перехід на наступний етап.

В іншому випадку визначається, скільки екземплярів i -го класу покриває дерево рішень j -го агента. Для всіх екземплярів, що відносяться до класу i , розраховується значення вихідної змінної відповідно до дерева рішень j -го агента, і на підставі одержуваних даних збільшується лічильник $cntMatch$, у якому зберігається кількість екземплярів, що покриваються отриманим деревом рішень. Після цього перевіряється умова (5)

$$cntMatch \geq inCntMatchMin_i, \quad (5)$$

де $inCntMatchMin_i$ – гранична мінімальна кількість екземплярів i -го класу, яка повинна визначатися деревом рішень.

Якщо зазначена умова виконується, то вважається, що дерево рішень ідентифікує необхідну кількість екземплярів, та j -ий агент завершив своє переміщення. Якщо усі агенти завершили своє переміщення ($j = cntAgents$) та перевірена приналежність до усіх класів ($i = k$), то виконується перехід на наступний етап. В іншому випадку виконується перехід до етапу 2.

Етап 4. Формування дерев рішень та оцінювання їх якості. Дерев рішень формуються випадковим чином шляхом усіляких накладень отриманих агентами розв'язків. При цьому сполучаються такі дерева рішень, які виконують відбір для однакових лінгвістичних термів вихідної змінної та з однаковим коренем (6):

$$p_1^{q,i} = p_1^{q,j}, \quad i \neq j, \quad (6)$$

де $p_1^{q,i}$ та $p_1^{q,j}$ – кореневі вузли дерев рішень i -го та j -го агентів, які використовуються для прогнозування q -го класу екземплярів.

Для оцінювання якості дерев рішень використовується вхідна навчальна вибірка, для кожного екземпляра якої визначається клас за відповідним деревом рішень. Грунтуючись на дані про клас екземплярів, отримані за допомогою дерева рішень, і класі екземплярів, виходячи із заданої навчальної вибірки, розраховують оцінку якості дерева рішень (7)

$$Q = \frac{cntMatch}{N}, \quad (7)$$

де $cntMatch$ – кількість екземплярів, для яких клас був визначений вірно за допомогою заданого дерева рішень; Q – якість прогнозування класу екземплярів на основі відповідної бази правил.

Якщо виконується умова (8)

$$Q_{high} \geq Q_{threshold}, \quad (8)$$

де Q_{high} – якість прогнозування дерева рішень, яка характеризується найкращою точністю прогнозування; $Q_{threshold}$ – прийнятна якість прогнозування, то відбувається перехід до етапу формування результуючого дерева рішень.

Етап 5. Додавання феромонів. Додавання феромонів виконується з метою підвищення пріоритетності тих термів, включення яких у дерево рішень сприяє підвищенню якості прогнозування результуючих дерев рішень. У зв'язку з цим кількість додаваного коефіцієнта пріоритетності прямо пропорційна якості прогнозування дерева рішень, у яке входить заданий лінгвістичний терм. При цьому додавання феромонів пропонується виконувати тільки для тих термів, які входять у дерева рішень, для яких виконується умова (9)

$$Q_{DT} \geq \delta \cdot Q_{high}, \quad (9)$$

де δ – коефіцієнт, що визначає, наскільки близько якість прогнозування дерева рішень DT повинно наближатися до кращої якості прогнозування етапу формування результуючого дерева рішень, щоб можна було застосовувати процедуру додавання феромонів для вузлів, що входять у дане дерево рішень DT .

Таким чином, додавання феромонів виконується для кожного терму, що входить у дерево рішень DT (10):

$$\tau_p^q(t) = \tau_p^q(t) + Q_{DT} \cdot \tau_p^q(t), \quad \forall p \in R, \quad \forall R \subset DT, \quad (10)$$

де $\tau_p^q(t)$ – кількість феромонів для терму p у графові пошуку для класу q , що визначається за допомогою відповідного дерева рішень.

Етап 6. Випаровування феромонів. Для виключення гірших термів, тобто таких, які при включенні їх у дерева рішень, знижують якість прогнозування, застосовують процедуру випаровування феромонів, яка виконується наприкінці кожної ітерації та застосовується для всіх вузлів у всіх графах пошуку. Випаровування феромонів виконується відповідно до формули (11)

$$\tau_p^q(t+1) = \rho \cdot \tau_p^q(t), \quad \forall p = \overline{1, T}, \quad q = \overline{1, K}, \quad (11)$$

де ρ – коефіцієнт випаровування, який задається при ініціалізації.

Якщо $t < t_{\max}$, то виконати перехід до етапу перезапуску агентів, а якщо ні, то вважається, що виконано максимально припустиму кількість ітерацій, і виконується перехід до етапу формування результуючого дерева рішень.

Етап 7. Перезапуск агентів. Усі дані про переміщення агентів у всіх графах пошуку обновляються, агенти розміщуються у випадкові точки графів пошуку, після чого виконується перехід до другого етапу.

Етап 8. Формування результуючого дерева рішень. Краще знайдене дерево рішень модифікується за допомогою традиційного мультиагентного методу з непрямым зв'язком між агентами. При цьому створюється граф пошуку з вузлів, що входять в обране краще дерево рішень. При цьому евристичними мірами пріоритетності вузлів є зважені значення феромонів для кожного терму, обчислені на основі отриманих матриць феромонів для кожного лінгвістичного терму вихідної змінної. Після цього виконується пошук агентами з непрямым зв'язком між ними за традиційною схемою мультиагентного пошуку з непрямым зв'язком. На основі отриманих результатів, з дерева рішень віддаляються ребра з найменшою кількістю феромонів, що дозволяє підвищити інтерпретабельність та логічну прозорість сформованого дерева рішень.

Запропонований мультиагентний метод ідентифікації дерев рішень з непрямым зв'язком між агентами було програмно реалізовано у середовищі пакету Matlab 7.0.

Для експериментів використовувалися тестові дані, які були взяті з загальнодоступних репозиторіїв [10]. Розроблений метод порівнювався з методом CART [5]. Були отримані дерева рішень, які характеризувалися точністю класифікації 81,2% і 86,1% для методу CART та розробленого мультиагентного методу синтезу дерев рішень з непрямым зв'язком між агентами, відповідно.

Виходячи з отриманих результатів, можна відзначити, що запропонований мультиагентний метод ідентифікації дерев рішень забезпечує синтез дерев рішень, які дозволяють виконувати класифікацію з більшою точністю, ніж у випадку синтезу дерев рішень з використанням існуючих методів.

Висновки. У роботі вирішено завдання синтезу дерев рішень.

Наукова новизна роботи полягає в тому, що розроблено мультиагентний метод ідентифікації дерев рішень, який використовує при пошуку декілька графів пошуку для кожного класу об'єктів, що дозволяє синтезувати дерева рішень з високою точністю класифікації.

Практична цінність отриманих результатів полягає в тому, що запропонований мультиагентний метод синтезу дерев рішень був програмно реалізований, і може бути використаний для розв'язання практичних завдань класифікації та прогнозування.

На основі результатів проведених експериментів можна зробити висновок, що запропонований метод ідентифікації дерев рішень дозволяє синтезувати дерева рішень з високими узагальнюючими властивостями, а також дозволяє уникати надлишкового розширення дерева.

Список літератури: 1. Рассел С. Искусственный интеллект: современный подход / С. Рассел, П. Норвиг. – М.: Вильямс, 2006. – 1408 с. 2. Субботін С.О. Неітеративні, еволюційні та мультиагентні методи синтезу нечіткологічних і нейромережних моделей: монографія / С.О. Субботін, А.О. Олійник, О.О. Олійник; під заг. ред. С.О. Субботіна. – Запоріжжя: ЗНТУ, 2009. – 375 с. 3. Rokach L. Data Mining with Decision Trees. Theory and Applications / L. Rokach, O. Maimon. – London: World Scientific Publishing Co, 2008. – 264 p. 4. Classification and regression trees / L. Breiman, J.H. Friedman, R.A. Olshen, C.J. Stone. – California: Wadsworth & Brooks, 1984. – 368 p. 5. Quinlan J.R. Simplifying decision trees / J.R. Quinlan // International Journal of Man-Machine Studies. – 1987. – № 27 (221). – P. 221–234. 6. Quinlan J.R. Decision trees and decision making / J.R. Quinlan // IEEE Transactions on Systems, Man and Cybernetics. – 1990. – № 2 (20). – P. 339–346. 7. Quinlan J.R. Induction of decision trees / J.R. Quinlan // Machine Learning. – 1986. – № 1. – P. 81–106. 8. Прогрессивные технологии моделирования, оптимизации и интеллектуальной автоматизации этапов жизненного цикла авиадвигателей: монографія / А.В. Богуслаев, Ал.А. Олейник, Ал.А. Олейник, Д.В. Павленко, С.А. Субботин. – Запоріжжя: ОАО "Мотор Сич", 2009. – 468 с. 9. Dorigo M. Optimization, Learning and Natural Algorithms / M. Dorigo. – Milano:

Politecnico di Milano, 1992. – 140 p. **10.** *UCI Machine Learning Repository* [electronic resource] / Center for Machine Learning and Intelligent Systems. – Access mode: <http://archive.ics.uci.edu/ml/datasets.html>.

Статья представлена д.т.н. проф. ЗНУ Гоменюком С.І.

УДК 004.93

Агентный метод синтеза деревьев решений / Гофман Е.А., Олейник А.А., Субботин С.А. // Вестник НТУ "ХПИ". Тематический выпуск: Информатика и моделирование. – Харьков: НТУ "ХПИ", – 2011. – № 17. – С. 16 – 25.

Рассмотрена задача построения моделей в виде деревьев решений. Проанализирован мультиагентный метод с косвенной связью между агентами. Предложен мультиагентный метод идентификации деревьев решений. Разработано программное обеспечение, реализующее предложенный метод. Библиогр.: 10 назв.

Ключевые слова: агент, дерево решений, модель, мультиагентный метод.

UDC 004.93

Agent-based method for the synthesis of decision trees / Gofman E.A., Oleynik A.A., Subbotin S.A. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 16 – 25.

The problem of constructing models in the form of decision trees is considered. Multiagent method of indirect communication between agents is analyzed. Multiagent method of decision trees identification is proposed. The software that implements the proposed method is developed. Refs.: 10 titles.

Key words: agent, decision tree, model, multi-agent method.

Поступила в редакцию 20.02.2011

В.Д. ДМИТРИЕНКО, д.т.н., проф. НТУ "ХПИ", Харьков,
А.Ю. ЗАКОВОРОТНЫЙ, к.т.н., ст. преп. НТУ "ХПИ", Харьков,
В.И. НОСКОВ, д.т.н., доц. НТУ "ХПИ", Харьков

ЛИНЕАРИЗАЦИЯ МАТЕМАТИЧЕСКОЙ МОДЕЛИ ДИЗЕЛЬ-ПОЕЗДА С ТЯГОВЫМ АСИНХРОННЫМ ПРИВОДОМ

Рассматривается синтез линейной математической модели дизель-поезда с тяговым асинхронным приводом на основе динамической линеаризации модели объекта управления средствами геометрической теории управления. На основании последовательности инволютивных распределений получена линейная математическая модель в форме Бруновского. Библиогр.: 15 назв.

Ключевые слова: линейная математическая модель, тяговый асинхронный привод, геометрическая теория управления, инволютивные распределения.

Постановка проблемы и анализ литературы. Тяговый подвижной состав железных дорог Украины является одним из основных потребителей электроэнергии и топлива. Поэтому снижение энергозатрат при перевозке пассажиров и грузов является одной из важнейших задач для Украинских железных дорог. Одним из путей уменьшения энергозатрат – это оптимизация управления тяговым подвижным составом. Вопросам оптимизации законов управления подвижным составом за последние десятилетия занимались многие ученые [1-10]. Однако в большинстве этих исследований использовались модели, описываемые системами обыкновенных нелинейных дифференциальных уравнений 2-3 порядка, а для асинхронного тягового привода – пятого порядка. Использование таких упрощенных моделей, с одной стороны, позволило решить ряд задач оптимального управления, но, с другой стороны, слишком упрощенное описание объекта управления не позволяет исследовать целый ряд процессов, влияющих на энергетические затраты тягового подвижного состава. Кроме того, даже при упрощенном описании тягового асинхронного привода системой нелинейных дифференциальных уравнений возникают серьезные трудности при синтезе оптимальных регуляторов с помощью большинства известных методов теории оптимального управления [11, 12]. В связи с этим в работах [10, 13] была предпринята попытка получить удобный математический инструмент для решения задачи управления тяговым приводом с помощью геометрической теории управления. При этом удалось получить законы оптимального управления для объектов, которые описывались системами обыкновенных нелинейных дифференциальных уравнений 5-6 порядка.

Однако при этом модель привода имела только один эквивалентный тяговый двигатель, что существенно ограничило возможности модели для поиска оптимальных законов управления реальным приводом. В связи с этим важно уточнение модели привода, разработка метода динамической линеаризации полученной модели (получение линейной модели в форме Бруновского) и поиск оптимальных законов управления с помощью этой модели.

Целью статьи является синтез линейной математической модели дизель-поезда с тяговым асинхронным приводом на основе динамической линеаризации модели объекта управления средствами геометрической теории управления.

Движение дизель-поезда по перегону может быть описано системой обыкновенных дифференциальных уравнений вида:

$$\frac{dS}{dt} = kV; \quad (1)$$

$$\frac{dV}{dt} = M_{T1} + M_{T2} - a_{21}V - a_{22}V^2; \quad (2)$$

$$\frac{d\Psi_{dj}}{dt} = -\alpha_j \Psi_{dj} + \alpha_j L_{mj} i_{dj}; \quad j = 1, 2; \quad (3)$$

$$\frac{di_{dj}}{dt} = -\gamma_j i_{dj} + p\Omega_j i_{qj} + \alpha_j L_{mj} \frac{i_{qj}^2}{\Psi_{dj}} + \alpha_j \beta_j \Psi_{dj} + \frac{1}{\sigma_j L_{sj}} u_{dj}; \quad j = 1, 2; \quad (4)$$

$$\frac{di_{qj}}{dt} = -\gamma_j i_{qj} - p\Omega_j i_{dj} - \alpha_j L_{mj} \frac{i_{dj} i_{qj}}{\Psi_{dj}} - p\beta_j \Omega_j \Psi_{dj} + \frac{1}{\sigma_j L_{sj}} u_{qj}; \quad j = 1, 2; \quad (5)$$

$$\frac{d\rho_j}{dt} = p \frac{V}{k_{\Omega j}} + \alpha_j L_{mj} \frac{i_{qj}}{\Psi_{dj}}; \quad j = 1, 2, \quad (6)$$

где S – расстояние, отсчитываемое от начала перегона; t – время; k , a_{21} , a_{22} – постоянные коэффициенты; V – скорость движения состава; M_{T1} , M_{T2} – соответственно тяговые моменты привода головного и последнего вагона дизель-поезда (здесь и далее полагаем, что индекс "1" относится к головному вагону дизель-поезда, а индекс "2" – к последнему); $M_{Tj} = k_j \mu_j \Psi_{dj} i_{qj}$, $j = 1, 2$; k_1 , k_2 – постоянные коэффициенты; $\mu_j = p L_{mj} / J_j L_r$; p – число пар полюсов статора тягового электродвигателя; L_{m1} , L_{m2} , L_{r1} , L_{r2} – соответственно индуктивность контура намагничивания и полная индуктивность ротора эквивалентного

тягового двигателя головного и последнего вагона; J_1, J_2 – моменты инерции, приведенные к валам тяговых двигателей;

$\Psi_{dj} = \sqrt{\Psi_{urj}^2 + \Psi_{vrj}^2}$ ($j=1, 2$) – соответственно потокоцепление ротора эквивалентного тягового двигателя головного и последнего вагона; Ψ_{urj} ,

Ψ_{vrj} – потокоцепления ротора эквивалентного тягового двигателя по

осям u и v ; $i_{qj} = i_{vsj} \cos p_j - i_{usj} \sin p_j$ – ток статора эквивалентного двигателя по оси q , $j = 1, 2$; i_{vsj} , i_{usj} – статорные токи по осям u и v

эквивалентных двигателей; $\rho_j = \arcsin\left(\Psi_{vrj} / \sqrt{\Psi_{urj}^2 + \Psi_{vrj}^2}\right)$ или

$\rho_j = \arccos\left(\Psi_{urj} / \sqrt{\Psi_{urj}^2 + \Psi_{vrj}^2}\right)$; $\alpha_j = 1/T_{rj}$; T_{rj} – постоянная времени

ротора j -го эквивалентного двигателя; $i_{dj} = i_{usj} \cos p_j - i_{vsj} \sin p_j$ – ток статора j -го двигателя по оси d в системе координат d, q ;

$\gamma_j = \frac{R_{rj} L_{mj}^2}{\sigma_j L_{sj} L_{rj}^2} + \frac{R_{sj}}{\sigma_j L_{sj}}$ ($j=1, 2$); R_{rj}, R_{sj} – активные сопротивления

соответственно ротора и статора j -го эквивалентного двигателя; σ_j, L_{sj} –

соответственно полный коэффициент рассеяния и полная индуктивность статора j -го двигателя; $\Omega_j = V/k_{\Omega j}$ – угловая скорость j -го двигателя;

$k_{\Omega j}$ ($j=1, 2$) – постоянные коэффициенты; $\beta_j = \frac{L_{mj}}{\sigma_j L_{sj} L_{rj}}$, $j=1, 2$;

$u_{dj} = u_{usj} \cos p_j + u_{vsj} \sin p_j$, $j=1, 2$; $u_{qj} = u_{usj} \cos p_j - u_{vsj} \sin p_j$, $j=1, 2$.

Введем в правые части уравнений (4), (5) объекта управления новые управления:

$$u_{1j} = p\Omega_j i_{qj} + \alpha_j L_{mj} i_{qj}^2 / \Psi_{dj} + \alpha_j \beta_j \Psi_{dj} + u_{dj} / (\sigma_j L_{sj}), \quad j=1, 2; \quad (7)$$

$$u_{2j} = -p\Omega_j i_{dj} - \alpha_j L_{mj} i_{dj} i_{qj} / \Psi_{dj} - p\beta_j \Omega_j \Psi_{dj} + u_{qj} / (\sigma_j L_{sj}), \quad j=1, 2. \quad (8)$$

Обозначив $x_1 = S$; $x_2 = V$; $x_3 = \Psi_{d1}$; $x_4 = i_{d1}$; $x_5 = i_{q1}$; $x_6 = \rho_1$; $x_7 = \Psi_{d2}$; $x_8 = i_{d2}$; $x_9 = i_{q2}$; $x_{10} = \rho_2$; $a_{235} = k_1 \mu_1$; $a_{279} = k_2 \mu_2$; $a_{33} = -\alpha_1$; $a_{34} = \alpha_1 L_{m1}$; $a_{44} = -\gamma_1 = a_{55}$; $a_{62} = p/k_{\Omega 1}$; $a_{635} = \alpha_1 L_{m1}$; $a_{77} = -\alpha_2$; $a_{78} = \alpha_2 L_{m2}$; $a_{88} = -\gamma_2 = a_{99}$; $a_{10,2} = p/k_{\Omega 2}$; $a_{10,7,9} = \alpha_2 L_{m2}$

и подставив управления (7), (8) в уравнения (4), (5) получим следующую модель движения дизель-поезда по перегону:

$$\begin{aligned}
 \frac{dx_1}{dt} &= x_2; & \frac{dx_6}{dt} &= a_{62}x_2 + a_{635} \frac{x_5}{x_3}; \\
 \frac{dx_2}{dt} &= a_{235}x_3x_5 + a_{279}x_7x_9 - a_{21}x_2 - a_{22}x_2^2; & \frac{dx_7}{dt} &= a_{77}x_7 + a_{78}x_8; \\
 \frac{dx_3}{dt} &= a_{33}x_3 + a_{34}x_4; & \frac{dx_8}{dt} &= a_{88}x_8 + u_{12}; \\
 \frac{dx_4}{dt} &= a_{44}x_4 + u_{11}; & \frac{dx_9}{dt} &= a_{99}x_9 + u_{22}; \\
 \frac{dx_5}{dt} &= a_{55}x_5 + u_{21}; & \frac{dx_{10}}{dt} &= a_{10,2}x_2 + a_{10,7,9} \frac{x_9}{x_7}.
 \end{aligned} \tag{9}$$

Для установления возможности преобразования системы нелинейных дифференциальных уравнений (9) к форме Бруновского определим выполнение условий инволютивности распределений M^0, M^1, M^2 для рассматриваемой системы [12]. С этой системой дифференциальных уравнений связаны векторные поля

$$X(x) = \begin{pmatrix} f_1 = x_2 \\ f_2 = a_{235}x_3x_5 + a_{279}x_7x_9 - a_{21}x_2 - a_{22}x_2^2 \\ f_3 = a_{33}x_3 + a_{34}x_4 \\ f_4 = a_{44}x_4 \\ f_5 = a_{55}x_5 \\ f_6 = a_{62}x_2 + a_{635} \frac{x_5}{x_3} \\ f_7 = a_{77}x_7 + a_{78}x_8 \\ f_8 = a_{88}x_8 \\ f_9 = a_{99}x_9 \\ f_{10} = a_{10,2}x_2 + a_{10,7,9} \frac{x_9}{x_7} \end{pmatrix}; Y_1 = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 1 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}; Y_2 = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 1 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}; Y_3 = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 1 \\ 0 \\ 0 \end{pmatrix}; Y_4 = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 0 \\ 1 \\ 0 \end{pmatrix},$$

где $x = (x_1, x_2, \dots, x_{10})$.

Векторные поля Y_1, Y_2, Y_3, Y_4 постоянны, поэтому распределение $M^0 = \text{span}\{Y_1, Y_2, Y_3, Y_4\}$ – инволютивно и $\dim M^0 = 4$, где $\text{span}\{Y_1, Y_2, Y_3, Y_4\}$ – линейная оболочка векторов $Y_1, \dots, Y_4$; $\dim M^0$ – размерность распределения M^0 .

Рассмотрим распределение $M^1 = \text{span}\{Y_1, Y_2, Y_3, Y_4, L_X Y_1, L_X Y_2, L_X Y_3, L_X Y_4\}$, где $L_X Y_j$ ($j = \overline{1, 4}$) – производные Ли вдоль векторного поля X векторных полей Y_j ($j = \overline{1, 4}$):

$$\begin{aligned}
 L_X Y_1 &= [X, Y_1] = \frac{\partial Y_1}{\partial \mathbf{x}} X - \frac{\partial X}{\partial \mathbf{x}} Y_1 = -\frac{\partial X}{\partial \mathbf{x}} Y_1 = \\
 &= - \begin{vmatrix} \frac{\partial f_1}{\partial x_1} & \frac{\partial f_1}{\partial x_2} & \dots & \frac{\partial f_1}{\partial x_{10}} \\ \dots & \dots & \dots & \dots \\ \frac{\partial f_{10}}{\partial x_1} & \frac{\partial f_{10}}{\partial x_2} & \dots & \frac{\partial f_{10}}{\partial x_{10}} \end{vmatrix} \cdot |0, 0, 0, 1, 0, 0, 0, 0, 0, 0|^T = \\
 &= |0, 0, -a_{34}, -a_{44}, 0, 0, 0, 0, 0, 0|^T; \\
 L_X Y_2 &= [X, Y_2] = \frac{\partial Y_2}{\partial \mathbf{x}} X - \frac{\partial X}{\partial \mathbf{x}} Y_2 = \left| 0, -a_{235}x_3, 0, 0, -a_{55}, -\frac{a_{635}}{x_3}, 0, 0, 0, 0 \right|^T; \\
 L_X Y_3 &= [X, Y_3] = \frac{\partial Y_3}{\partial \mathbf{x}} X - \frac{\partial X}{\partial \mathbf{x}} Y_3 = |0, 0, 0, 0, 0, 0, -a_{78}, -a_{88}, 0, 0|^T; \\
 L_X Y_4 &= [X, Y_4] = -\frac{\partial Y_4}{\partial \mathbf{x}} X - \frac{\partial X}{\partial \mathbf{x}} Y_4 = \\
 &= \left| 0, -a_{279}x_7, 0, 0, 0, 0, 0, 0, -a_{99}, -\frac{a_{10,7,9}}{x_7} \right|^T.
 \end{aligned}$$

Проверим инволютивность распределения M^1 . Для этого необходимо выполнение условия $\text{rank}(Y_1, Y_2, Y_3, Y_4, L_X Y_1, L_X Y_2, L_X Y_3, L_X Y_4, [X_i, X_j]) = 8$ где X_i, X_j – векторные поля из семейства $(Y_1, Y_2, Y_3, Y_4, L_X Y_1, L_X Y_2, L_X Y_3, L_X Y_4)$.

Так как

$$\begin{aligned}
 [L_X Y_1, L_X Y_2] &= \frac{\partial(L_X Y_2)}{\partial \mathbf{x}} L_X Y_1 - \frac{\partial(L_X Y_1)}{\partial \mathbf{x}} L_X Y_2 = \frac{\partial(L_X Y_2)}{\partial \mathbf{x}} \cdot L_X Y_1 = \\
 &= \frac{\partial(L_X Y_2)}{\partial \mathbf{x}} \cdot |0, 0, -a_{34}, -a_{44}, 0, 0, 0, 0, 0, 0|^T = \left| 0, a_{34}a_{235}, 0, 0, 0, -\frac{a_{34}a_{625}}{x_3^2}, 0, 0, 0, 0 \right|^T,
 \end{aligned}$$

то матрица $B = (Y_1, Y_2, Y_3, Y_4, L_X Y_1, L_X Y_2, L_X Y_3, L_X Y_4, [L_X Y_1, L_X Y_2])$ имеет ранг, равный 9, т.е. условие инволютивности распределения M^1 не выполняется. При этом $\det(B) = 2 \cdot a_{34}^2 \cdot a_{235} / x^3 \cdot a_{675} \cdot a_{78} \cdot a_{10,7,9} / x_7$.

Проверим инволютивность распределений $M_k^1 = \text{span}\{Y_1, Y_2, Y_3, Y_4, L_X Y_k\}$, $k = \overline{1, 4}$. Очевидно, что

$$\begin{aligned} [Y_1, Y_2] &= [Y_1, Y_3] = [Y_1, Y_4] = [Y_2, Y_3] = [Y_2, Y_4] = [Y_3, Y_4] = [Y_1, L_X Y_1] = \\ &= [Y_2, L_X Y_1] = [Y_3, L_X Y_1] = [Y_4, L_X Y_1] = |0, 0, 0, 0, 0, 0, 0, 0, 0|^T, \end{aligned}$$

поэтому распределение M_1^1 инволютивно. Имеем также

$$\begin{aligned} [Y_1, L_X Y_2] &= \frac{\partial(L_X Y_2)}{\partial x} Y_1 - \frac{\partial Y_1}{\partial x} L_X Y_2 = |0, 0, 0, 0, 0, 0, 0, 0, 0|^T; \\ [Y_2, L_X Y_2] &= \frac{\partial(L_X Y_2)}{\partial x} Y_2 - \frac{\partial Y_2}{\partial x} L_X Y_2 = \\ &= [Y_3, L_X Y_2] = [Y_4, L_X Y_2] = |0, 0, 0, 0, 0, 0, 0, 0, 0|^T; \\ [Y_k, L_X Y_3] &= [Y_k, L_X Y_4] = |0, 0, 0, 0, 0, 0, 0, 0, 0|^T, \quad k = \overline{1, 4}. \end{aligned}$$

Все подраспределения $M_k^1 = \text{span}\{Y_1, Y_2, Y_3, Y_4, L_X Y_k\}$, $k = \overline{1, 4}$, являются инволютивными, поэтому дополнительную переменную (или интегратор) можно вводить в любой канал управления. Однако введение одного, двух или трех интеграторов не позволяет решить проблему получения инволютивного распределения M^1 для расширенной системы. Распределение M^1 становится инволютивным при введении по два интегратора в каналы управления, связанные со вторым и четвертым управлениями.

Введем следующие обозначения

$$\begin{aligned} u_{11} &= u_1, \quad u_{21} = y_7; \quad x_i = y_i, \quad i = \overline{1, 6}; \\ \frac{dy_7}{dt} &= y_8, \quad \frac{dy_8}{dt} = u_2; \quad u_{12} = u_3; \quad x_{k+2} = y_k, \quad k = \overline{9, 12}; \quad u_{22} = y_{13}; \quad \frac{dy_{13}}{dt} = y_{14}; \\ \frac{dy_{14}}{dt} &= u_4. \end{aligned}$$

В новых обозначениях расширенная модель объекта управления имеет следующий вид:

$$\begin{aligned}
\frac{dy_1}{dt} &= y_2; & \frac{dy_8}{dt} &= u_2; \\
\frac{dy_2}{dt} &= a_{235}y_3y_5 + a_{279}y_9y_{11} - a_{21}y_2 - a_{22}y_2^2; & \frac{dy_9}{dt} &= a_{77}y_9 + a_{78}y_{10}; \\
\frac{dy_3}{dt} &= a_{33}y_3 + a_{34}y_4; & \frac{dy_{10}}{dt} &= a_{88}y_{10} + u_3; \\
\frac{dy_4}{dt} &= a_{44}y_4 + u_1; & \frac{dy_{11}}{dt} &= a_{99}y_{11} + y_{13}; \\
\frac{dy_5}{dt} &= a_{55}y_5 + y_7; & \frac{dy_{12}}{dt} &= a_{10,2}y_2 + a_{10,7,9} \frac{y_{11}}{y_9}; \\
\frac{dy_6}{dt} &= a_{62}y_2 + a_{635} \frac{y_5}{y_3}; & \frac{dy_{13}}{dt} &= y_{14}; \\
\frac{dy_7}{dt} &= y_8; & \frac{dy_{14}}{dt} &= u_4.
\end{aligned} \tag{10}$$

С расширенной моделью объекта управления связаны векторные поля:

$$\begin{aligned}
Y(y) &= |g_1, g_2, g_3, g_4, g_5, g_6, g_7, g_8, g_9, g_{10}, g_{11}, g_{12}, g_{13}, g_{14}|^T; \\
Y_1^* &= |0, 0, 0, 1, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0|^T; \\
Y_2^* &= |0, 0, 0, 0, 0, 0, 0, 1, 0, 0, 0, 0, 0, 0|^T; \\
Y_3^* &= |0, 0, 0, 0, 0, 0, 0, 0, 0, 1, 0, 0, 0, 0|^T; \\
Y_4^* &= |0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 1, 0|^T,
\end{aligned}$$

где $g_1 = y_2$; $g_2 = a_{235}y_3y_5 + a_{279}y_9y_{11} - a_{21}y_2 - a_{22}y_2^2$; $g_3 = a_{33}y_3 + a_{34}y_4$; $g_4 = a_{44}y_4$; $g_5 = a_{55}y_5 + y_7$; $g_6 = a_{62}y_2 + a_{635} \cdot y_5/y_3$; $g_7 = y_8$; $g_8 = 0$; $g_9 = a_{77}y_9 + a_{78}y_{10}$; $g_{10} = a_{88}y_{10}$; $g_{11} = a_{99}y_{11} + y_{13}$; $g_{12} = a_{10,2}y_2 + a_{10,7,9} \cdot y_{11}/y_9$; $g_{13} = y_{14}$; $g_{14} = 0$.

Для расширенной модели объекта управления распределение $M^{0*} = \text{span}\{Y_1^*, Y_2^*, Y_3^*, Y_4^*\}$ инволютивно и $m_0 = \dim M^{0*} = 4$. Проверим инволютивность распределения $M^{1*} = \text{span}\{Y_1^*, Y_2^*, Y_3^*, Y_4^*, L_Y Y_1^*, L_Y Y_2^*, L_Y Y_3^*, L_Y Y_4^*\}$. Поскольку имеем

$$L_Y Y_1^* = -\frac{\partial Y}{\partial y} Y_1^* = -|0, 0, a_{34}, a_{44}, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0|^T;$$

$$\begin{aligned}
L_Y^* Y_2^* &= -\frac{\partial Y}{\partial y} Y_2^* = -|0, 0, 0, 0, 0, 0, 1, 0, 0, 0, 0, 0, 0|^T; \\
L_Y^* Y_3^* &= -\frac{\partial Y}{\partial y} Y_3^* = -|0, 0, 0, 0, 0, 0, 0, 0, a_{78}, a_{88}, 0, 0, 0|^T; \\
L_Y^* Y_4^* &= -\frac{\partial Y}{\partial y} Y_4^* = -|0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 1, 0|^T,
\end{aligned}$$

т.е. все вектора распределения M^{1*} имеют постоянные компоненты, то распределение M^{1*} инволютивно $m_1 = \dim M^{1*} = 8$.

Проверим инволютивность распределения $M^{2*} = \text{span}\{M^{1*}, [Y^*, M^{1*}]\} = \text{span}\{Y_1^*, Y_2^*, Y_3^*, Y_4^*, L_Y Y_1^*, L_Y Y_2^*, L_Y Y_3^*, L_Y Y_4^*, L_Y^2 Y_1^*, L_Y^2 Y_2^*, L_Y^2 Y_3^*, L_Y^2 Y_4^*\}$. Имеем

$$\begin{aligned}
L_Y^2 Y_1^* &= [Y, L_Y Y_1^*] = -\frac{\partial Y}{\partial y} L_Y Y_1^* = \\
&= \left| 0, a_{34}a_{235}y_5, a_{33}a_{34} + a_{34}a_{44}, a_{44}^2, 0, \frac{a_{34}a_{635}y_5}{y_3^2}, 0, 0, 0, 0, 0, 0, 0 \right|^T; \\
L_Y^2 Y_2^* &= -\frac{\partial Y}{\partial y} L_Y Y_2^* = |0, 0, 0, 0, 1, 0, 0, 0, 0, 0, 0, 0, 0|^T; \quad L_Y^2 Y_3^* = -\frac{\partial Y}{\partial y} L_Y Y_3^* = \\
&= \left| 0, a_{78}a_{279}y_{11}, 0, 0, 0, 0, 0, 0, a_{77}a_{78} + a_{78}a_{88}, a_{88}, 0, -\frac{a_{78}a_{10,7,9}y_{11}}{y_9^2}, 0, 0 \right|^T; \\
L_Y^2 Y_4^* &= -\frac{\partial Y}{\partial y} L_Y Y_4^* = |0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 1, 0, 0|^T; \\
[L_Y^2 Y_1^*, L_Y^2 Y_3^*] &= \frac{\partial(L_Y^2 Y_3^*)}{\partial y} L_Y^2 Y_1^* - \frac{\partial(L_Y^2 Y_1^*)}{\partial y} L_Y^2 Y_3^* = \\
&= |0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0|^T.
\end{aligned}$$

Проверяем другие пары:

$$\begin{aligned}
[L_Y^2 Y_1^*, L_Y^2 Y_2^*] &= \frac{\partial(L_Y^2 Y_2^*)}{\partial y} L_Y^2 Y_1^* - \frac{\partial(L_Y^2 Y_1^*)}{\partial y} L_Y^2 Y_2^* = \\
&= \left| 0, a_{34}a_{235}, 0, 0, 0, -\frac{a_{34}a_{635}}{y_3^2}, 0, 0, 0, 0, 0, 0, 0 \right|^T;
\end{aligned}$$

$$\begin{aligned}
[L_Y^2 Y_1^*, L_Y^2 Y_4^*] &= |0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0|^T; \\
[L_Y^2 Y_3^*, L_Y^2 Y_4^*] &= \frac{\partial(L_Y^2 Y_4^*)}{\partial y} L_Y^2 Y_3^* - \frac{\partial(L_Y^2 Y_3^*)}{\partial y} L_Y^2 Y_4^* = \\
&= - \left| 0, a_{78} a_{279}, 0, 0, 0, 0, 0, 0, 0, 0, -\frac{a_{78} a_{10,7,9}}{y_9^2}, 0, 0 \right|^T; \\
[L_Y^2 Y_3^*, L_Y^2 Y_2^*] &= \frac{\partial(L_Y^2 Y_2^*)}{\partial y} L_Y^2 Y_3^* - \frac{\partial(L_Y^2 Y_3^*)}{\partial y} L_Y^2 Y_2^* = \\
&= |0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0, 0|^T
\end{aligned}$$

Поскольку ранг матрицы $R_2 = (Y_1^*, Y_2^*, Y_3^*, Y_4^*, L_Y Y_1^*, L_Y Y_2^*, L_Y Y_3^*, L_Y Y_4^*, L_Y^2 Y_1^*, L_Y^2 Y_2^*, L_Y^2 Y_3^*, L_Y^2 Y_4^*, [L_Y^2 Y_3^*, L_Y^2 Y_4^*])$ не равен 12, то распределение M^{2*} не является инволютивным.

Динамическая линеаризация при наличии нескольких управлений возможна при наличии инволютивности более простых подраспределений:

$$\begin{aligned}
M_1^{2*} &= \text{span}\{Y_1^*, Y_2^*, Y_3^*, Y_4^*, L_Y Y_1^*, L_Y^2 Y_1^*\}, \\
M_2^{2*} &= \text{span}\{Y_1^*, Y_2^*, Y_3^*, Y_4^*, L_Y Y_2^*, L_Y^2 Y_2^*\}, \\
M_3^{2*} &= \text{span}\{Y_1^*, Y_2^*, Y_3^*, Y_4^*, L_Y Y_3^*, L_Y^2 Y_3^*\}, \\
M_4^{2*} &= \text{span}\{Y_1^*, Y_2^*, Y_3^*, Y_4^*, L_Y Y_4^*, L_Y^2 Y_4^*\}.
\end{aligned}$$

Подраспределения M_2^{2*} и M_4^{2*} являются инволютивными, поскольку все их вектора имеют постоянные компоненты.

Исследования распределений:

$$\begin{aligned}
M_2^{3*} &= \text{span}\{M_2^{2*}, [X^*, M_2^{2*}]\} = \text{span}\{Y_1^*, Y_2^*, Y_3^*, Y_4^*, L_Y Y_2^*, L_Y^2 Y_2^*, L_Y^3 Y_2^*\}; \\
M_4^{3*} &= \text{span}\{M_4^{2*}, [X^*, M_4^{2*}]\} = \text{span}\{Y_1^*, Y_2^*, Y_3^*, Y_4^*, L_Y Y_4^*, L_Y^2 Y_4^*, L_Y^3 Y_4^*\},
\end{aligned}$$

показывает, что они также являются инволютивными.

На основании теории о линейных эквивалентах для нелинейных аффинных систем с t управлениями [14, 15] получим, что каноническая форма Бруновского имеет четыре клетки, и индекс управляемости $k_{\max}$ для данного объекта равен четырем. Поскольку число рассматриваемых инволютивных распределений M_j , $j = \overline{0, k_{\max} - 1}$, то условия для получения линейного эквивалента для рассматриваемого объекта выполнены.

В результате получим математическую модель в форме Бруновского:

$$\begin{aligned} \frac{dz_i}{dt} &= z_{i+1}, \quad i = 1, 2, 3, 5, 6, 7, 9, 11, 12, 13; \\ \frac{dz_4}{dt} &= v_1; \quad \frac{dz_8}{dt} = v_2; \quad \frac{dz_{10}}{dt} = v_3; \quad \frac{dz_{14}}{dt} = v_4, \end{aligned} \quad (11)$$

где v_k ($k = \overline{1, 4}$) – управления.

Поскольку модель объекта в форме Бруновского имеет четыре клетки, то необходимо определить четыре функции $T_j(y)$ ($j = \overline{1, 4}$) – преобразования от расширенной модели объекта управления к модели в форме Бруновского. Известно, что такие функции существуют и методика получения их известна [13 – 15]. Математическое моделирование объекта управления в форме Бруновского показало его работоспособность.

Выводы. Таким образом, впервые средствами дифференциальной геометрии получена работоспособная линейная математическая модель дизель-поезда, с тяговым асинхронным приводом, эквивалентная нелинейной математической модели, описываемой системой нелинейных обыкновенных дифференциальных уравнений десятого порядка. Полученная линейная модель в канонической форме Бруновского может быть использована для синтеза оптимальных законов управления тяговым подвижным составом.

Список литературы: 1. Ковальский А.Н. Синтез автоматического управления поездом метрополитена (САУ-М) и ее модернизация / А.Н. Ковальский // Труды МИИЖТ. – Вып. 276. – М.: МИИЖТ, 1968. – С. 3 – 13. 2. Петров Ю.П. Оптимальное управление движением транспортных средств / Ю.П. Петров. – Л.: Энергия, 1969. – 96 с. 3. Шинская Ю.В. Расчет оптимальных режимов ведения поездов метрополитена методом динамического прогнозирования / Ю.В. Шинская // Труды ЛИИЖТ. – Вып. 315. – Л.: ЛИИЖТ, 1970. – С. 18 – 23. 4. Легостаев Е.Н. Автоматизация управления движением поездов на метрополитенах / Е.Н. Легостаев, И.П. Исаев, А.Н. Ковальский. – М.: Транспорт, 1976. – 96 с. 5. Кудрявцев Я.Б. Принцип максимума и оптимальное управление движением поезда / Я.Б. Кудрявцев // Вісник ВНИИЖТ. – 1977. – № 1. – С. 57 – 61. 6. Костромин А.М. Оптимизация управления локомотивом. / А.М. Костромин. – М.: Транспорт, 1979. – 119 с. 7. Носков В.И. Моделирование и оптимизация систем управления и контроля локомотивов / В.И. Носков, В.Д. Дмитриенко, Н.И. Заповольский, С.Ю. Леонов. – Харьков: ХФИ "Транспорт Украины", 2003. – 248 с. 8. Дмитриенко В.Д. Математическое моделирование и оптимизация системы управления тяговым электроприводом / В.Д. Дмитриенко, В.И. Носков, М.В. Липчанский // Системи обробки інформації. – Харків: ХУПС. – 2004. – Вип. 11 (39). – С. 55–62. 9. Дмитриенко В.Д. Определение оптимальных режимов ведения дизель-поезда с использованием нейронных сетей АРТ / В.Д. Дмитриенко, В.И. Носков, М.В. Липчанский, А.Ю. Заковоротный // Вісник НТУ "ХПІ". – Харків: НТУ "ХПІ". – 2004. – № 46. – С. 90 – 96. 10. Дмитриенко В.Д. Синтез оптимальных законов управления тяговым

электроприводом методами дифференциальной геометрии и принципа максимума // В.Д. Дмитриенко, А.Ю. Заковоротный // Системы обработки информации. – Харьков: ХУПС. – 2009. – Вып. 4 (78). – С. 42–51. **11.** Методы классической и современной теории автоматического управления: Учебник в 5-ти томах. Т. 4: Теория оптимизации систем автоматического управления / Под ред. К.А. Пупкова и И.Д. Егунова. – М.: МГТУ им. Н.Э. Баумана, 2004. – 744 с. **12.** Методы классической и современной теории автоматического управления: Учебник в 5-и томах. Т. 5: Методы современной теории управления / Под ред. К.А. Пупкова, Н.Д. Егунова. – М.: Изд-во МГТУ им. Н.Э. Баумана, 2004. – 784 с. **13.** Дмитриенко В.Д. Синтез оптимальных законов управления движением дизель-поезда с помощью математической модели в форме Бруновского / В.Д. Дмитриенко, А.Ю. Заковоротный, Н.В. Мезенцев // Информационно-керуючі системи на залізничному транспорті. – Харків: УкрДАЗТ. – 2010. – Вып. 5-6. – С. 7–13. **14.** Краснощёченко В.И. Нелинейные системы: геометрический метод анализа и синтеза / В.И. Краснощёченко, А.П. Грищенко. – М.: Изд-во МГТУ им. Н.Э. Баумана. – 2005. – 520 с. **15.** Qiang Lu. Nonlinear control systems and power system dynamics. / Lu Qiang, Sun Yuangzhang, Mei Shengwei. – 2001. – 376 с.

УДК 517.9:629.42

Лінеаризація математичної моделі дизель-поїзда з тяговим асинхронним приводом / Дмитрієнко В.Д., Заковоротний О.Ю., Носков В.І. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 26 – 36.

Розглядається синтез лінійної математичної моделі дизель-поїзда з тяговим асинхронним приводом на основі динамічної лінеаризації моделі об'єкта керування засобами геометричної теорії керування. На підставі послідовності інволютивних розподілів отримана лінійна математична модель у формі Бруновського. Бібліогр.: 15 назв.

Ключові слова: лінійна математична модель, тяговий асинхронний привод, геометрична теорія керування, інволютивні розподіли.

UDC 517.9:629.42

Linearization mathematical model diesel train with asynchronous traction drive / Dmitrienko V.D., Zakovorotnyi A.Y., Noskov V.I. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 26 – 36.

A synthesis of linear mathematical model diesel train with asynchronous traction drive based on dynamic object model linearization control of the means of geometric control theory. Based on the sequence of involutive transformations received linear mathematical model in the form of Brunovski. Refs.: 15 titles.

Keywords: linear mathematical model, asynchronous traction drive, geometric control theory, involutive transformations.

Поступила в редакцію 01.03.2011

Е.Г. ЖИЛЯКОВ, д.т.н., проф. НИУ "БелГУ", Белгород,
А.В. БОЛДЫШЕВ, аспирант НИУ "БелГУ", Белгород,
А.В. КУРЛОВ, аспирант НИУ БелГУ, Белгород,
А.А. ФИРСОВА, аспирант НИУ "БелГУ", Белгород,
А.В. ЭСАУЛЕНКО, магистр НИУ "БелГУ", Белгород

ОБ ИЗБИРАТЕЛЬНОМ ПРЕОБРАЗОВАНИИ ЧАСТОТНЫХ КОМПОНЕНТ РЕЧЕВЫХ СИГНАЛОВ В ЗАДАЧЕ СЖАТИЯ

В статье приведены результаты вычислительных экспериментов по апробации метода сжатия речевых данных на основе избирательного преобразования частотных компонент речевых сигналов, полученных с помощью нового метода субполосного частотного анализа/синтеза. Ил.: 1. Табл.: 2. Библиогр.: 8 назв.

Ключевые слова: частотные компоненты речевых сигналов, сжатие речевых данных, субполосный частотный анализ/синтез.

Постановка проблемы и анализ литературы. Одной из особенностей звуков русской речи является сосредоточенность энергии в достаточно узких частотных диапазонах, суммарная ширина которых гораздо меньше частоты дискретизации [1 – 5]. Эта особенность может быть использована в различных направлениях области обработки речевых сообщений, в том числе, в задаче сжатия речевых данных, для этого необходимо точно определить, в каком количестве частотных интервалов сосредоточена основная доля энергии [4, 6]. Чтобы определить количество частотных интервалов, в которых сосредоточена основная доля энергии, можно воспользоваться понятием частотной концентрации, которая определяется минимальным количеством частотных интервалов, в которых сосредоточена заданная доля энергии отрезка речевого сигнала [6]:

$$W_{NR}^t = l_{NR}^{tm} / R, \quad (1)$$

где l_{NR}^{tm} – минимальное количество частотных интервалов, в которых сосредоточена заданная доля энергии отрезка речевого сигнала, так что имеет место

$$l_{NR}^{tm} = \min d_{NR}^{tm}; \quad (2)$$

R – количество частотных интервалов, на которые разбивается частотный диапазон;

t – обозначает один из анализируемых речевых отрезков, порождаемых звуком русской речи;

N – значение длины анализируемого отрезка;

m – доля общей энергии, задаваемая для определения минимального количества частотных интервалов, в которых она сосредоточена.

Для правых частей (2) выполняется неравенство

$$\sum_{k=1}^{d_{NR}^m} P_{(k)N} \geq m \|\vec{x}_N\|^2 = m \sum_{i=1}^N x_i^2, \quad (3)$$

где $\vec{x}_N = (x_1, \dots, x_N)^T$ – анализируемый отрезок речевых данных;

T – операция транспонирования.

Индекс в скобках, у слагаемых суммы слева соотношения (3) означает, что части энергий P_{kN} упорядочиваются по убыванию:

$$P_{(k)N} \in \{P_{rN}, r = 1, \dots, R\}; P_{(k+1)N} \leq P_{(k)N}, k = 1, \dots, R. \quad (4)$$

В качестве примера в таблице 1 приведено минимальное количество частотных интервалов, составляющих заданную долю энергии, для звука русской речи "И".

Таблица 1. Минимальное количество частотных интервалов и частотная концентрация при различных значениях параметра m (звук "И")

m	I_{NR}^m	W_{NR}^t
0,98	5	0,3175
0,96	3	0,1875
0,9 – 0,94	2	0,125
0,86 – 0,88	1	0,0625

При заданной доле энергии 0,98 частотная концентрация для приведенного примера составляет 0,3125. Для большинства звуков русской речи величина частотной концентрации составляет порядка 0,35 и только для шумоподобных звуков – порядка 0,55 – 0,60 [6]. На основании сведений о количестве и расположении частотных интервалов, в которых сосредоточена заданная доля энергии, можно осуществить сжатие речевых данных за счет хранения только составляющих речевого сигнала, соответствующих этим частотным интервалам.

Целью статьи является анализ эффективности сжатия речевых данных за счет использования предлагаемого метода субполосного частотного анализа/синтеза.

Одним из способов получения составляющих речевого сигнала, соответствующих выбранным частотным интервалам, является субполосный частотный анализ/синтез. В настоящее время наибольшее

распространение получил метод на основе банка КИХ-фильтров, однако, этот метод обладает рядом недостатков, которые приводят к увеличению погрешностей восстановления данных [3]. В ряде публикаций [1 – 3] описывается метод субполосного анализа/синтеза, оптимальный с точки зрения минимума среднеквадратической погрешности аппроксимации трансформант Фурье исходного отрезка речевого сигнала в заданном частотном интервале, а также показаны преимущества этого метода перед современными аналогами. В основе предлагаемого метода лежит математический аппарат с использованием субполосных матриц вида:

$$A_r = \{a_{ik}^r\} = \{\sin(v_r(i-k)) - \sin(v_{r-1}(i-k))\} / \pi(i-k), \quad i, k = 1, \dots, N, \quad (5)$$

где v_r и v_{r-1} верхняя и нижняя границы частотного интервала.

Эта матрица является симметричной и неотрицательно определённой, поэтому она обладает полной системой ортонормальных собственных векторов, соответствующих неотрицательным собственным числам [3, 7]. На основе этих матриц можно вычислять точные значения долей энергий отрезков речевых сигналов в выбранных частотных интервалах

$$P_r = \bar{x} A_r \bar{x}^T. \quad (6)$$

Данный математический аппарат можно использовать для получения частотных компонент исходного анализируемого отрезка речевого сигнала, которые отражают заданную долю энергии, сосредоточенную в выбранных частотных интервалах l . Для этого будет использован метод субполосного преобразования, основанный на формировании составной матрицы, которая вычисляется как сумма субполосных матриц, соответствующих выбранным частотным интервалам, составляющих заданную долю энергии m

$$A_\Sigma = \sum_{r=1}^{d_{NR}^m} A_{(r)}, \quad (7)$$

где $A_{(r)}$ – субполосные матрицы, соответствующие тем частотным интервалам, которые составляют заданную долю энергии m .

Составная матрица обладает полной системой ортонормальных собственных векторов (8), соответствующих неотрицательным собственным числам (9):

$$Q_\Sigma = \{\bar{q}_{\Sigma 1}, \bar{q}_{\Sigma 2}, \dots, \bar{q}_{\Sigma N}\}, \quad (8)$$

$$L_{\Sigma} = \text{diag}(\lambda_{\Sigma 1}, \dots, \lambda_{\Sigma N}) . \quad (9)$$

Собственные числа количественно равны сосредоточенным в выбранных частотных интервалах долям энергий соответствующих собственных векторов и удовлетворяют условию

$$0 \leq \lambda_{\Sigma k} \leq 1, \quad k = 1, \dots, N . \quad (10)$$

Субполосное преобразование осуществляется следующим образом

$$\vec{y}_{\Sigma} = \sqrt{L_{\Sigma}} Q_{\Sigma}^T \vec{x} , \quad (11)$$

где $\sqrt{L_{\Sigma}}$ – корень из диагонального элемента, соответствующего определенному собственному вектору.

С точки зрения сжатия речевых данных, можно поставить задачу нахождения минимального количества собственных значений составной матрицы, при оставлении которых будет достигаться максимальная степень сжатия. Сжатие исходных речевых данных будет осуществляться за счет хранения вектора субполосного преобразования, размерностью равной минимальному количеству ненулевых собственных значений. При этом важным условием является минимизация погрешности восстановления исходного отрезка речевых данных, т.е. обеспечение высокого качества воспроизведения исходного речевого сообщения. Выражение (6) можно представить как

$$P_{\Sigma} = \sum_{i=1}^{J_{\Sigma}} y_{\Sigma i}^2 + \sum_{i=J_{\Sigma}+1}^N y_{\Sigma i}^2 ,$$

где $\sum_{i=1}^{J_{\Sigma}} y_{\Sigma i}^2$ – первое слагаемое, в котором λ_i – собственные значения суммарной матрицы, величина которых достаточно большая;

$\sum_{i=J_{\Sigma}+1}^N y_{\Sigma i}^2$ – второе слагаемое, в котором λ_i – собственные значения суммарной матрицы, величина которых достаточно мала (близка к 0).

Доля энергии второго слагаемого, настолько мала, что ею можно пренебречь без получения существенных искажений. Таким образом, для оценки минимального количества собственных значений J_{Σ} , необходимых для восстановления исходного отрезка речевого сигнала без существенных потерь, можно использовать следующее выражение

$$\sum_{i=1}^{J_{\Sigma}} \lambda_i / \sum_{i=1}^N \lambda_i \geq c ,$$

где c – некий порог, который показывает, какую долю составляют собственные значения, величина которых близка к 0.

Вычислительные эксперименты. В качестве исходных данных использовано большое количество речевого материала, который был получен в результате записи естественной речи различных дикторов. Для проведения экспериментальных исследований были выбраны следующие параметры: порог $c = 0.89 \div 0.995$, длительность отрезка речевых данных $N = 160$, количество частотных интервалов $R = 16$, заданная доля энергии отрезка речевого сигнала $m = 0.86 \div 0.98$. В качестве примера в табл. 2 приведено минимальное количество собственных значений J_{Σ} для звука русской речи "И".

Таблица 2. Звук "И", $N = 160$, $R = 16$

$c \backslash m$	0.86 – 0.88	0.9 – 0.94	0.96	0.98
0.89 – 0.9	10	18	28	46
0.91	10	19	29	47
0.92 – 0.93	10	19	29	48
0.94 – 0.95	11	20	30	49
0.96	11	20	31	51
0.97	12	21	32	52
0.98	12	22	33	54
0.99	13	22	34	56
0.995	14	23	36	58

Как видно из приведенных результатов вычислительных экспериментов, минимальное количество собственных значений составной матрицы для данного звука в среднем составляет порядка 40, что позволяет говорить о возможности четырехкратного сокращения объема памяти, требуемого для хранения данного звука. Однако, выбираемое количество собственных значений будет сказываться на качестве воспроизведения речевого сообщения. Ниже на рис. 1 приведена полученная степень сжатия для всех звуков русской речи. Степень сжатия определялась следующим образом [8]:

$$K = N / J_{\Sigma} .$$

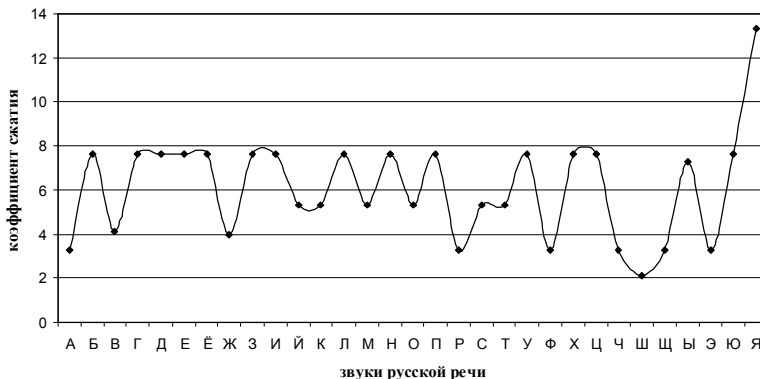


Рис. Степень сжатия для различных звуков русской речи
($c = 0.92$, $m = 0.92$)

Выводы. Проведенные вычислительные эксперименты показали высокую эффективность разработанного метода с позиции сжатия речевых данных при сохранении приемлемого для пользователя качества воспроизведения. Предлагаемый подход к сжатию речевых сигналов в среднем позволяет достичь сжатия в 6 – 7 раз без существенной потери качества воспроизведения. Этот показатель можно увеличить за счет использования известных алгоритмов обнаружения и удаления пауз.

Список литературы: 1. Жиялков Е.Г. Методы обработки речевых данных в информационно-телекоммуникационных системах на основе частотных представлений: монография / Е.Г. Жиялков, С.П. Белов, Е.И. Прохоренко // Белгород: Изд-во БелГУ, 2007. – 136 с. 2. Шелухин О.И. Цифровая обработка и передача речи / О.И. Шелухин, Н.Ф. Лукьянцев. – М.: Радио и связь, 2000. – 456 с. 3. Жиялков Е.Г. Вариационные методы анализа и построения функций по эмпирическим данным: монография / Е.Г. Жиялков. – Белгород: Изд-во БелГУ, 2007. – 160 с. 4. Прохоренко Е.И. Новый метод оптимального субполосного преобразования в задаче сжатия речевых данных / Е.И. Прохоренко, А.В. Болдышев, А.А. Фирсова, А.В. Эсауленко // Журнал "Вопросы радиоэлектроники", серия ЭВТ. – Вып. №1. – М.: 2010. – С. 49 – 55. 5. Ковалгин Ю.А. Цифровое кодирование звуковых сигналов / Ю.А. Ковалгин, Э.И. Вологдин. – Изд-во: Корона Принт, 2004. – 240 с. 6. Болдышев А.В. О различиях распределения энергии звуков русской речи и шума / А.В. Болдышев, А.А. Фирсова // Материалы 12-ой Международной конференции и выставки "Цифровая обработка сигналов и её применение. – "DSPA'2010". – М.: 2010. – С. 204 – 207. 7. Гантмахер Ф.Р. Теория матриц / Ф.Р. Гантмахер. – М.: Физматлит, 2004. – 560 с. 8. Сизиков В.С. Математические методы обработки результатов измерений: учебник для вузов / В.С. Сизиков. – СПб.: Политехника, 2001. – 240 с.

УДК 621.391

Про виборче перетворення частотних компонент мовних сигналів в завданні стиснення / Жіялков Е.Г., Болдышев А.В., Курлов А.В., Фірсова А.А., Есауленко А.В. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 37 – 43.

У статті приведені результати обчислювальних експериментів по апробації методу стиснення мовних даних на основі виборчого перетворення частотних компонент мовних сигналів, отриманих за допомогою нового методу субсмугового частотного аналізу/синтезу. Іл.: 1. Табл.: 2. Бібліогр.: 8 назв.

Ключові слова: частотні компоненти мовних сигналів, стиснення мовних даних, субполосний частотний аналіз/синтез.

UDK 621.391

About electoral transformation frequency component of vocal signals in the task of compression / Zhilyakov E.G., Boldyshev A.V., Kurlov A.V., Firsova A.A., Esaulenko A.V. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 37 – 43.

In this article the results of calculable experiments are resulted on approbation of method of compression of vocal data on the basis of electoral transformation frequency component of vocal signals, got by the new method of subband frequency analysis/synthesis. Figs.: 1. Tabl.: 2. Refs.: 8 titles.

Keywords: frequency components of speech signals, compress speech data, sub-band frequency analysis/synthesis.

Поступила в редакцію 03.02.2011

Е.Г. ЖИЛЯКОВ, д.т.н., проф. НИУ "БелГУ", Белгород,
Е.И. ПРОХОРЕНКО, к.т.н., доц. НИУ "БелГУ", Белгород,
А.В. БОЛДЫШЕВ, аспирант НИУ "БелГУ", Белгород,
А.А. ФИРСОВА, аспирантка НИУ "БелГУ", Белгород,
М.В. ФАТОВА, магистр НИУ "БелГУ", Белгород

СЕГМЕНТАЦИЯ РЕЧЕВЫХ СИГНАЛОВ НА ОСНОВЕ АНАЛИЗА ОСОБЕННОСТЕЙ РАСПРЕДЕЛЕНИЯ ДОЛЕЙ ЭНЕРГИИ ПО ЧАСТОТНЫМ ИНТЕРВАЛАМ¹

В статье рассмотрены существующие подходы к сегментации речевых сигналов. Представлены результаты оценки особенностей распределения энергии речевых сигналов. Предлагается способ сегментации речи на основе учета особенностей распределения долей ее энергии по частотным интервалам. Ил.: 3. Табл.: 1. Библиогр.: 8 назв.

Ключевые слова: сегментация, речевые сигналы, распределение долей энергии.

Постановка проблемы и анализ литературы. Одной из проблем обработки речевых сигналов является сегментация сигнала на звуки. Точность алгоритмов сегментации определяет надежность и эффективность использования в дальнейшем таких алгоритмов, как распознавание речи, синтез, сжатие.

Сегментация речи – это процесс поиска границ между элементами речевого сообщения: фразами, словами, слогами, фонемами. Сегментация может осуществляться вручную или автоматически. Сегментация вручную является надежным, но трудоемким способом, особенно если это касается большого объема обрабатываемой информации. Также сегментация вручную невозможна при реализации обработки сигнала в режиме реального времени [1, 2, 3].

Наиболее интересной представляется автоматическая сегментация речевых сигналов. Эффективность алгоритма сегментации определяется точностью определения границы между различными звуками. Существует два подхода автоматической сегментации. Первый состоит в том, что при обработке речевого сигнала известна последовательность фонем, необходимо только определить границы между ними. Второй подход не использует априорную информацию о речевом сообщении, и сегментация осуществляется на основе изменения характера речевого сигнала. Можно выделить также третий подход, объединяющий два

¹ Исследования выполнены при поддержке гранта РФФИ № 0-07-00326-а.

перечисленных: использование априорной информации и анализ изменений характера сигнала [1, 4].

Все существующие алгоритмы сегментации речи основываются на статических или динамических характеристиках речи. Анализ статических характеристик не всегда приводит к точной сегментации. Оценка динамических характеристик сигнала позволяет увеличить точность сегментации [1, 4].

Одними из основных методов сегментации являются [1 – 4]:

- 1) сегментация по усредненному нормированному спектру;
- 2) сегментация по динамическим детекторам;
- 3) сегментация по корреляции между равноотстоящими спектрами;
- 4) сегментация с использованием дискретного вейвлет-преобразования.

Цель статьи – разработка нового метода сегментации речевых сигналов, основанного на учете особенностей распределения долей энергии по частотным интервалам для каждого звука речи [5].

Результаты исследований. Каждый звук имеет свое особенное распределение долей энергии по частотному диапазону. Звуки, соответствующие буквам русского алфавита, сосредоточены в узком частотном интервале, в то время как спектр шума распределен по всей области частот более равномерно. Эту особенность можно использовать для определения начала и конца слова или словосочетания.

На всей длительности любого звука можно выделить несколько участков, имеющих некоторые особенности. К таким участкам относятся начало звука, середина и конец. Это вызвано тем, что в слитной речи происходит переход одного звука в другой, для этого речевой аппарат человека некоторое время перестраивается. Для некоторых звуков, соответствующих таким буквам, как "е", "ё", "ю", "я", можно выделить большее число участков. Это связано с тем, что они состоят из двух звуков, плавно переходящих из одного в другой. Каждый из участков имеет свои особенности распределения энергии. Распределение энергии разных участков одного фрагмента звука отличается незначительно. Эти особенности могут быть использованы для сегментации речи.

Анализ распределения энергии отрезков сигналов по частотным интервалам предлагается проводить на основе точного метода [7]. В этом случае полный набор долей энергии отрезка сигнала определяется следующим образом:

$$P_r = \vec{x}^T A_r \vec{x}, \quad (1)$$

где $\vec{x}$ – анализируемый отрезок сигнала; r ($r = 1, \dots, R$) – номер частотного интервала; R – число частотных интервалов, на которые разбивается частотная ось; $A_r = \{a_{ik}^r\}$ – субполосная матрица, определяемая для каждого из R частотных интервалов, с элементами вида

$$a_{ik}^r = (\sin(v_{r+1}(i-k)) - \sin(v_r(i-k))) / (\pi(i-k)), \quad i, k = 1, \dots, N, \quad (2)$$

где v_r, v_{r+1} – границы r -го частотного интервала, причем:

$$0 \leq v_r < v_{r+1} \leq \pi, \quad v_{r+1} - v_r = \pi / R, \quad r = 1, \dots, R, \quad (3)$$

N – длительность анализируемого отрезка речевого сигнала.

Одной из характеристик, отражающей особенности звуков русской речи, является величина частотной концентрации, которая оценивается с использованием следующего выражения [8]:

$$W_{NR} = f_{NR}^m / R, \quad (4)$$

где f_{NR}^m – минимальное количество частотных интервалов (частотная концентрация), в которых сосредоточена заданная доля энергии m звукового отрезка, т.е.

$$f_{NR}^m = \min d_{NR}^m. \quad (5)$$

Здесь выполняется неравенство

$$\sum_{k=1}^{d_{NR}^m} P_{(k),N} \geq m \|\vec{x}_N\|^2 = m \sum_{i=1}^N x_i^2, \quad (6)$$

где $\vec{x}_N$ – анализируемый отрезок сигнала; m – заданное значение доли энергии сигнала; $P_{(k),N}$ – упорядоченные по убыванию доли энергий сигнала, попадающие в заданные частотные интервалы, т.е.

$$P_{(k),N} \in \{P_{rN}, r = 1, \dots, R\}, \quad P_{(k+1),N} \leq P_{(k),N}, \quad k = 1, \dots, R. \quad (7)$$

Для оценки возможности сегментации с использованием свойства частотной концентрации звуков русской речи было проведено большое количество экспериментов по оценке частотной концентрации различных фонем при различных значениях числа интервалов, на которые разбивается ось частот ($R = 4, 8, 16, 32, 64$), и значениях длины окна анализа ($N = 64, 128, 256$). В качестве исходного материала был использован фрагмент лекции, содержащий большое количество

различных фонем, записанный с частотой дискретизации $f_d = 8$ кГц с 16-битовым представлением в монорежиме.

Результаты экспериментов показали, что увеличение количества интервалов, на которые разбивается частотная ось, приводит к уточнению величины частотной концентрации отрезка сигнала.

В таблице представлены результаты оценки величины частотной концентрации для различных звуков русской речи.

Таблица. Распределение долей частотных интервалов, в которых сосредоточено 95% энергии при $N = 128$, $R = 32$ для различных звуков русской речи

гласные										
звук	а	е	ё	и	о	у	ы	э	ю	я
W_{NR}	0,31	0,13	0,09	0,09	0,19	0,09	0,16	0,31	0,13	0,09
сонорные согласные										
звук	й		л		м		н		р	
W_{NR}	0,19		0,19		0,19		0,16		0,34	
звонкие согласные										
звук	б		в		г		д		ж	
W_{NR}	0,22		0,28		0,19		0,16		0,25	
глухие согласные										
звук	к	п	с	т	ф	х	ц	ч	ш	щ
W_{NR}	0,22	0,16	0,25	0,28	0,16	0,25	0,19	0,44	0,47	0,34

Из таблицы видно, что для некоторых гласных и согласных звуков величины частотной концентрации совпадают. Особенно это проявляется для сонорных согласных.

На рис. 1 представлен фрагмент речевого сигнала, соответствующий звукосочетанию "шеч", выделенному из слова "шахматно-шашечный". Звук разбит на 16 равных окон анализа по 128 отсчетов.

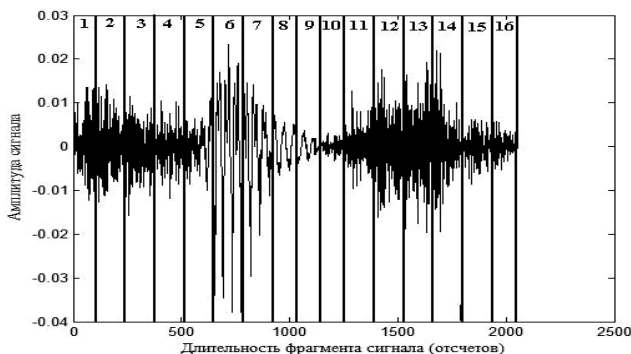


Рис. 1. Фрагмент речевого сигнала (слог "шеч" – безударный)

На рис. 2 представлен график распределения величины частотной концентрации для звуко сочетания "шеч" из слова "шахматно-шашечный", при длине окна анализа $N = 128$ и количестве частотных интервалов $R = 32$, для различных значений доли энергии.

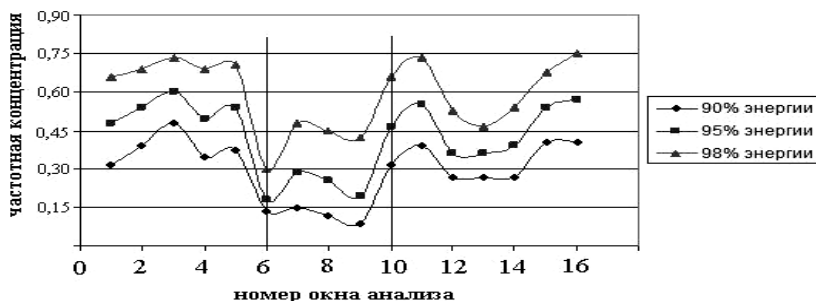


Рис. 2. Распределение величины частотной концентрации для звуко сочетания "шеч" ($N = 128$ и $R = 32$)

Анализ результатов экспериментов, представленных на рис. 2, показывает, что при выборе 95% энергии при переходе от 5 окна к 6-му, а также от 9-го к 10-му, наблюдается наибольшее изменение величины частотной концентрации. Эта особенность может быть использована для определения перехода между звуками. На рис. 1 видно, что окна 5 и 9 соответствуют переходу соответственно от звука "ш" к звуку "е", и от звука "е" к звуку "ч". Таким образом, увеличение разности частотной концентрации между соседними окнами может быть использовано для определения границы перехода между звуками.

На рис. 3 представлены результаты сегментации речевого сигнала с использованием различия величины частотной концентрации для различных звуков речи.

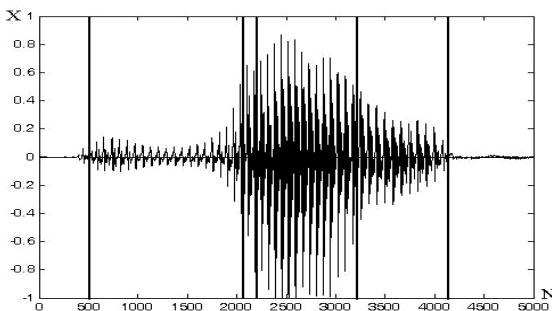


Рис. 3. Разбиение слова "заяц" на сегменты при $N = 128$, $R = 32$

Анализ рис. 3 показывает, что предлагаемый алгоритм позволил отделить звуки "з" и "а", "а" и "я", "я" и "ц". Проявилась дополнительная граница на фрагменте, соответствующем переходу между звуками "з" и "а".

Выводы. В результате проведенных экспериментов было выявлено, что длина фонем изменяется в пределах 1000 – 4000 отсчетов и зависит от типа звуко сочетания: открытый слог, закрытый слог, ударный слог, безударный слог и т.д. Сравнение ударных и безударных слогов показало, что если гласный стоит под ударением, то длительность слога возрастает примерно в 1,25 раза. Спектры соответствующих звуков в ударном и безударном слогах отличаются незначительно. Использование представленного метода позволяет выявить место перестройки речевого аппарата с согласной на гласную и с гласной на согласную. Таким образом, данный метод может быть использован как один из элементов алгоритма сегментации речевого сигнала на отдельные звуки.

Список литературы: 1. *Сорокин В.Н.* Сегментация речи на кардинальные элементы / *В.Н. Сорокин, А.И. Цыплихин* // Информационные процессы. – 2006. – Т. 6. – № 3. – С. 177–207. 2. *Дремин И.М.* Вейвлеты и их использование / *И.М. Дремин, О.В. Иванов, В.А. Нечитайло* // Успехи физических наук. – 2001. – Т. 171. – №5. – С. 465–500. 3. *Ермоленко Т.Н.* Алгоритмы сегментации с применением быстрого вейвлет-преобразования / *Т.Н. Ермоленко, В.И. Шевчук* // Статьи, принятые к публикации на сайте международной конференции Диалог'2003. www.dialog-21.ru. 4. *Сорокин В.Н.* Сегментация и распознавание гласных / *В.Н. Сорокин, А.И. Цыплихин* // Информационные процессы. – 2004. – Т. 4. – № 2 – С. 202–220 5. *Жиляков Е.Г.* Методы обработки речевых данных в информационно-телекоммуникационных системах на основе частотных представлений: монография / *Е.Г. Жиляков, С.П. Белов, Е.И. Прохоренко*. – Белгород: Изд-во БелГУ, 2007. – 136 с. 6. *Шелухин О.И.* Цифровая обработка и передача речи / *О.И. Шелухин, Н.Ф. Лукьянцев*. Под ред. О.И. Шелухина. – М.: Радио и связь, 2000. – 456 с. 7. *Жиляков Е.Г.* Вариационные методы анализа и построения функций по эмпирическим данным: монография / *Е.Г. Жиляков*. – Белгород: Изд-во БелГУ, 2007. – 160 с. 8. *Фирсова А.А.* О различиях распределения энергии звуков русской речи и шума / *А.В. Болдышев, А.А. Фирсова* // Материалы 12-ой Международной конференции и выставки "Цифровая обработка сигналов и её применение. – "DSPA'2010". – Москва. – 2010. – С. 204–207.

УДК 621.391

Сегментація мовних сигналів на основі аналізу особливостей розподілу часток енергії за частотним інтервалам / **Є.Г. Жиляков, Є.І. Прохоренко, А.В. Болдышев, А.А. Фірсова, М.В. Фатова** // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 44 – 50.

У статті розглянуті існуючі підходи до сегментації мовних сигналів. Представлені результати оцінки особливостей розподілу енергії мовних сигналів. Пропонується спосіб сегментації мови на основі врахування особливостей розподілу часток її енергії за частотними інтервалами. Іл.: 3. Табл.: 1. Бібліогр.: 8 назв.

Ключові слова: сегментація, мовні сигнали, розподіл часток енергії.

UDC 621.391

Segmentation of speech signals based on analysis of distribution of shares on energy frequency band/ Zhilyakov E.G., Prokhorenko E.I., Boldyshev A.V., Firsov A.A., Fatova M.V. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 44 – 50.

The paper considers existing approaches to segmentation of speech signals. Presents the results of evaluation of the features of the energy distribution of speech signals. Provides a method of speech segmentation on the basis of features of the distribution of shares of its energy on the frequency range. Figs.: 3. Tabl.: 1. Ref.: 8 titles.

Key words: segmentation, voice signals, the distribution of shares of energy.

Поступила в редакцию 01.02.2011

В.Г. ИВАНОВ, д.т.н., проф., зав. каф. НЮАУ им. Я. Мудрого, Харьков,
Ю.В. ЛОМОНОСОВ, к.т.н., доц. НЮАУ им. Я. Мудрого, Харьков,
М.Г. ЛЮБАРСКИЙ, д.ф.-м.н., проф. НЮАУ им. Я. Мудрого, Харьков,
Н.А. КОШЕВА, к.т.н., доц. НЮАУ им. Я. Мудрого, Харьков,
М.В. ГВОЗДЕНКО, ст. преп. НЮАУ им. Я. Мудрого, Харьков,
Н.И. МАЗНИЧЕНКО, ст. преп. НЮАУ им. Я. Мудрого, Харьков

СЖАТИЕ ДАННЫХ В СИСТЕМЕ ХААРА С УЧЕТОМ ОБЪЕМА ПЕРВИЧНОЙ ДИСКРЕТИЗАЦИИ

Формализован процесс вычисления коэффициентов Хаара непрерывных функций с учетом объема первичной дискретизации сигналов и получены сравнительные оценки эффективности этих преобразований в задачах сжатия сообщений. Ил.: 2. Библиогр.: 15 назв.

Ключевые слова: сжатие; система Хаара. дискретизация сигналов.

Постановка проблемы и анализ литературы. Методы сжатия данных стали перманентной составляющей практически всех алгоритмов и компьютерных систем, и сетей хранения, обработки, передачи или поиска мультимедийной информации [1 – 6].

При этом одной из основных проблем проектирования устройств сжатия данных является проблема эффективного дискретного представления непрерывных сообщений $f(t)$ как функции дискретного времени $f(t_i)$, которая характеризуется совокупностью координат V_k , полученных в результате разложения контролируемого параметра по системе каких-либо ортогональных базисных функций.

При этом повышение коэффициента сжатия в устройствах связано со свойством сходимости аппроксимирующего ряда из этих коэффициентов. Чем лучше сходимость, тем меньше членов ряда, а следовательно, и координат требуется для аппроксимации исходной функции при заданной точности и приемлемой сложности вычислений.

Таким критериям отвечают алгоритмы сжатия исходных данных на основе базовых вейвлетов Хаара.

Идея вейвлетной декомпозиции сигналов является основой для многих алгоритмов и методов фильтрации, кодирования или распознавания данных [2, 6 – 11]. Однако кодирование сигналов на основе вейвлет-преобразований Хаара с учетом как объема первичного

квантования так и аппроксимирующих свойств рядов Хаара, исследовано не достаточно.

Цель статьи. Формализовать процесс вычисления коэффициентов Хаара непрерывных функций с учетом объема первичной дискретизации и получить сравнительные оценки эффективности этих преобразований в задачах сжатия данных.

Решение задачи. Известно [12], что любая непрерывная на отрезке $[0, 1]$ функция $f(x)$, отображающая сигнал, подлежащий обработке, разлагается в равномерно сходящийся ряд

$$f(x) = \sum_{k=1}^{\infty} C_k \chi_k(x), \quad (1)$$

где C_k и $\chi_k(x)$ – соответственно коэффициенты и функции Хаара.

Легко показать, что учитывая свойства системы функций $\chi_k(x)$, коэффициенты Хаара могут быть вычислены по формуле:

$$C_{mj} = 2^{\frac{m-1}{2}} \left[\int_{l_{mj}^-} f(x) dx - \int_{l_{mj}^+} f(x) dx \right], \quad (2)$$

где $l_{mj}^{\mp}$ – соответственно левая и правая половина двоичного отрезка определения функций Хаара.

При формировании сообщения, содержащего "n" координат Хаара, оценка поведения параметра без учета влияния шумов первичного квантования получается в виде:

$$f^*(x) = \sum_{k=1}^n C_k \chi_k(x). \quad (3)$$

Используя известное положение теории квадратурных формул [13], заменим интегралы в (2) их частичными суммами:

$$S_M = \sum_{n=1}^M f(x_n) \Delta x_n, \quad (4)$$

где $\Delta x_n = \frac{B-A}{M}$, $x_n = A + \frac{B-A}{M}n$, а A и B – соответственно верхний и нижний пределы интегрирования.

Тогда выражение (2) можно переписать следующим образом:

$$C_{mj} = 2^{\frac{m-1}{2}} \left[\sum_{n=0}^{M-1} f(x_n) \Delta x_n - \sum_{n=0}^{M-1} f(x_n^*) \Delta x_n^* \right]. \quad (5)$$

Так как для системы Хаара

$$\Delta x_n = \frac{2j-1}{2^m} - \frac{2(j-1)}{2^m} \quad \text{и} \quad \Delta x_n^* = \frac{2j}{2^m} - \frac{2j-1}{2^m},$$

то (5) запишется в виде

$$C_{mj} = 2^{-\frac{m-1}{2}} \left[\frac{2^{j-1}}{2^m} \frac{2^{j-2}}{M} \sum_{n=0}^{M-1} f\left(n \frac{2^{j-1}}{2^m} + \frac{2^{j-2}}{2^m}\right) - \frac{2^j}{2^m} \frac{2^{j-1}}{M} \sum_{n=0}^{M-1} f\left(n \frac{2^m}{2^m} + \frac{2^{j-1}}{2^m}\right) \right], \quad (6)$$

где M – число отсчетов функции на $l_{mj}^{\mp}$ интервалах.

После соответствующих преобразований выражение (6) приводится к виду:

$$C_{mj} = \frac{2^{-\frac{m+1}{2}}}{M} \sum_{n=0}^{M-1} \left[f\left(y - \frac{1}{2^m}\right) - f(y) \right], \quad (7)$$

$$\text{где } y = \frac{\frac{n}{M} + 2j-1}{2^m}.$$

Полученная формула (7) позволяет вычислить коэффициенты Хаара функции $f(x)$. Причем, частота первичной дискретизации, т.е. расстояние между f_n и f_{n+1} , будет разной при вычислении коэффициентов с различными порядковыми номерами, а число выборок на двоичных отрезках $l_{mj}^{\mp}$ будет постоянно при любом m .

Используя среднеквадратичный критерий, запишем мощность ошибки при аппроксимации сигнала $f(x)$ рядом Хаара:

$$\varepsilon = \frac{1}{T} \int \left[f(x) - \sum_{k=1}^N C_k \chi_k(x) \right]^2 dt. \quad (8)$$

После преобразований будем иметь:

$$\varepsilon = P - \sum_{k=1}^N \left[\frac{2^{-(m+1)}}{M} \sum_{n=0}^{M-1} f^2\left(y - \frac{1}{2^m}\right) - 2 \sum_{n=0}^{M-1} f\left(y - \frac{1}{2^m}\right) \cdot \sum_{n=0}^{M-1} f(y) + \sum_{n=0}^{M-1} f^2(y) \right], \quad (9)$$

где P определяет мощность сигнала.

Выражение (9) позволяет оценить ошибку представления сигнала в базисе Хаара при конечном числе членов ряда (1) с учетом шумов первичного квантования, т.е. с различным числом отсчетов функции, необходимых для определения коэффициентов Хаара.

Однако, применение на практике формулы (7) несет в себе ряд неудобств, так как вычисление коэффициентов Хаара с большими порядковыми номерами, т.е. когда $l_{mj}^{\mp}$ мало, не выгодно с точки зрения вычислительных затрат, а с технической стороны требует перестройки частоты временного квантователя.

Если функционально изменять величину M в (7) при переходе от значения m к $(m+1)$ таким образом, чтобы ошибка ε в выражении (9) оставалась в допустимых пределах, то число вычислительных операций может быть снижено.

Запишем рекуррентное соотношение для значения M при переходе к вычислению коэффициента Хаара с большим порядковым номером в виде $\frac{M}{2^m}$ и представим после преобразований выражение (7) в виде

$$C_{mj} = \frac{2^{\frac{3m-1}{2}}}{M} \sum_{n=0}^{\frac{M-1}{2^m}} \left[f\left(z - \frac{1}{2^m}\right) - f(z) \right], \quad (10)$$

$$\text{где } z = \frac{\frac{n}{M} 2^m + 2j - 1}{2^m}.$$

Полученная формула (10) позволяет существенно снизить число отсчетов, а следовательно, и вычислительные затраты при ортогональном разложении сигналов в ряд Хаара.

Если M четное, то использование выражения (10) позволяет нелинейно изменять число отсчетов функции $f(x)$ на двоичных отрезках, оставляя длительность между отсчетами постоянной, что очень существенно в практическом плане, тогда как вычисления по формуле (7) требуют изменения частоты выборки на каждом двоичном отрезке. Если, например, на всем интервале определения сигнала имеется 128 отсчетов функции, то при вычислении первого коэффициента по формуле (10) будут учитываться все отсчеты, при вычислении второго коэффициента – 64 отсчета и т.д.

Как уже отмечалось, эффективность устройств сжатия данных существенным образом зависит от экстремальных свойств ортогональных рядов по выбранной системе базисных функций. Учитывая тот факт, что не существует базиса, который обеспечивал бы наилучшую сходимость для всех классов аппроксимируемых функций, можно сказать, что на передний план выдвигаются такие свойства базисных функций, которые обеспечивают минимальную сложность вычислений коэффициентов.

В этой связи перспективным видится использование рассматриваемого базиса Хаара.

Число отсчетов при представлении сигналов в этом базисе, а следовательно, и вычислительных операций по формулам (7) и (10) будет равно соответственно $M(2^{m+1}-1)+2^m$ и $(4M+2^m)$, где m определяет высший порядок используемых функций Хаара.

В случае использования тригонометрического базиса Фурье для представления сигнала $f(x)$, объём вычислительных операций составит M^2 умножений и столько же сложений.

Если предварительно вычислить значение суммы в (7) и (10) на двоичных отрезках минимальной длительности, то коэффициенты можно определить с использованием алгоритма быстрого преобразования Хаара [14]. Причем вычислительные затраты в этом случае составят соответственно $M \cdot 2^m + (2^m - 1)$ и $M + 2(2^m - 1)$ операций типа сложения – вычитания.

Один из наиболее простых вариантов сжатия сигнала, заключается в ограничении его спектра [15].

При разложении непрерывного сигнала $S(t)$ по системе функций $\{\eta(k, t)\}$ образуется спектр $\{C(k)\}$.

Ограничение спектра заключается в том, что фильтр пропускает все составляющие спектра $\{C_k\}$ с порядком не выше $k \leq n$ и задерживает все составляющие при $k > n$.

Сигнал на выходе фильтра и выходной спектр связаны соотношениями:

$$\begin{aligned} S(t_1) &= \sum_{k=0}^n \tilde{C}(k) \eta(k, t_1), \\ C(k) &= \frac{1}{P_r} \frac{2}{T} \int_0^{\frac{T}{2}} S(t) \eta(k, t) dt, \end{aligned} \quad (11)$$

где P_r означает мощность базисных функций.

Таким образом, процесс фильтрации можно записать в виде

$$\tilde{S}(t_1) = \sum_0^{\frac{T}{2}} S(t) D_n(t, t_1) dt, \quad (12)$$

где $D_n(t, t_1)$ – ядро Дирихле, которое полностью определяется базисом $\{\eta(k, t)\}$ и фильтром (числом n удерживаемых членов ряда Фурье). Эффективность фильтра будет тем выше, чем меньше значение n в

выражении (12) при одинаковых среднеквадратических ошибках восстановления. Однако, такой критерий не учитывает вычислительных особенностей базисов таких, например, как Фурье и Хаара.

В этой связи представляется небезынтересным введение такого критерия, как зависимость среднеквадратичной ошибки аппроксимации (8) от времени обработки.

На рис. 1 и 2 приведены рассчитанные на ЭВМ графики зависимости среднеквадратичной ошибки представления реализаций стационарного процесса с периодической функцией корреляции от числа удерживаемых членов ряда фильтра Фурье и Хаара, которые показывают сравнительную эффективность обработки при различном числе первичных отсчетов M . Для рис. 1 и 2 приняты следующие обозначения: 1 – базис Фурье; кривые 2, 3, 4 – базис Хаара при различных значениях M в выражениях (7) и (10), причем $M(2) > M(3) > M(4)$.

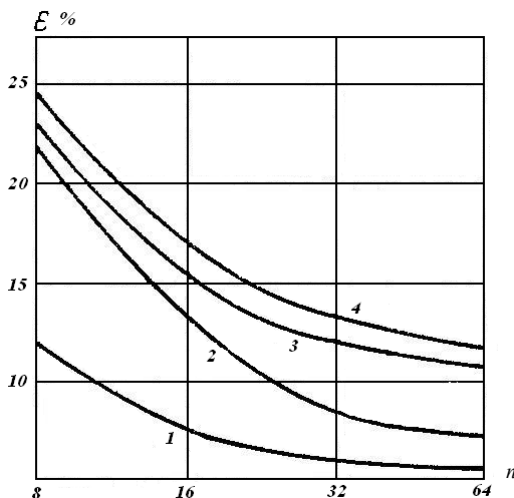


Рис. 1. Зависимости среднеквадратичной ошибки представления сигналов коэффициентами ряда в базисе Фурье (1) и в базисе Хаара (2, 3, 4)

Анализ полученных результатов показывает, что эффективность обработки в базисе Хаара (рис. 2, кривая 2) с большим значением первичных отсчетов M выше, чем когда используется меньшее число первичных отсчетов. Эта эффективность растет с увеличением числа удерживаемых коэффициентов ряда.

Представление сигналов в базисе Хаара с использованием нелинейного изменения числа отсчетов M согласно выражению (10)

(рис. 2, кривая 4), при всех значениях нормированного времени t дает ошибку фильтрации почти в два раза меньшую, чем использование представления в соответствии с выражением (7). Т.е. при одинаковом значении нормированного времени t , число удерживаемых коэффициентов ряда Хаара будет большим при использовании формулы (10). Причем кривая 4 на рис. 2 имеет более крутой спад в начале нормированного времени обработки, что говорит о низкочастотном характере спектра Хаара. Графики под номером 5 и 6 на рис. 2 показывают эффективность обработки в базисе Хаара с использованием алгоритма быстрого преобразования и выражений (7) и (10) соответственно.

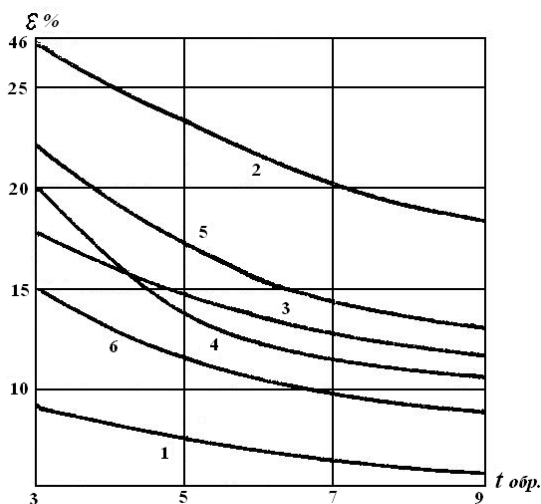


Рис. 2. Зависимость среднеквадратичной ошибки фильтров Фурье (1) и Хаара (2, 3, 4, 5) от времени обработки

Эти результаты распространяются на класс случайных стационарных процессов с периодической функцией корреляции, т.е. предложенное эмпирическое нелинейное преобразование отсчетов $M = \varphi(m)$ гарантирует неувеличение среднеквадратичной ошибки фильтрации ε .

Выводы. Практическая ценность полученных результатов заключается в том, что они позволяют вычислить коэффициенты и оценить ошибку представления сигналов рядом Хаара как с учетом числа первичных отсчетов, так и числа удерживаемых коэффициентов ряда на

інтервалі аналізу і дозволяють знайти прийнятний компроміс між заданим часом і якістю обробки.

Як і слід було очікувати, для даного класу сигналів фільтр Фур'є виявився переважніше фільтра Хаара, так як тригонометричні функції з кратними періодами є оптимальним базисом в разі стаціонарного процесу з періодичною кореляційною функцією.

Список літератури: 1. *Мюррей Д.* Енциклопедія форматів графічних файлів / *Д. Мюррей, Уан Райнер.* – К.: Видавнича група ВНУ, 1997. – 672 с. 2. *Гонсалес Р.* Цифрова обробка зображень / *Р. Гонсалес, Р. Вудс.* – М.: Техносфера, 2005. – 1072 с. 3. *Соломон Д.* Сжатие данных, изображений и звука / *Д. Соломон.* – М.: Техносфера, 2004. – 368 с. 4. *Ричардсон Ян.* Виточковий кодирование. H. 264 и MPEG-4 – стандарты нового поколения / *Ян. Ричардсон.* – М.: Техносфера, 2005. – 368 с. 5. *Иванов В.Г.* Фурье и вейвлет-анализ изображений в плоскости JPEG технологий / *В.Г. Иванов, М.Г. Любарский, Ю.В. Ломоносов* // Проблемы управления и информатики. – Київ, 2004. – № 5. – С. 111-124. 6. *Форсайт Д.* Компьютерное зрение. Современный подход / *Д. Форсайт, Д. Понс.* – М.: Вильямс, 2004. – 928 с. 7. *Иванов В.Г.* Сокращение содержательной избыточности изображений на основе классификации объектов и фона / *В.Г. Иванов, М.Г. Любарский, Ю.В. Ломоносов* // Проблемы управления и информатики. – Київ, 2007. – № 3. С. 93-102. 8. *Уэлстид С.* Фракталы и вейвлеты для сжатия изображений в действии: Учеб. пособ. / *С. Уэлстид.* – М.: Триумф, 2003. – 320 с. 9. *Кинтцель Тим.* Руководство программиста по работе со звуком / *Тим Кинтцель.* – М.: ДМК Пресс, 2000. – 432 с. 10. *Кравченко В.Ф.* Wavelet-системы и их применение в обработке сигналов / *В.Ф. Кравченко, В.А. Рвачев* // Зарубежная радиоэлектроника. – М.: Радио и связь, 1996. – № 4. – С. 3-20. 11. *Иванов В.Г.* Применение вейвлет-анализа к сжатию звуковых сигналов / *В.Г. Иванов, М.Г. Любарский, Ю.В. Ломоносов* // Вісник Національного технічного університету "Харківський політехнічний інститут". – Харків: НТУ "ХПІ", 2003. – Т. 1. – № 7. – С. 39-50. 12. *Соболь И.М.* Многомерные квадратичные формулы и функции Хаара / *И.М. Соболь.* – М.: Наука, 1970. – 288 с. 13. *Микеладзе Ш.Е.* Численные методы математического анализа / *Ш.Е. Микеладзе.* – М.: Гос. изд-во технико-теоретич. лит., 1953. – 527 с. 14. *Иванов В.Г.* Формальное описание дискретных преобразований Хаара / *В.Г. Иванов* // Проблемы управления и информатики. – 2003. – № 5, – С. 68-75. 15. *Трахтман А.М.* Введение в обобщенную спектральную теорию сигналов / *А.М. Трахтман.* – М.: Сов. Радио, 1972. – 352 с.

УДК 004.627

Стиснення даних в системі Хаара з урахуванням об'єму первинної дискретизації / *Иванов В.Г., Ломоносов Ю.В., Любарский М.Г., Кошева Н.А., Гвозденко М.В., Мазниченко Н.І.* // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 51 – 59.

Формалізований процес обчислення коефіцієнтів Хаару безперервних функцій з урахуванням об'єму первинної дискретизації сигналів і отримані порівняльні оцінки ефективності цих перетворень в завданнях стиснення повідомлень. Іл.: 2. Бібліогр.: 15 назв.

Ключові слова: стиснення, система Хаару, дискретизація сигналів.

UDC 004.627

Compression of data in the system of Haar taking into account the volume of primary digitization / *Ivanov V.G., Lyubarsky M.G., Lomonosov U.V., Kosheva N.A., Gvozdenko M.V., Maznichenko N.I.* // Herald of the National Technical University "KhPI".

Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 51 – 59.

The process of calculation of coefficients of Haar of continuous functions is formalized taking into account the volume of primary digitization of signals and the comparative estimations of efficiency of these transformations are got to the tasks of compression of reports. Figs.: 2. Refs.: 15 titles.

Keywords: compression, system of Haar, digitization of signals.

Поступила в редакцию 10.02.2011

Д.Е. ИВАНОВ, к.т.н., доц., с.н.с. отдела теории управляющих систем ИПММ НАН Украины, Донецк

ПРИМЕНЕНИЕ АЛГОРИТМОВ СИМУЛЯЦИИ ОТЖИГА В ЗАДАЧАХ ИДЕНТИФИКАЦИИ ЦИФРОВЫХ СХЕМ

В статье рассматривается опыт применения эволюционного алгоритма симуляции отжига к решению задач идентификации, возникающих при проектировании цифровых схем. Описывается общая структура алгоритма симуляции отжига, его компоненты решения задач построения идентифицирующих последовательностей и их оптимизации. Проводится сравнительный анализ эффективности с генетическими алгоритмами. Ил.: 2. Табл.: 2. Библиогр.: 20 назв.

Ключевые слова: эволюционные алгоритмы, идентификация, цифровая схема, симуляция отжига, генетический алгоритм.

Постановка проблемы и анализ литературы. Эволюционные алгоритмы находят широкое применение в задачах идентификации цифровых схем [1 – 3]. Наибольшее внимание авторы уделяют генетическим алгоритмам (ГА) построения входных идентифицирующих последовательностей: разработаны генераторы проверяющих тестов [4, 5], инициализирующих последовательностей [6, 7] и верифицирующих эквивалентность двух заданных схем [8]. Преимуществом указанных выше алгоритмов является то, что они позволяют обрабатывать большие последовательностные схемы и получать приемлемые для разработчика результаты в терминах полноты и времени работы. Это свойство объясняется тем, что алгоритмы данного рода не рассматривают внутреннюю структуру схемы и не строят деревьев решений, а используют итеративный анализ свойств потенциальных решений. Также известны работы по оптимизации тестовых последовательностей, в частности, с целью уменьшения рассеивания тепла при самотестировании [9]. Известно, что все указанные выше задачи являются NP-полными [10].

В настоящее время усиливается интерес к другим эволюционным подходам, среди которых можно выделить алгоритм симуляции отжига (СО) [11 – 12]. Однако большинство исследователей ограничиваются применением данного алгоритма к решению модельных задач: задача о рюкзаке, задача раскроя и т.д. Автор данной статьи в последнее время разработал ряд практических алгоритмов решения задач, которые возникают при проектировании цифровых схем [13 – 15].

Целью данной работы является обобщение имеющегося опыта применения алгоритмов СО к решению задач идентификации цифровых

схем, а также сравнение их поисковых свойств на данных задачах с ГА.

Общая схема алгоритма симуляции отжига. Алгоритм СО впервые был предложен в виде статистической модели в [11] для нахождения состояния группы атомов в остывающем слитке. Применение алгоритма СО как оптимизирующей стратегии было начато после публикации [12], в которой авторы на модельных задачах показали, что данный подход является оптимизирующей стратегией в широком смысле.

Общая структура алгоритма СО приведена на рис. 1.

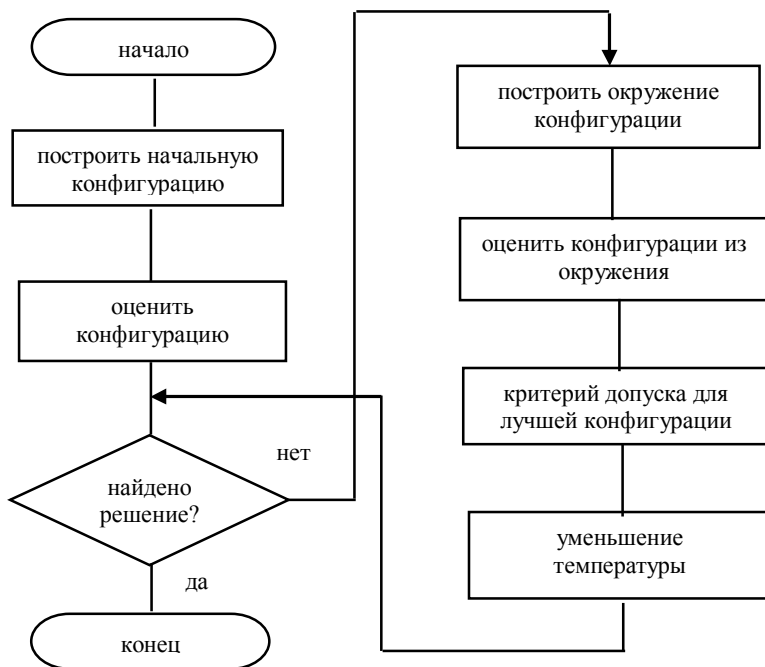


Рис. 1. Блок-схема алгоритма СО

Для задания алгоритма вводятся следующие понятия: конфигурация K_i шага i – потенциальное решение задачи; окружение $O(K_i)$ – соседние к данной конфигурации решения, которые строятся с помощью правил возмущения; оценка конфигурации $C(K_i)$, показывающая, насколько данная конфигурация хорошо решает поставленную задачу, обычно ищется решение с максимальной оценкой: $C(K_i) \rightarrow \max$;

расписание температур $\{T_j\}$ – убывающая последовательность. Понятие температуры играет особую роль в алгоритме. Её значение влияет на возможность принятия ухудшающих изменений: чем выше температура – тем чаще принимаются возмущения, которые ухудшают оценку конфигурации, и наоборот. Применение данного механизма позволяет алгоритму избегать сходимости к локальным экстремумам. В целом видно, что алгоритм представляет собой итеративный поиск, который может быть остановлен как при нахождении нужного решения, так и при достижении предельного числа итераций. Решить поставленную задачу с помощью алгоритма СО означает задать перечисленные компоненты и подобрать значения эвристических констант.

Алгоритмы СО построения входных идентифицирующих последовательностей. В зависимости от сложности решаемой задачи мы различаем одно- и двухуровневые схемы применения алгоритма СО. В первом случае происходит однократный вызов алгоритма СО для нахождения решения задачи. В случае двухуровневой схемы происходит итеративный процесс: определение локальной цели (верхний уровень) и вызов алгоритма СО для достижения локальной цели (нижний уровень). По одноуровневой схеме построены алгоритмы нахождения инициализирующих и верифицирующих эквивалентность последовательностей, по двухуровневой – генератор проверяющих тестов. Дадим более формальную постановку данных задач.

1) Построение инициализирующих последовательностей (ИП).

Пусть Q – множество всех состояний последовательностной схемы в трёхзначном алфавите моделирования $E_3 = \{0, 1, u\}$; пусть Z – множество всех возможных определённых состояний схемы, а Σ – множество всех возможных входных последовательностей s_i . Для выбранного трёхзначного моделирования $Z \subset Q$. Пусть $x \in Q$ начальное неопределённое состояние. Функция $F: Q \times \Sigma \rightarrow Q$ обозначает все состояния, достижимые схемой при поступлении на её вход последовательностей из Σ при использовании трёхзначного алфавита моделирования.

Определение 1. Последовательность $s \in \Sigma$ называется инициализирующей для заданной схемы, если в финальном состоянии $q = F(x, s)$ все состояния определены, т.е. $q \in Z$.

2) При построении последовательностей, проверяющих эквивалентность двух заданных схем, задача, по существу, сводится к доказательству несуществования входной последовательности, различающей две заданные схемы, т.е. ищется контрпример,

показывающий, что схемы различны.

Определение 2. Входная последовательность $s \in \Sigma$ называется различающей для цифровых схем A_0 и A_1 , когда выходные реакции схем A_0 и A_1 различны: $A_0(s) \neq A_1(s)$.

3) Задача построения тестов. Пусть задано исправное устройство A_0 и конечное множество его неисправных модификаций $A = \{A_1, A_2, \dots, A_n\}$, и $A_0 \neq A_i$ для $i = 1, \dots, n$. Традиционно, мы используем множество одиночных константных неисправностей.

Определение 3. Тестом, проверяющим все неисправности модификаций A , называется такая входная последовательность, которая является проверяющей для всех $A_i \in A$, $i = 1, \dots, n$.

Исходя из целей указанных задач – построение входных последовательностей – выбирается кодирование конфигураций и возмущающие операции (рис. 2). Отметим, что такое кодирование полностью соответствует кодированию особей в ГА, а возмущающие операции соответствуют операциям мутации в ГА.

Общим в данных задачах является то, что необходимо искать входную последовательность, для которой качество решения задачи проверяется путём моделирования работы цифровой схемы на данной последовательности. Для задачи построения ИП функция оценки имеет вид:

$$C(K_i) = C(s) = f(n_1, n_2, n_3) = (c_1 \times n_1 + c_2 \times n_2) \times c_3^{n_3},$$

где s – входная последовательность; n_1 – отношение числа инициированных триггеров к их общему числу; n_2 – активность схемы или число событий моделирования; n_3 – длина последовательности; c_1 , c_2 , c_3 – нормализующие константы.

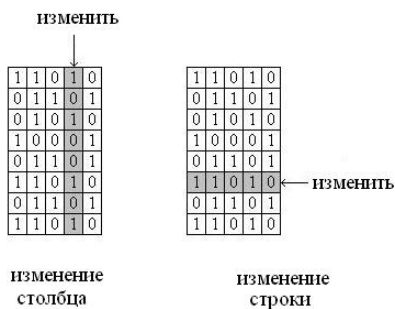
Для задачи верификации эквивалентности оценка вычисляется следующим образом:

$$C(K_i) = C(s) = f(n_1, n_2, n_3) = n_1 + c_1 \times n_2 + c_2 \times n_3,$$

где n_1 , n_2 , n_3 – число различных значений на внешних выходах, псевдовыходах, выходах логических элементов двух анализируемых схем соответственно.



а) конфигурация



б) возмущающие операции

Рис. 2. Представление решений и операции над ними

В задаче построения тестов проверяющая последовательность ищется для каждой из неисправностей целевого множества. Это заставляет применять в данном алгоритме двухуровневую схему. На верхнем уровне в множестве непроверенных неисправностей ищется такая неисправность $f_{\text{цел}}$, какую можно активировать с помощью псевдослучайной генерации, т.е. различие в поведении исправной и неисправной схемы распространено на внешние псевдовыходы схемы.

Далее на нижнем уровне для неисправности $f_{\text{цел}}$ с помощью алгоритма СО строится тест. Оценочная функция в этом случае имеет вид:

$$C(K_i, f_{\text{цел}}) = \sum_{i=1}^{i=\text{длина}} H^i \times h(v_i, f_{\text{цел}}),$$

где H^i – близкие к 1 константы, $h(v, f_{\text{цел}})$ определяет качество отдельного входного вектора v :

$$h(v, f_{\text{цел}}) = f_1(v, f) + c_1 \times f_2(v, f) = n_1 + c_1 \times n_2,$$

где c_1 – константа нормирования, равная отношению числа вентилях схемы к числу триггеров; функции $f_1(v, f)$ и $f_2(v, f)$ определяют различие на множестве выходов логических элементов и псевдовыходов соответственно.

Алгоритмы СО оптимизации рассеивания тепла тестовых последовательностей. Алгоритмы СО могут быть применены не только к задачам построения идентифицирующих последовательностей, но и к их оптимизации. В [15] авторами разработана стратегия генерации энергоэффективных тестов. Она основана на понятии избыточного тестирования и состоит из трёх этапов.

На первом этапе производится генерация множества избыточных тестов $S = \{s_1, s_2, \dots, s_l\}$. Тест называется избыточным с заданным фактором избыточности r , если каждая неисправность в целевом множестве проверяется не мене, чем r подпоследовательностями. Данный этап реализован с помощью двухуровневого алгоритма СО построения тестов путём добавления для каждой неисправности соответствующего счётчика.

На втором этапе для каждой из построенных последовательностей оценивается рассеивание тепловой энергии:

$$E(s_{\text{вх}}) = 0,5 \times V^2 \times C \times f \times A(s_{\text{вх}}),$$

где: V – напряжение работы схемы; C – физическая ёмкость на выходе вентиля; f – частота работы схемы; $A(s_{\text{вх}})$ – число событий при моделировании на заданной входной последовательности.

Третий этап заключается в выборе такого подмножества последовательностей $S' \subset S$, чтобы полнота теста была не хуже начальной $P(S') = P(S)$, а рассеивание тепла минимальным

$E(S') \rightarrow \min$. Очевидно, что выбор различных подмножеств последовательностей будет влиять и на полноту теста и на параметр рассеивания тепла. Задача выбора такого подмножества является NP-полной и её решение строится на основании алгоритма СО. В отличие от предыдущих задач, в данной в качестве конфигурации K_i выступает множество номеров подпоследовательностей из S , которые фактически задают входную последовательность. Список номеров в разных конфигурациях изменяется в ходе эволюции, также как и число номеров в списке. Используются три классические операции возмущения для множеств: обмен элементами, удаление элемента и добавление случайного элемента.

Алгоритм делится на две фазы, в каждой из которых применяется свой вид функции оценки. В фазе 1 из множества S выбираются такие подмножества, для каждого из которых выполнено первое условие. Функция оценки при этом имеет вид

$$C(K_i) = P(S) - P(S').$$

При обнаружении набора подпоследовательностей, для которого условие выполнено, алгоритм СО переходит к фазе 2 оптимизации данного набора. Здесь развитие конфигурации идёт таким образом, чтобы уменьшить рассеивание тепла для подмножества последовательностей, и оценка имеет вид $C(K_i) = E_{\max} - E_i$, где $E_{\max} = E(S)$.

Эффективность алгоритма СО измеряется уменьшением рассеивания тепла для конфигураций на входе и выходе фазы 2.

Сравнение алгоритмов СО и ГА. Разработка алгоритмов, рассмотренных выше, даёт разработчикам возможность выбора между алгоритмами СО и ГА. Близость данных стратегий позволяет ставить вопрос о сравнительной эффективности их поисковых свойств. В работе [16] проводился практический сравнительный анализ ГА и СО, из которого следует, что метод СО на большинстве задач не проигрывает ГА, а на многих – выигрывает. Для сравнения практической эффективности алгоритмов СО и ГА изучались результаты численных экспериментов на схемах из контрольного каталога ISCAS-89 [17]. Аккуратное сравнение различных алгоритмов является сложной проблемой [18] и определить абсолютно лучший алгоритм не представляется возможным, поскольку сравнение производится сразу по нескольким критериям.

Результаты сравнительного анализа алгоритмов СО и ГА [7] построения ИП приведены в табл. 1. Видно, что оба подхода показывают

приблизительно равные результаты.

Таблица 1. Сравнение алгоритмов СО и ГА построения инициализирующих последовательностей.

Параметр сравнения	Результаты сравнения		
	лучше ГА	лучше СО	ничья
Число инициализированных триггеров	4	0	34
Длина инициализирующей последовательности	1	3	34
Рассмотрено число точек в пространстве поиска	14	18	6

При проведении экспериментов верификации эквивалентности применялась следующая стратегия. В исходную схему вносились миноритарные изменения (заменялся тип одного случайно выбранного логического вентиля) и для двух схем строилась различающая последовательность. Эксперимент с каждой схемой проводился 25 раз. В среднем алгоритм СО построил различающую последовательность в 96.71% экспериментов, алгоритм ГА [8] – в 94.7%, алгоритм VEGA – 88,35% [19], алгоритм AQUILA – 65,00% [20]. Последние два алгоритма являются структурными и для них не приводятся результаты работы с большими схемами.

При сравнении эффективности алгоритмов построения энергоэффективных тестов СО и ГА [9] для третьего этапа использовались одинаковые результаты, полученные на этапах 1 и 2. Результаты сравнения приведены в табл. 2.

Таблица 2. Сравнение алгоритмов СО и ГА построения энергоэффективных тестов

	Результат ГА, в среднем	Результат СО, в среднем	Лучше ГА, раз	Лучше СО, раз	Ничья, раз
Уменьшение рассеивания тепла	85.59%	86.34%	0	22	2
Уменьшение длины результирующей последовательности	12.04 раз	12.89 раз	3	15	6
Время работы фазы 3	-	-	4	5	15

Видно, алгоритм СО в подавляющем числе экспериментов показал лучшие результаты как в терминах уменьшения рассеивания тепла, так и уменьшения длинны тестовой последовательности.

Выводы. В статье рассмотрены основные подходы применения стратегии СО к решению задач построения идентифицирующих последовательностей и их оптимизации. Из анализа результатов машинных экспериментов видно, что алгоритмы, построенные на основе данной стратегии, показывают очень хорошие эксплуатационные характеристики. Более того, несмотря на упрощение эволюции – одно решение в алгоритмах СО против популяции решений в ГА – данные алгоритмы часто находят лучшие решения. На основании этого можно сделать вывод о целесообразности применения данной стратегии к решению других задач в цикле проектирования цифровых схем.

Список литературы: 1. *Goldberg D.E.* Genetic Algorithm in Search, Optimization, and Machine Learning / *D.E. Goldberg*. – Addison-Wesley, 1989. – 432 p. 2. *Скобцов Ю.А.* Основы эволюционных вычислений / *Ю.А. Скобцов*. – Донецк: ДонНТУ, 2008. – 326 с. 3. Неітеративні, еволюційні та мультиагентні методи синтезу нечіткої логічних і нейромережних моделей: Монографія / Під заг. ред. *С.О. Суботіна*. – Запоріжжя: ЗНТУ, 2009. – 375 с. 4. *Corno F.* GATTO: a Genetic Algorithm for Automatic Test Pattern Generation for Large Synchronous Sequential Circuits / *F. Corno, P. Prinetto, M. Rebaudengo, M. Sonza Reorda* // IEEE Transactions on Computer-Aided Design, August 1996. – Vol. 15. – № 8. – P. 943-951. 5. *Skobtsov A.* Genetic algorithms in test generation for digital circuits / *Y.A. Skobtsov, D.E. Ivanov, V.Y. Skobtsov* // Proceedings of the 8th Biennial Baltic Electronics Conference. – Tallinn Technical University, 2002. – P. 291-294. 6. *Corno F.* A Genetic Algorithm for the Computation of Initialization Sequences for Synchronous Sequential Circuits / *F. Corno, P. Prinetto, M. Rebaudengo, M. Sonza Reorda, G. Squillero* // ATS97: The Sixth IEEE Asian Test Symposium, Akita (JP). – 1997. – P. 56–61. 7. *Иванов Д.Е.* Построение инициализирующих последовательностей синхронных цифровых схем с помощью генетических алгоритмов / *Д.Е. Иванов, Ю.А. Скобцов, А.И. Эль-Хатиб* // Проблеми інформаційних технологій. – Херсон: ХНТУ. – 2007. – № 1. – С. 158–164. 8. *Иванов Д.Е.* Генетический подход проверки эквивалентности последовательностных схем / *Д.Е. Иванов* // "Радіоелектроніка. Інформатика. Управління". – Запоріжжя: ЗНТУ. – 2009. – № 1 (20). – С. 118–123. 9. *Иванов Д.Е.* Генетический алгоритм оптимизации рассеивания тепловой энергии входных тестовых последовательностей / *Д.Е. Иванов* // Наукові праці Донецького національного технічного університету. Серія: "Обчислювальна техніка та автоматизація". – Випуск 18 (169). – Донецьк: ДонНТУ, 2010. – С. 206–215. 10. *Krishnamurthy B.* On the Complexity of Estimating the Size of a Test Set / *B. Krishnamurthy, S. B. Akers* // IEEE Trans. on Computers. – August 1984. – P. 750–753. 11. *Metropolis N.* Equation of State Calculation by Fast Computing Mashines / *N. Metropolis, A.W. Rosenbluth, M.N. Rosenbluth, A.H. Teller, E. Teller* // Journal of Chem.Phys. – 1953. – Vol. 21. – № 6. – P. 1087–1092. 12. *Kirkpatrick S.* Optimization by simulating annealing / *S. Kirkpatrick, C.D. Gelatt, M.P. Vecchi* // Science. – 1983. – Vol. 220. – P. 671–680. 13. *Иванов Д.Е.* Алгоритм построения инициализирующих последовательностей цифровых схем, основанный на стратегии симуляции отжига / *Д.Е. Иванов, Р. Зуауи* // Искусственный интеллект, 2009. – № 4. – С. 415–424. 14. *Иванов Д.Е.* Верификация эквивалентности цифровых схем с использованием стратегии симуляции отжига / *Д.Е. Иванов, Р. Зуауи* // "Науковий вісник Чернівецького університету". – Вип. 479. – "Комп'ютерні системи та компоненти", 2009. – С. 33–41. 15. *Иванов Д.Е.* Алгоритм симуляции отжига оптимизации рассеивания тепла диагностических тестов / *Д.Е. Иванов, Р. Зуауи* // "Радіоелектронні і комп'ютерні системи", 2010. – № 7 (48). – С. 170–175. 16. *Inberg L.* Simulated Annealing: Practice versus Theory / *L. Inberg* // Journal of Mathematical Computer Modelling. – 1993. – Vol. 18. – № 1. – P. 29–57. 17. *Brgles F.* Combinational profiles of sequential benchmark circuits / *F. Brgles, D. Bryan, K. Kozminski*

// International symposium of circuits and systems, ISCAS-89. – 1989. – P. 1929–1934.
18. Corno F. Comparing topological, symbolic and GA-based ATPGs: an experimental approach / F. Corno, P. Prinetto, M. Rebaudengo, M. Sonza Reorda // Proceedings of the IEEE International Test Conference on Test and Design Validity, Washington (USA). – 1996. – P. 39–47.
19. Corno F. VEGA: A Verification Tool Based on Genetic Algorithms / F. Corno, M. Sonza Reorda, G. Squillero // ICCD98, International Conference on Circuit Design, Austin, Texas (USA). – 1998. – P. 321–326. **20.** Cheng K.-T. AQUILA: An Equivalence Checking System for Large Sequential Designs / K.-T. Cheng, S.-Yu. Huang, K.-Ch. Chen, F. Brewer, C.-Y. Huang // IEEE Transactions on Computers. – 2000. – Vol. 49. – № 5. – P. 443–464.

*Статья представлена д.т.н., проф., зав. каф. АСУ ДонНТУ
 Скобцовым Ю.А.*

УДК 681.518:681.326.7

Застосування алгоритмів симуляції відпалювання в задачах ідентифікації цифрових схем / Іванов Д.Є. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 60 – 69.

У статті розглянуто досвід застосування еволюційного алгоритму симуляції відпалювання до вирішення задач ідентифікації, що виникають при проектуванні цифрових схем. Описано загальну структуру алгоритму симуляції відпалювання, його компоненти рішення задач побудови ідентифікуючих послідовностей та їх оптимізації. Проведено порівняльний аналіз ефективності з генетичними алгоритмами. Іл.: 2. Табл.: 2. Бібліогр.: 20 назв.

Ключові слова: еволюційні алгоритми, ідентифікація, цифрова схема, симуляція відпалювання, генетичний алгоритм.

UDC 681.518:681.326.7

The use of simulation annealing algorithms in problems of identification of digital circuits / Ivanov D.E. / Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 60 – 69.

The article describes the experience of use of evolutionary algorithm of the simulated annealing to solve the identification problems in the design of digital circuits. It is described the overall structure of the simulated annealing algorithm and its components for algorithm of constructing of the identifying sequences and to optimize them. A comparative analysis of the effectiveness with genetic algorithms is performed. Figs.: 2. Tabl.: 2. Refs.: 20 titles.

Key words: evolutionary algorithm, identification, digital circuit, simulated annealing, genetic algorithm.

Поступила в редакцію 11.02.2011

А.С. КУИМОВА, аспирантка Волжской государственной академии водного транспорта, Нижний Новгород,

Ю.С. ФЕДОСЕНКО, д.т.н., проф., зав. каф. ИСУиТ Волжской государственной академии водного транспорта, Нижний Новгород

СИНТЕЗ ОГРАНИЧЕННЫХ ПО СТРУКТУРЕ ОПТИМАЛЬНО-КОМПРОМИССНЫХ СТРАТЕГИЙ ОБСЛУЖИВАНИЯ ПОТОКА ОБЪЕКТОВ

Вводится модель однофазного обслуживания детерминированного потока объектов. Формулируется бикритериальная задача синтеза стратегий обслуживания. Показано, что учет ограничений на структуру стратегий обслуживания выводит задачу из класса *NP*-трудных и позволяет построить полиномиальный алгоритм синтеза стратегий, оптимальных по Парето. Приводится пример. Библиогр.: 8 назв.

Ключевые слова: стратегии обслуживания, *NP*-трудность, оптимальность по Парето.

Постановка проблемы и анализ литературы. В работе [1] рассмотрена каноническая модель обслуживания потока объектов стационарным процессором, описывающая практически важную проблему оптимизации использования дискретных ресурсов. Вместе с тем, естественное требование адекватности математического описания реальным ситуациям, возникающим в технических, экономических и организационных системах, нередко предопределяет необходимость различных усложнений канонической модели. Так, в данной работе формулируется значимая для приложений (в частности, в автоматизированных производственных системах и на внутреннем водном транспорте [2]) модель, в которой стратегии обслуживания оцениваются по значениям не одного, а двух независимых критериев. При этом в качестве парадигмы решения сформулированной экстремальной задачи принята концепция Парето [3, 4], предполагающая выявление и предъявление лицу, принимающему решение, всего множества оптимально-компромиссных стратегий обслуживания. Данная оптимизационная задача относится к числу *NP*-трудных [5]. Поэтому актуальной является проблема построения таких модификаций бикритериальной модели обслуживания, которые при сохранении адекватности описания порождают полиномиально разрешимые подклассы задачи синтеза стратегий обслуживания. Одна из таких модификаций формализует ограничение на допустимую величину опережений в обслуживании.

Цель статьи – разработка бикритериальной модели однопроцессорного обслуживания потока объектов, порождающей

полиномиально разрешимую задачу синтеза оптимально-компромиссных стратегий, и построение решающего алгоритма.

Математическая модель и постановка задачи. Задан конечный детерминированный поток объектов $O_n = \{o_1, o_2, \dots, o_n\}$, которые подлежат однофазному обслуживанию на стационарном процессоре P . Для каждого объекта o_i ($i = \overline{1, n}$) определены целочисленные характеристики: t_i – момент поступления в очередь на обслуживание ($0 \leq t_1 \leq t_2 \leq \dots \leq t_n$); τ_i – норма длительности обслуживания; a_i – штраф за единицу времени простоя в ожидании обслуживания; d_i – мягкий директивный срок завершения обслуживания ($d_i \geq t_i + \tau_i$). Если обслуживание объекта o_i начинается в момент времени t_i^* ($t_i^* \geq t_i$), то величина индивидуального штрафа по этому объекту определяется значением линейной функции вида $a_i(t_i^* - t_i)$. Обслуживание объекта o_i ($i = \overline{1, n}$) может быть начато свободным процессором в любой момент времени t ($t \geq t_i$) и осуществляется без прерываний; одновременное обслуживание двух и более объектов и неоправданные простои процессора запрещены; процессор готов к обслуживанию объектов потока O_n в момент времени $t = 0$.

Очевидно, любая стратегия S обслуживания потока O_n представляет собой перестановку $\{i_1, i_2, \dots, i_n\}$ совокупности индексов $N = \{1, 2, \dots, n\}$; при этом объект с индексом i_k в стратегии S обслуживается k -м по очереди ($k = \overline{1, n}$). Введем следующее определение: стратегию S_q назовем q -стратегией, если не существует объекта, который при реализации стратегии S_q допускает пропуск вперед на обслуживание более q других объектов. Множество таких стратегий будем обозначать как $Z(q)$. Пусть $t^*(i_k, S_q)$ и $\bar{t}(i_k, S_q)$ соответственно моменты начала и завершения обслуживания объекта с индексом i_k при реализации стратегии S_q . Считаем, что реализация компактна, т.е. выполняются соотношения

$$t^*(i_k, S) = \max[\bar{t}(i_{k-1}, S_q), t_{i_k}], \quad k = \overline{2, n}; \quad t^*(i_1, S) = t_{i_1};$$

$$\bar{t}(i_k, S_q) = t^*(i_k, S_q) + \tau_{i_k}, \quad k = \overline{1, n}.$$

Качество каждой стратегии будем оценивать по значениям пары критериев $K_1(S_q)$ и $K_2(S_q)$, первый из которых представляет собой суммарный штраф за простои объектов в ожидании обслуживания, а второй – оценивает максимальное по продолжительности нарушение

директивного срока завершения обслуживания среди всех объектов потока O_n . С учетом введенных выше обозначений имеем

$$K_1(S_q) = \sum_{k=1}^n a_{i_k} (t^*(i_k, S_q) - t_{i_k}); \quad K_2(S_q) = \max_{1 \leq k \leq n} (\bar{t}(i_k, S_q) - d_{i_k}, 0).$$

Формализация парадигмы Парето на множестве $Z(q)$ приводит к следующей задаче выделения в плоскости $(K_1(S_q), K_2(S_q))$ полной совокупности эффективных точек и формирования им соответствующих оптимально-компромиссные стратегии обслуживания

$$\{ \min_{S_q \in Z(q)} (K_1(S_q)), \min_{S_q \in Z(q)} (K_2(S_q)) \}. \quad (1)$$

Синтез стратегий обслуживания. Алгоритм решения задачи (1) построим на основе идеологии динамического программирования [6, 7]. С этой целью введём операцию $\otimes$ между произвольным вектором $x = (x_1, x_2)$ и множеством Y векторов $y = (y_1, y_2)$ той же размерности: через $Y \otimes x$ обозначим совокупность всех векторов $v = (v_1, v_2)$, где первая компонента представима в виде $v_1 = y_1 + x_1$, а вторая определяется по правилу $v_2 = \max(y_2, x_2)$. Будем обозначать: для произвольного множества двумерных векторов-оценок M через $eff[M]$ максимальное по включению подмножество недоминируемых в M векторов; через $F(t)$ подмножество индексов объектов, которые поступают в систему обслуживания в момент времени t . При обслуживании объектов потока O_n , решения принимаются в те моменты времени, когда процессор свободен и необходимо выбрать следующий объект. Очевидно, что набор (t, j, M_j^t, Q) определяет текущее состояние системы в момент t принятия решения, где j – наименьший из индексов в множестве индексов необслуженных объектов, а M_j^t – множество индексов обслуженных объектов, обошедших на момент времени t объект с индексом j ($|M_j^t| \leq q$), Q – множество индексов объектов, ожидающих обслуживания в этот момент времени. Введем обозначения: $\{M\}_r$ – совокупность индексов объектов из множества M с индексами большими r ; $D(t, \Delta)$ – совокупность индексов объектов, прибывающих в систему на интервале $[t + 1, t + \Delta]$, $\Delta \geq 1$, т.е.

$$D(t, \Delta) = \bigcup_{i=1}^{\Delta} F(t + i).$$

В множество Q считается включенным индекс 0 фиктивного объекта с параметрами $t_0 = t$, $\tau_0 = \min \{ \tau \mid D(t, \tau) \neq \emptyset \}$, $a_0 = 0$, $d_0 = \infty$. Выбор такого объекта на обслуживание означает простой процессора с момента времени t до момента поступления в систему очередного объекта. Если в

текущем состоянии фиктивный объект не был выбран на обслуживание, то он исключается из множества Q и в дальнейшем не рассматривается.

Введем оператор R раскрытия состояния системы обслуживания: если (t, j, M_j^t, Q) является некоторым произвольным состоянием, то $R(t, j, M_j^t, Q)$ – множество порождаемых им состояний, т.е. непосредственно следующих за состоянием (t, j, M_j^t, Q) . Последовательно раскрывая все состояния, начиная от начального, получим в итоге совокупность финальных состояний. Заметим, что одно и то же состояние может сформироваться при раскрытии различных состояний. Присвоим начальному состоянию номер 0 и последовательно пронумеруем все состояния в порядке их порождения оператором R . Для одного и того же состояния, если его можно получить несколькими способами, используется только один номер, присвоенный первоначально; каждое рассматриваемое состояние может характеризоваться несколькими оценками.

Пусть $W_q(t, j, M_j^t, Q)$ – задача, получаемая из исходной задачи (1) при условии, что обслуживание начинается в момент времени t ; при этом надо обслужить объекты с индексами из множества Q и все объекты, поступающие в систему позднее момента времени t . Через $E_q(t, j, M_j^t, Q)$ обозначим совокупность эффективных по Парето двумерных оценок, получаемую при решении задачи $W_q(t, j, M_j^t, Q)$. Тогда полная совокупность эффективных оценок задачи (1) будет представлять множество $E_q(0, 1, M_1^0, F(0))$.

Ведем обозначения: Q_t – совокупность индексов всех необслуженных к моменту t объектов, т.е. $Q_t = Q \cup \{k : o_k \in O_n \text{ и } t_k > t\}$; $\xi[W]$ – минимальный из положительных индексов объектов множества W . Очевидно, что для любого $\theta \geq 0$, $j = \alpha$ и $Q = \{\alpha\}$, где α – произвольный индекс объекта из потока O_n , имеет место соотношение

$$E_\delta(t_n + \theta, \alpha, M_\alpha^{t_n + \theta}, \{\alpha\}) = (\alpha_\alpha(t_n + \theta - t_\alpha), \max(\bar{t}(\alpha, S_q) - d_\alpha, 0)). \quad (2)$$

Если в состоянии (t, j, M_j^t, Q) в качестве следующего обслуживаемого выбран объект с индексом $\alpha \in Q$ ($\alpha \neq j$) и при этом $|M_j^t| < q$, то очередным моментом принятия решения будет $t + \tau_\alpha$ и выбор индекса следующего объекта для обслуживания будет осуществляться из множества $(Q \setminus \{\alpha\}) \cup D(t, \tau_\alpha)$; при этом $M_j^t \cup \{\alpha\}$ – множество объектов, опередивших объект с индексом j . Если на обслуживание выбран объект с индексом j , то в следующий момент принятия решения $t + \tau_j$, наименьшим индексом объекта будет являться $j^* = \xi[Q \setminus \{j\}]$, множеством его опередивших объектов будет $\{M_j^t\}_{j^*}$, а множество ожидающих

обслуживания объектов определяется как $(Q \setminus \{j\}) \cup D(t, \tau_j)$. Следовательно, при $|M_j^t| < q$ имеет место соотношение

$$\begin{aligned} E_q(t, j, M_j^t, Q) = & \text{eff}[\text{eff}[\bigcup_{\alpha \in Q \setminus \{j\}} ((a_\alpha(t - t_\alpha), \max(\bar{t}(\alpha, S_q) - d_\alpha, 0)) \otimes \\ & \otimes E_q(t + \tau_\alpha, j, M_j^t \cup \{\alpha\}, Q \setminus \{\alpha\} \cup D(t, \tau_\alpha))] \cup ((a_j(t - t_j), \max(\bar{t}(j, S_q) - d_j, 0)) \otimes E_q(t + \tau_j, \xi[Q_t \setminus \{j\}], \{M_j^t\}_{\xi[Q_t \setminus \{j\}]}, Q \setminus \{j\} \cup D(t, \tau_j))]. \end{aligned} \quad (3)$$

Если в состоянии (t, j, M_j^t, Q) имеет место равенство $|M_j^t| = q$, то следующим можно обслуживать лишь объект с индексом j . Поэтому

$$\begin{aligned} E_q(t, j, M_j^t, Q) = & \text{eff}[(a_j(t - t_j), \max(\bar{t}(j, S) - d_j, 0)) \otimes \\ & \otimes E_q(t + \tau_j, \xi[Q_t \setminus \{j\}], \{M_j^t\}_{\xi[Q_t \setminus \{j\}]}, Q \setminus \{j\} \cup D(t, \tau_j))]. \end{aligned} \quad (4)$$

Выражения (2) – (4) суть решающие рекуррентные соотношения. Непосредственным анализом [8] устанавливается, что вычислительная сложность описанного алгоритма оценивается величиной $O(n^4)$.

Результаты вычислительных исследований и выводы. Целью вычислительных экспериментов являлось определение скоростных и объёмных характеристик разработанного алгоритма. Для значений параметров потока O_n были установлены следующие диапазоны изменения: $t_{i-1} \leq t_i \leq t_{i-1} + 5$, $1 \leq a_i \leq 11$, $1 \leq \tau_i \leq 15$, $t_i + \tau_i \leq d_i \leq t_i + \tau_i + 4$, $i = \overline{1, n}$. Таким образом была обеспечена достаточно высокая плотность потока и гарантировалось получение заведомо пессимистических оценок характеристик алгоритма. Размерность потока изменялась от 6 до 13 с единичным шагом; аналогично параметр q принимал значения от 1 до $\lfloor n/2 \rfloor$. В соответствующих узлах решетки (n, q) генерировались по равномерному распределению данные для серии из 100 частных задач типа (1).

В результате вычислительных исследований на компьютере с процессором Intel Core 2 Duo, 3,16 ГГц (ОЗУ 2 Гб) установлено, что продолжительность синтеза полной совокупности эффективных оценок и соответствующих им оптимально-компромиссных стратегий обслуживания не превышает 68 с, объем требуемой оперативной памяти не превышает 1,5 Гб, что позволяет рекомендовать разработанный алгоритм для включения в автоматизированные системы диспетчеризации процессов рассмотренного типа.

Список литературы: 1. Коган Д.И. Задача диспетчеризации: анализ вычислительной сложности и полиномиально разрешимые подклассы / Д.И. Коган, Ю.С. Федосенко / Дискретная математика. – 1996. – Т. 8. – Вып. 3. – С. 135-147. 2. Коган Д.И. Проблема синтеза оптимального расписания обслуживания бинарного потока объектов mobile-процессором / Д.И. Коган, Ю.С. Федосенко, А.В. Шеянов / Труды III Международной конференции "Дискретные модели в теории управляющих систем", Москва, 1998. – М.: Изд-во МГУ им. М.В. Ломоносова, 1998. – С. 43-46. 3. Подиновский В.В. Парето-оптимальные решения многокритериальных задач / В.В. Подиновский, В.Д. Ногин. – М.: Физматлит, 2007. – 256 с. 4. Лотов В.А. Многокритериальные задачи принятия решений / В.А. Лотов, И.И. Поспелова. – М.: МАКС Пресс, 2008. – 197 с. 5. Гэри М. Вычислительные машины и труднорешаемые задачи / М. Гэри, Д. Джонсон. – М.: Мир, 1982. – 416 с. 6. Беллман Р. Прикладные задачи динамического программирования / Р. Беллман, С. Дрейфус. – М.: Наука, 1965. – 457 с. 7. Коган Д.И. Динамическое программирование и дискретная многокритериальная оптимизация / Д.И. Коган. – Н. Новгород: Изд-во ННГУ, 2005. – 260 с. 8. Кормен Т. Алгоритмы: построение и анализ / Т. Кормен, Ч. Лейзерсон, Р. Ривест, К. Штайн. – М.: Вильямс, 2007. – 1296 с.

УДК 681.31+519.8

Синтез обмежених за структурою оптимально-компрісних стратегій обслуговування потоку об'єктів / Федосенко Ю.С., Куимова А.С. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 70 – 75.

Вводиться модель однофазного обслуговування детермінованого потоку об'єктів. Формулюється бікритеріальна завдання синтезу стратегій обслуговування. Показано, що врахування обмежень на структуру стратегій обслуговування виводить завдання з класу NP-важких та дозволяє побудувати поліноміальний алгоритм синтезу стратегій обслуговування, оптимальних за Парето. Наводиться приклад. Бібліогр.: 8 назв.

Ключові слова: стратегії обслуговування, NP-важкий, оптимальність за Парето.

UDC 681.31+519.8

The synthesis of structure-bounded optimal compromise service policies of object flow / Fedosenko Yu.S., Kuimova A.S. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 70 – 75.

In this paper a model of single-phase service of deterministic objects flow is discussed. A bicriteria problem of service policies synthesis is formulated. It is shown that if restrictions on service policies structure are imposed the problem will become non NP-hard and it will be possible to design a polynomial algorithm for Pareto-optimal service policies synthesis. Example is introduced. Refs.: 8 titles.

Key words: service policies, NP-hardness, Pareto optimality.

Поступила в редакцію 15.02.2011

А.І. КУРСІН, зав. лаб. НТУ "ХПІ", Харків

МЕТОД НОРМАЛІЗАЦІЇ СЛОВОФОРМ ЗІ СЛОВНИКОМ ПОЧАТКОВИХ ФОРМ СЛІВ БЕЗ ГРАМАТИЧНОЇ ІНФОРМАЦІЇ

Запропоновано оригінальний метод нормалізації словоформ флексивних мов з використанням словника початкових форм слів, в якому слова не мають кодів типу словозміни. Такий алгоритм хоча і використовує словник, але не потребує граматичного кодування слів у ньому, хоча і вираховує можливі початкові форми за таблицею псевдо-закінчень, але досягає більшої точності, перевіряючи результати за словником. Табл.: 2. Бібліогр.: 8 назв.

Ключові слова: нормалізація словоформ, словник початкових форм.

Постановка проблеми і аналіз літератури. Неодмінною складовою більшості систем автоматичної обробки природномовної текстової інформації є етап морфологічного аналізу словоформ. На цьому етапі встановлюється початкова форма слова, його граматичний клас та граматична інформація щодо форми слова. Для великого класу задач, наприклад, інформаційного пошуку, буває достатньо лише встановлення початкової форми слова (ця операція ще називається нормалізацією словоформ або лемматизацією).

Традиційно алгоритми морфологічного аналізу розділяються на словникові і безсловникові. Словниковий алгоритм [1 – 3] оперує словником слів або основ з приписаною їм інформацією щодо граматичних категорій слова та типу словозміни [4]. Одним з суттєвих недоліків такого підходу є великі трудозатрати на створення такого словника, бо граматичне кодування слів попри всі спроби автоматизації виконується в значній мірі вручну. Безсловниковий алгоритм [5, 6] замість словника використовує списки флексій, квазі-флексій, квазі-суфіксів, притаманних словам даної мови, і робить аналіз, перевіряючи кінець словоформи за цими списками. Безумовною перевагою таких алгоритмів є те, що множина слів, які можуть бути ними оброблені, не обмежується рамками словників, що використовуються. Але безсловникові алгоритми зазвичай відрізняє менша точність аналізу, більша кількість хибних варіантів, хоча й дещо більша швидкість роботи.

Мета статті – описати оригінальний алгоритм нормалізації словоформ, який поєднує властивості обох типів алгоритмів морфологічного аналізу та частково позбавлений недоліків їх обох. Такий алгоритм хоча і використовує словник, але не потребує граматичного

кодування слів у ньому, хоча і вираховує можливі початкові форми за таблицею псевдо-закінчень, але досягає більшої точності, перевіряючи результати за словником.

Опис задачі. Скажемо декілька слів про задачу, яка дала поштовх до розробки алгоритму. Створювалася програма підтримки електронних словників з функцією пошуку в них довільних форм слів. Виявилось, що словники, якими мала бути наповнена система, містять багато слів відсутніх у словникові наявної процедури морфологічного аналізу, а введення їх до цього словника разом з граматичним кодуванням було неприйнятним через стислі строки відведені для виконання робіт. Тоді з'явилася ідея, використовуючи морфологічні таблиці словникового алгоритму, генерувати всі можливі початкові форми для даної довільної форми слова, а потім "просіювати" їх по списку початкових форм слів словника.

Опис алгоритму. Класичний словниковий алгоритм морфологічного аналізу [1, 2] разом зі словником використовує морфологічні таблиці, зокрема таблицю парадигм. Кожний з рядків цієї таблиці відповідає певній парадигмі або типу словозміни, а кожна позиція в рядку – певній формі слова і містить відповідну флексію. На базі цієї таблиці ми створили власну, кожний рядок якої містить одну з флексій непрямих форм слів та всі флексії початкових форм, що зустрічаються спільно з нею в парадигмах. Фрагмент такої таблиці показаний в табл. 1.

Таблиця 1. Флексії непрямой форми та початкових форм

Флексія непрямой форми	Флексії початкових форм
-ет	-ть, -ять, -ти, -уть, -вать, -ать, -ить, -отъ, -ь
-ете	-ть, -ять, -ти, -уть, -вать, -ать, -ить, -отъ, -ь
-етесь	-ться, -яться, -утся, -атся, -отся, -тись, -ваются, -иься, -ься
-ется	-ться, -яться, -утся, -атся, -отся, -тись, -ваются, -иься, -ься
-ех	-и
-ехот	-иста
-ехстах	-иста
-ец	-цы, -ьца, -ца, -ьцо
-ешь	-ть, -ять, -ти, -уть, -вать, -ать, -ить, -отъ, -ь
-ешься	-ться, -яться, -утся, -атся, -отся, -тись, -ваются, -иься, -ься

Для певної словоформи алгоритм знаходить в таблиці всі можливі для неї флексії непрямих форм. Та будує потенційні варіанти початкових форм слова, використовуючи флексії з відповідних рядків таблиці.

Приклад такого списку для російської словоформи "девочкой" наведений в табл. 2.

Таблиця 2. Список для словоформи "девочкой"

девочкий	девочкоить	девочкойя
девочкий	девочкыть	девочка
девочкая	девочкетъ	девочко
девочкое	девочкоять	

Потім проводиться пошук початкових форм зі списку в індексі початкових форм слів словника. Звісно, згенерований список початкових форм містить багато хибних варіантів. Але зазвичай лише один з них збігається з якою-небудь початковою формою зі словника. Ми були самі здивовані результатами попередніх тестів, які показали, що алгоритм "не сходиться", тобто дає більше одного можливого результату, або не дає жодного збігу, лише для 6% російських словоформ і 4% українських. До цих відсотків входили також і дійсні омоформи, тобто такі випадки, де і людина не може зробити однозначну нормалізацію, не вдаючись до контексту.

Згодом морфологічна таблиця даного алгоритму була вдосконалена за матеріалами граматичного словника Залізняка [4]. Всі непрямі закінчення довжини менш ніж 4 були розширенні до цієї довжини додаванням останніх літер основ, з якими вони вживаються. Таке ж розширення було проведене і з деякими окремими закінченнями, які давали суттєву кількість хибних результатів. Короткі слова (переважно службові слова, займенники, тощо) були винесені в окремий список стоп-слів (300 – 400 одиниць). Також, зважаючи на задачу, в якій застосовувався даний алгоритм, виявилось доцільним ввести певний пріоритет закінчень у таблиці. Наприклад, непрямі форми дієприкметників (напр., "пораненими") спершу приводяться до чоловічого роду однини ("поранений"), і лише за відсутності такої форми в словнику, проводиться пошук неозначеної форми дієслова ("поранити").

В результаті зазначених вдосконалень кількість випадків "несходимості" алгоритму зменшилася для російської мови – до 3,3%, з яких 2,7% складають дійсні омоформи.

Зауваження щодо реалізації. Будь-який алгоритм, що стосується пошуку інформації у словнику і претендує на практичну реалізацію, мусить мати достатньо швидку пошукову частину. В нашому випадку кількість пошукових операцій у словнику зростає завдяки великій кількості хибних початкових форм, які приймають участь у пошуку (в середньому в 15,1 рази для російської мови). Тому треба було вжити

заходи для забезпечення достатньої швидкості пошуку. Ми використовували індексацію слів словника за допомогою В-дерева [7] з ключами мінливої довжини зі стисненням методом Купера [8]. Така структура поєднує можливість швидкого пошуку строкових ключів з компактним об'ємом індексу. Збільшення кількості пошукових ключів для однієї словоформи нівелюється пакетним пошуком в індексі відразу всього списку початкових форм. Швидкість нормалізації в тестах складала приблизно 25 мс на тисячу словоформ (AMD Athlon 2.2 ГГц).

На базі описаного алгоритму ми розробили процедури нормалізації для російської та української мов, а також для англійської мови. В останньому випадку замість використання таблиць флексій працює підпрограма, яка генерує можливі початкові форми слів за правилами англійської мови. Ці процедури нормалізації використовуються в ряді інформаційно-пошукових та лінгвістичних систем, як в своєму первинному варіанті – для пошуку довільної форми слова у списку початкових форм слів, так і у вигляді окремого модуля нормалізації, який комплектується достатньо великим словником слів відповідної мови. Перевага практичного використання даного алгоритму в тому, що укладачі словника позбавляються кодування ключових слів щодо типу їх словозміни. Проблема обробки незнайомих слів ми не вирішуємо кардинально, але процедура поповнення словника новими словами також значно спрощується. Завдяки використанню достатньо швидкого пошукового індексу, наша процедура нормалізації може замінити собою пошук ключових слів в індексі інформаційно-пошукової системи.

Висновки. Описаний нами алгоритм доводить, що достатньо точну нормалізацію словоформ флексивних мов можна проводити і без кодування слів у словнику щодо типу словозміни. Це робить перспективним застосування алгоритму в інформаційно-пошукових системах та системах обробки текстової інформації.

Список літератури: 1. Белоногов Г.Г. Алгоритм морфологического анализа русских слов / Г.Г. Белоногов, Ю.Г. Зеленков // Вопр. инф. теории и практики. – М.: 1985. – № 53. – С. 62–93. 2. Антошкина Ж.Г. Морфологический процессор русского языка / Ж.Г. Антошкина // Бюллетень машинного фонда русского языка. – М.: 1996. – Вып. 3. – С. 53–57. 3. Сегалович И. Русский морфологический анализ и синтез с генерацией моделей словоизменения для не описанных в словаре слов / И. Сегалович, М. Маслов // Труды международного семинара Диалог'98 по компьютерной лингвистике и её приложениям. – Том 2. – 1998. – С. 547–552. 4. Зализняк А.А. Грамматический словарь русского языка / А.А. Зализняк. – М.: Русские словари, 2003. – 795 с. 5. Sheremetyeva S. Rapid Deployment Morphology / S. Sheremetyeva, W. Jin, S. Nirenburg // Machine Translation. – 1998. – Vol. 13. – № 4. – С. 239–268. 6. Ножов И.М. Прикладной морфологический анализ без словаря / И.М. Ножов // КИИ-2000. Труды конференции. – М.: 2000. – Т.1. – С. 424–429. 7. Bayer R. Organization and Maintenance of Large Ordered Indexes / R. Bayer, E. McCreight // Acta

Informatica. – 1972. – № 1 (3). – P. 173–189. **8. Cooper W.S.** The storage problem / *W.S. Cooper* // Mechanical Translation. – 1958. – Vol. 5. – № 2. – P. 74–83.

Стаття представлена д.т.н., проф. НТУ "ХПІ" Серковим О.А.

УДК 81.322.3:81.366

Метод нормализации словоформ со словарём начальных форм слов без грамматической информации / Курсин А.И. // Вестник НТУ "ХПИ". Тематический выпуск: Информатика и моделирование. — Харьков: НТУ "ХПИ". – 2011. – № 17. – С. 76 – 80.

Предложен оригинальный метод нормализации словоформ флективных языков с использованием словаря начальных форм слов, в котором у слов отсутствуют коды типа словоизменения. Такой алгоритм хотя и использует словарь, но не нуждается в грамматической кодировке слов в нем, хотя и высчитывает возможные начальные формы за таблицей псевдо-окончаний, но достигает большей точности, проверяя результаты по словарю. Табл.: 2. Библиогр.: 8 назв.

Ключевые слова: нормализация словоформ, словарь начальных форм.

UDC 81.322.3:81.366

Word form normalization method using a dictionary of citation word forms without grammatical codes / Kursin A.I. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 76 – 80.

An original method of word form normalization if proposed for inflectional languages. The method eliminates the necessity of inflection type codes for words in a dictionary used for normalization. Such algorithm uses a dictionary though, but does not need grammatical code of words in him, though calculates possible initial forms after the table of pseudo-end, but arrives at greater exactness, checking up results after a dictionary. Tabl.: 2. Refs.: 8 titles

Keywords: word form normalization, dictionary of initial forms.

Надійшла до редакції 22.04.2010

В.В. ЛЮБЧЕНКО, к.т.н., доц. ОНПУ, Одеса,
О.С. ШИНКАРЮК, бакалавр ОНПУ, Одеса

МЕТОД БУДУВАННЯ НАВЧАЛЬНОЇ ТРАЄКТОРІЇ В УМОВАХ МОБІЛЬНОГО НАВЧАННЯ

Розглянуто вплив на вимоги до навчального матеріалу обмеженості у часі сеансів навчання в умовах мобільного навчання. Модифіковано алгоритм визначення компонент сильної зв'язності для застосування його до змішаних графів з метою отримання декомпозиції навчального курсу. Запропоновано метод будування навчальної траєкторії на основі множин пов'язаних навчальних концептів. Бібліогр.: 8 назв.

Ключові слова: мобільне навчання, компоненти сильної зв'язності, декомпозиція навчального курсу, навчальна траєкторія.

Постановка проблеми. Стрімкий розвиток та популяризація мобільних пристроїв, таких як смартфони, нетбуки, планшетні комп'ютери, карманні персональні комп'ютери (КПК) і мобільні телефони, призвів до виникнення ідеї про їх використання в процесі навчання для створення освітнього середовища. *Мобільне навчання (m-learning)* – це форма навчання, яка комбінує можливості мобільних обчислювальних пристроїв з можливостями електронного навчання [1].

Мобільне навчання повинно давати можливість об'єктам навчання отримувати необхідну інформацію, що стосується навчального курсу, в будь-який час та з будь-якого місця. Цією інформацією можуть бути адаптовані навчальні матеріали, приклади практичного застосування теоретичного матеріалу, посилання на додаткові ресурси тощо. Процес навчання здійснюється відповідно до *навчальної траєкторії* – логічної послідовності вивчення навчального курсу, яка дозволяє досягти всі навчальні цілі цього курсу. Але при будуванні навчальної траєкторії слід враховувати особливість мобільного навчання – обмеження у часі сеансів навчання, які пов'язані з його технічними (наприклад, обмеження заряду джерела енергії мобільного пристрою) та організаційними (наприклад, мобільність об'єкта навчання) умовами.

Аналіз літератури. Будування навчальної траєкторії є дослідженим питанням в галузі електронного навчання. Знання об'єкта навчання в будь-якій області найчастіше представляються *оверлейною моделлю*, яка заснована на структурній моделі предметної області [2, 3] представлений як мережа концептів. Концепти пов'язані один з одним, утворюючи в такий спосіб семантичну мережу, яка відбиває структуру предметної області навчального курсу. Ідея оверлейної моделі полягає в тому, щоб

представити знання об'єкту навчання з певної теми як перекриття або накладення на модель предметної області. Для кожного концепту моделі предметної області, індивідуальна оверлейна модель зберігає визначене значення, яке є оцінкою рівня знань об'єкту навчання щодо цього концепту.

Слід зазначити, що всі роботи з цієї тематики як одиниці, на яких виконується будування навчальної траєкторії, розглядають окремі концепти навчального курсу або окремі об'єкти навчального матеріалу [4]. Такий підхід є цілком виправданим в умовах, коли об'єкт навчання має постійний доступ до навчального ресурсу. Але в умовах мобільного навчання він призводить до певних утруднень. Для ефективного використання мобільних ресурсів об'єкт навчання має завантажувати на свій пристрій матеріали, які пов'язані з певною множиною навчальних концептів. Проте існуючи методи будування навчальної траєкторії не надають рекомендацій щодо формування подібних множин. Отже об'єкт навчання може створити заплановану в навчальному курсі послідовність вивчення концептів, чим погіршить власне розуміння навчального матеріалу.

Мета роботи – розробка методу будування навчальної траєкторії для використання в умовах мобільного навчання, який забезпечує поліпшення розуміння навчального матеріалу в умовах обмежених у часі сеансів навчання.

Математична модель структури навчального матеріалу. Як показано в [5], для аналізу структурованості і логічної зв'язності навчального матеріалу – головних факторів, що визначають його якість – доречно використовувати апарат асоціативних зв'язків і побудовану на їх основі асоціативну карту навчального курсу. *Асоціативна карта* – це змішаний граф, вершинам якого відповідають концепти навчального матеріалу, а ребрам/дугам – визначені на цих концептах асоціативні зв'язки. Цей граф є зваженим, вагові коефіцієнти ребер/дуг визначаються мірою асоціативного зв'язку $ass(c_i, c_j)$, який моделюється відповідними ребрами/дугами. Очевидно, що множинам пов'язаних концептів навчального курсу відповідають підграфи асоціативної карти.

Ключовою властивістю шуканих множин пов'язаних концептів – навчальних модулів – є їх незалежність. Для того, щоб отримати незалежний навчальний модуль слід дотримуватися принципів зв'язності і зчеплення [6]. *Принцип зв'язності* полягає в тому, що кожен навчальний модуль має бути сконцентрований на одній і тільки на одній цілі. Відбір і організація контенту і діяльності для навчального модуля зосереджена на його цілі. *Принцип зчеплення* стверджує, що навчальний модуль має

бути як найменше прив'язаний до інших модулів. Тобто зміст навчального модуля не повинен посилається та використовувати матеріал іншого навчального модуля таким чином, щоб створювати необхідні залежності, оскільки потім цей модуль не може бути використаний незалежно від інших.

Використання принципів зв'язності та зчеплення приводить до будування декомпозиції, яка забезпечує простоту додавання, зміни і видалення навчальних модулів, а також можливість їх повторного використання. Асоціативна карта може стати основою для будування такої декомпозиції.

Алгоритм будування декомпозиції. В умовах обмеження у часі сеансів навчання виникає потреба представляти навчальний матеріал компактними порціями, які з одного боку будуть достатньо цілісними, а з іншого – досить самостійними. Для виконання такої декомпозиції навчального матеріалу пропонується розкласти асоціативну карту на компоненти сильної зв'язності. Компонента сильної зв'язності орієнтованого графа $G=(V, E)$ є максимальною множиною вершин $T \subseteq V$, такою що пари вершин u та v з T досяжні одна з одною. Для ефективного пошуку компонент сильної зв'язності звичайно використовують алгоритм Косарайю, який дозволяє виконати пошук таких компонент за лінійний час та пам'ять [7].

Необхідно зазначити, що цей алгоритм розрахований на пошук компонент сильної зв'язності тільки в орієнтованому графі. В нашому випадку, як було зазначено вище, асоціативна карта є змішаним графом, тому виникає потреба в модифікації алгоритму. Сформулюємо алгоритм будування декомпозиції змішаного графу на компоненти сильної зв'язності:

1. Початкове перетворення змішаного графу.

а. Кожне ребро змішаного графу замінити на дві протилежно спрямовані дуги.

б. Якщо в графі між парою вершин v та u існує три дуги, дві з яких спрямовані в одному напрямку і одна в протилежному, замінити їх всі дугою в напрямку двох односпрямованих дуг.

Як результат отримуємо орієнтований граф.

2. Інвертування ребер орієнтованого графа. Як результат отримуємо звернений граф.

3. Виконання пошуку в глибину на зверненому графі. Як результат отримуємо вектор обходу вершин зверненого графу.

4. Виконання пошуку в глибину на орієнтованому графі, отриманому на кроці 1, з вибором не відвіданої вершини з максимальним

порядковим номером у векторі обходу, отриманому на кроці 3. Як результат отримуємо ліс з дерев, кожне з яких відповідає компоненті сильної зв'язності.

Метод будування навчальної траєкторії. В умовах мобільного навчання задача будування навчальної траєкторії є, фактично, задачею визначення порядку вивчення навчальних модулів. Для її рішення слід виконати такі дії:

1. Побудувати декомпозицію асоціативної карти навчального курсу на компоненти сильної зв'язності.

2. Згорнути множини вершин асоціативної карти, які входять до компонент сильної зв'язності, в мета-вершини.

3. Розрахувати вагові коефіцієнти дуг отриманого конденсату асоціативної карти як

$$ass(c_k^{SC}, c_l^{SC}) = \sum_{c_i \in c_k^{SC}} \sum_{c_j \in c_l^{SC}} ass(c_i, c_j),$$

де $ass(c_i, c_j)$ – міра асоціативного зв'язку між навчальними концептами c_i та c_j ; c_k^{SC}, c_l^{SC} – вершини конденсату асоціативної карти, що відповідають k -й та l -й компонентам сильної зв'язності.

4. Визначити порядок вивчення навчальних модулів.

В [8] сформульовано алгоритм будування навчальної траєкторії на основі асоціативної карти навчального курсу, який може бути застосований і для її конденсату з метою визначення порядку вивчення навчальних модулів.

Висновки. Як результат виконаної роботи запропоновано метод будування навчальної траєкторії в умовах мобільного навчання. Цей метод полягає в знаходженні на асоціативній карті навчального курсу компонент сильної зв'язності і будуванні навчальної траєкторії на конденсаті асоціативної карти.

Використання запропонованого методу будування навчальної траєкторії дозволяє підвищити ефективність процесу мобільного навчання. Об'єкт навчання при поєднанні з навчальним ресурсом завантажуватиме на свій мобільний пристрій навчальні матеріали, що стосуються відносно незалежної множини тісно пов'язаних навчальних концептів. Порядок вивчення навчальних модулів строго визначатиметься на основі асоціативних зв'язків, що існують між навчальними концептами, тобто на основі об'єктивних характеристик моделі предметної області навчального курсу, а не особистих преференцій об'єкту навчання.

Список літератури: 1. *El-Hussein M.O.M.* Defining Mobile Learning in the Higher Education Landscape [Електронний ресурс] / *M.O.M. El-Hussein, J.C. Cronje* // Educational Technology & Society. – 2010. – 13 (3). – Р. 12 – 21. – Режим доступу до журн.: http://findarticles.com/p/articles/mi_7100/is_3_13/ai_n56337671/. 2. *Nwana H.S.* Intelligent Tutoring Systems: an overview [Електронний ресурс] / *H.S. Nwana* // Artificial Intelligent Review. – 1990. – 4. – Р. 251 – 277. – Режим доступу до журн.: <http://www.compassproject.net/sadhana/teaching/readings/its.pdf>. 3. *Martins A.C.* User Modeling in Adaptive Hypermedia Educational [Електронний ресурс] / *A.C. Martins, L. Faria, C. Vaz de Carvalho, E. Carrapatoso* // Systems. Educational Technology & Society. – 2008. – 11 (1). – Р. 194 – 207. – Режим доступу до журн.: http://www.ifets.info/journals/11_1/14.pdf. 4. *De Bra P.* Adaptive Web-based Educational Hypermedia [Електронний ресурс] / *P. De Bra, L. Aroyo, A. Cristea* // Web Dynamics, Adaptive to Change in Content, Size, Topology and Use / Eds. Levene M., Poulouvasilis A. – Springer, 2004. – Р. 387 – 410. – Режим доступу до книги: <http://www.wis.win.tue.nl/~debra/dm-elearning.pdf>. 5. *Любченко В.В.* Асоціативна карта для аналізу якості навчальних курсів / *В.В. Любченко, І.І. Саприкін, Є.О. Сичов, О.С. Шинкарюк* // Електромашинобудування та електрообладнання. – Вип. 72. – К.: Техніка, 2009. – С. 208 – 211. 6. *Boyle T.* Design principles for authoring dynamic, reusable learning objects [Електронний ресурс] / *T. Boyle* // Australian Journal of Educational Technology. – 2003. – 19(1). – Р. 46 – 58. – Режим доступу до журн.: <http://www.ascilite.org.au/ajet/ajet19/boyle.html>. 7. *Sedgewick R.* Algorithms in Java / *R. Sedgewick* – Addison Wesley, 2002. – 768 р. 8. *Любченко В.* Модифікований алгоритм топологічного сортування для будування навчальної траєкторії / *В. Любченко, І. Саприкін* // Технічні вісті. – 2010. – № 1 (31), 2 (32). – С. 163 – 165.

Стаття представлена д.т.н., проф. ОНПУ Крісіловим В.А.

УДК 004.02

Метод построения учебной траектории в условиях мобильного обучения / *Любченко В.В., Шинкарюк А.С.* // Вестник НТУ "ХПИ". Тематический выпуск: Информатика и моделирование. – Харьков: НТУ "ХПИ". – 2011. – № 17. – С. 81 – 85.

Рассмотрено влияние на требования к учебному материалу ограниченности во времени учебных сеансов в условиях мобильного обучения. Модифицирован алгоритм определения компонент сильной связности для применения его в смешанных графах с целью получения декомпозиции учебного курса. Предложен метод построения учебной траектории на основе множеств связанных учебных концептов. Библиогр.: 8 назв.

Ключевые слова: мобильное обучение, компоненты сильной связности, декомпозиция учебного курса, учебная траектория.

UDC 004.02

Method of learning trajectory building for m-learning / *Liubchenko V.V., Shynkariuk O.S.* // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – № 17. – Р. 81 – 85.

Impact of time constraints of learning sessions in mobile learning on requirements for learning materials is considered. The algorithm for strong components discovering is modified for using on the mixed graphs to obtain the syllabus decomposition. A method of learning trajectory building based on a set of related educational concepts is proposed. Refs.: 8 titles.

Key words: mobile learning, strong components, syllabus decomposition, learning trajectory.

Надійшла до редакції 15.02.2011

А.Н. ЛЫСЕНКО, д.т.н., проф. каф. КЭВА, НТУУ "КПИ", Киев,
А.Ю. РОМАНОВ, аспирант, асс. каф. КЭВА, НТУУ "КПИ", Киев

РЕСУРСОЭФФЕКТИВНЫЙ РОУТЕР ДЛЯ МНОГОПРОЦЕССОРНОЙ СЕТИ НА ЧИПЕ

Рассмотрены различные подходы к организации сетей на чипе. Выявлен основной недостаток сетей на чипе с коммутацией пакетов – чрезмерно большие объемы входных и выходных буферов роутеров. Предложена новая архитектура роутера с улучшенными показателями потребляемых ресурсов и высоким быстродействием. Ил.: 5. Библиогр.: 12 назв.

Ключевые слова: сеть на чипе, коммутация пакетов, буфер, архитектура роутера.

Постановка проблемы. Развитие полупроводниковых технологий привело к тому, что огромное количество транзисторов, доступное на одном кристалле FPGA, позволяет разработчикам интегрировать десятки IP-ядер в одну многопроцессорную систему [1]. Такие системы на чипе (MPSoC), благодаря параллельной обработке потоков данных имеют большое быстродействие и пропускную способность, что и обусловило их перспективность и популярность. Однако с увеличением количества вычислителей на одном чипе значительно возросли требования к подсистеме связи по пропускной способности, потребляемым ресурсам и т.д., а такие классические подходы как односвязная архитектура и общая шина стали неэффективны [1 – 4]. Решением проблемы стало появление концепции объединения вычислительных ядер в сеть – NoC (Network-on-Chip). Сутью данного подхода является объединение IP-ядер, обычно представляющих собой процессорное ядро (PN) с локальной памятью и дополнительными устройствами, с помощью специализированных роутеров – SW (switch) [2] (рис. 1).

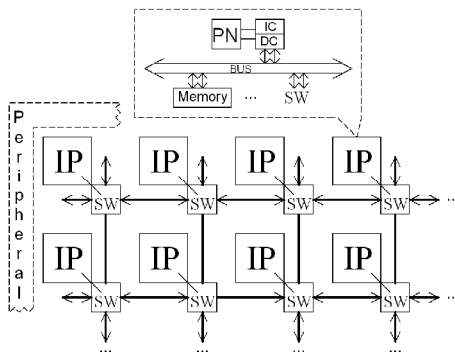


Рис. 1. Типовая структура сети на чипе

Преимущество данной концепции перед классическими состоит в том, что короткие соединения точка-точка между роутерами обеспечивают малую паразитную емкость и высокую частоту работы благодаря наличию промежуточных регистров на пути прохождения сигнала, а данные в различных сегментах сети передаются и коммутируются одновременно [3]. Это обеспечивает высокую производительность, пропускную способность и экономию ресурсов, что делает исследования и разработку новых архитектур сетей на чипе актуальными.

Анализ литературы. Наиболее распространенной является архитектура сетей на чипе с коммутацией пакетов, где обычно применяется wormhole-технология соединения с кредитным управлением потоком и XY – маршрутизацией [2, 4 – 8]. Главным недостатком таких систем является необходимость в буферизации соединений на входе, а иногда и на выходе, что приводит к большим затратам ресурсов. Уменьшение объема буферов малоэффективно и чревато потерей производительности. В работах [9, 10] предложены реализации сетей на чипе с коммутацией на уровне цепей, при этом экономия ресурсов достигается за счет отсутствия буферизации. В этих подходах падение быстродействия происходит из-за того, что на время передачи данных резервируется полный путь от источника до адресата и блокируются другие передачи [11].

Таким образом, на фоне многообразия различных вариантов реализаций сетей на чипе существует большая потребность в разработке новых архитектур и их составных элементов, обеспечивающих высокую пропускную способность при минимальных ресурсных затратах. Кроме этого, свой отпечаток накладывает тот факт, что данное направление является сравнительно новым в электронике и недостаточно разработанным, а большинство проектов – коммерческими, сориентированными на отсутствующую в Украине технологию ASIC [1, 3, 4, 6], а FPGA реализации – на платформу Xilinx [7, 8].

Цель статьи – разработка новой схемы роутера с низкими требованиями к объему буферных элементов и высоким быстродействием для сетей на чипе на базе FPGA платформы фирмы Altera.

Используемая архитектура сети на чипе. Авторами выбрана в качестве рабочей архитектура тороидальной сети на чипе с коммутацией пакетов. Выбор данной топологии обусловлен ее простотой и эффективностью, а также идентичностью структур роутеров [2, 5, 11, 12]. Применение технологии с коммутацией на уровне пакетов обусловлено

ее большей эффективности в сравнении с коммутацией на уровне цепей [11]. Поскольку предполагаемое количество IP-ядер составляет не более 16 и задержки синхросигнала будут незначительные, выбрана сеть с глобальной синхронизацией.

Предлагаемая архитектура роутера. Наиболее распространенная архитектура роутера для сетей на чипе приведена на рис. 2.

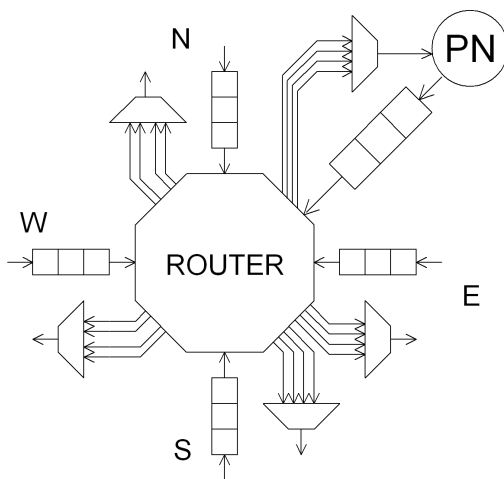


Рис. 2. Классическая структура роутера

Роутер представляет собой коммутационную матрицу, направляющую на соответствующие выходы потоки флитов, накапливаемые во входных буферных элементах (N, E, S, W) [2, 5, 7, 11]. Кроме входов, буферные элементы также могут присутствовать и на выходах роутера, их размер варьируется в зависимости от требуемой пропускной способности и длины пакетов. Использование такой архитектуры приводит к большим ресурсным затратам и энергопотреблению, а из-за неравномерного характера загрузки сети значительная часть буферных элементов может простаивать.

С целью устранения недостатков описанной выше архитектуры и уменьшения буферной памяти авторами предложено разделение функции коммутационной части роутера на входной и выходной блоки, соединенные между собой буферной памятью типа FIFO (First Input First Output) объемом в 16 флитов (рис. 3).

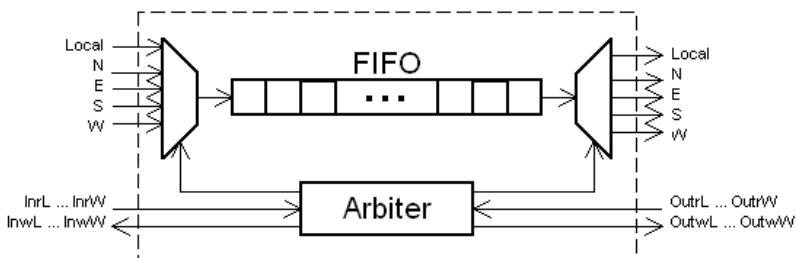


Рис. 3. Предложенная структура роутера

Управление коммутацией с локальным IP-ядром и соседними роутерами (N, E, S, W) осуществляет распределенный арбитр с помощью сигналов разрешения и подтверждения приема/передачи (Inr-InrW, InwL-InwW, OutrL-OutrW, OutwL-OutwW). Очевидно, что предложенная сеть на чипе относится к классу сервиса с лучшей пропускной способностью и FCFS (First Come First Serve) арбитражем [2, 5, 11].

Входной и выходной коммутационные модули роутера созданы на высокоуровневом языке описания аппаратуры Verilog, а модуль FIFO взят из библиотеки мегафункций среды проектирования Quartus II фирмы Altera, поскольку разрабатываемая сеть на кристалле ориентирована на FPGA данной фирмы. Кроме того, реализации стандартных мегафункций универсальны, легко настраиваемы, что облегчает процесс проектирования, и написаны на низкоуровневом языке AHDL. Это делает их более компактными, быстродействующими и оптимизированными под архитектуру ПЛИС в сравнении с аналогичными реализациями на других языках.

Особенности алгоритма передачи данных. Адрес назначения сообщений формируется по месту их генерации и является однонаправленным. Сами сообщения разбиваются на пакеты длиной от 1 до 8 флитов каждый. Флит состоит из 32 бит данных, дополнительного бита, единичное значение которого идентифицирует, что данный флит последний в пакете, а также 4 адресных бита (рис. 4).

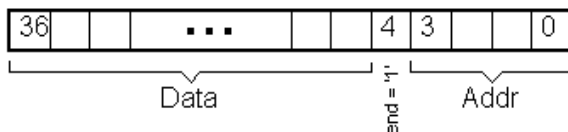


Рис. 4. Структура флита

Таким образом, в сети может быть до 16 роутеров, каждому из которых присвоен уникальный адрес.

Принятие решения о коммутации пакетов осуществляется по классическому детерминистическому алгоритму *XY*, т.е. сначала пакеты распространяются по сети в горизонтальном направлении, а потом – в вертикальном. Очевидно, что данный алгоритм охватывает только часть возможных путей прохождения пакетов (partial), хотя и направляет пакеты всегда с приближением к цели (profitable) [11].

Передача данных осуществляется с подтверждением (рис. 5).

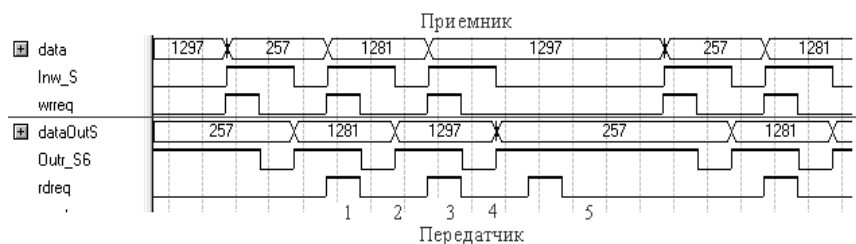


Рис. 5. Прием/передача данных

Передатчик проверяет наличие в FIFO данных, генерирует сигнал чтения и получает один флит (шаг 1), анализирует адрес назначения, согласно с алгоритмом *XY* выставляет данные на соответствующий порт и генерирует сигнал запроса на передачу (шаг 2). В свою очередь, приемник соседнего роутера при наличии места в FIFO анализирует сигналы запроса на передачу, фиксирует одно из направлений и коммутует данные на вход FIFO, генерируя сигнал записи (шаг 3). Одновременно передатчик сбрасывает сигнал передачи и готовится к следующей итерации, генерируя сигнал чтения из FIFO. По завершении приема приемник генерирует сигнал подтверждения передачи и готовности принимать следующий флит (шаг 4). Кроме того, он анализирует end-бит флита, и если это конец пакета, прежде чем продолжить прием данных с текущего направления, проверяет наличие запросов с других направлений, во избежание длительных блокировок (шаг 5).

Апробация. Разработанный роутер был синтезирован в среде проектирования Quartus II для FPGA Cyclone II, занимает 250 Les, 592 бита памяти. При этом статическое и динамическое потребление энергии составило соответственно 40,43 мВт и 81,44 мВт, а максимальная частота работы – 200 МГц. Так как минимальное время прохождения флита составляет 3 такта (рис. 5), а его информативная часть – 32 бита, то

пропускная способность маршрутизатора досягає $\sim 2,13$ Гбіт/с, що порівнювано з існуючими аналогами [7, 8], при значно меншій ресурсній витраті. При цьому в залежності від базової платформи FPGA тактова частота може бути значно вище (наприклад, 310 МГц для Stratix II).

Висновки. Предложено нову архітектуру маршрутизатора для багатопроцесорної мережі на чіпі з комутацією пакетів, відмінною особливістю якої є розділення функцій його комутаційної частини на вхідну та вихідну блоки, з'єднані між собою буферною пам'яттю типу FIFO. Це дозволило в порівнянні з вказаними вище аналогами в декілька разів зменшити витрати ресурсів на організацію буферних елементів, зберігши при цьому показники швидкодії.

Перспективним напрямком подальших досліджень є вдосконалення маршрутизатора шляхом введення віртуальних каналів, зменшення часу проходження пакета до 2 тактів, рознесення приймача та передавача на різні фронти тактової частоти.

Список літератури: 1. Пат. 03013591A1 U.S., МПК G06F17/50. Method to design network-on-chip (NOC) – based communication systems / *Murali S., Benini L., De Micheli G.*; заявитель и патентообладатель EPFL (CH). – заявл. 10.10.2006; опубл. 28.01.2009. 2. *Axel J.* Networks on Chip / *J. Axel, T. Hannu* // Kluwer Academic Publishers. – Dordrecht, 2003. – 303 p. 3. *Angiolini F.* A layout-aware analysis of networks-on-chip and traditional interconnects for mpsoes / *F. Angiolini, P. Meloni, M.S. Carta* // IEEE Transactions on Computer Aided Design of Integrated Circuits and Systems. – 2007. – Vol. 26. – № 3. – P. 421–434. 4. *Benini L.* Network-on-chip architectures and design methods / *L. Benini, D. Bertozzi* // Computers and Digital Techniques. – IEEE Proc., 2005. – Vol. 152. – No. 2. – P. 261–272. 5. *Bjerregaard T.* A survey of research and practices of Network-on-chip / *T. Bjerregaard, S. Mahadevan* // ACM Computing Surveys. – 2006. – Vol. 38(1). – 51 p. 6. *Rijkema E.* Trade-offs in the design of a router with both guaranteed and best-effort services for networks on chip / *E. Rijkema, K.G.W. Goossens, A. Radulescu*. – IEEE Proc., 2003. – Vol. 150. – № 5. – P. 294–302. 7. *Moraes F.G.* HERMES: an Infrastructure for Low Area Overhead Packet-switching Networks on Chip / *F.G. Moraes, N.L.V. Calazans, A.V. Mello* // Integration, the VLSI Journal. – 2004. – Vol. 38. – No. 1. – P. 69–93. 8. *Marescaux T.* Interconnection Networks Enable Fine-Grain Dynamic Multi-Tasking on FPGAs / *T. Marescaux* // FPL'02. – 2002. – P. 795–805. 9. *Phi-Hung P.* High Performance and Area-Efficient Circuit-Switched Network on Chip Design / *P. Phi-Hung* // IEEE CIT'06. – 2006. – P. 243. 10. *Wolkotte P.T.* An Energy-Efficient Reconfigurable Circuit-Switched Network-on-Chip / *P.T. Wolkotte, G.J.M. Smit, G.K. Rauwerda*. – IEEE Proc. IPDPS, 2005. – Vol. 4. – P. 155a. 11. *Ankur A.* Survey of Network on Chip (NoC) Architectures & Contributions / *A. Ankur, I. Cyril, S. Ravi* // Engineering, Computing & Architecture. – 2009. – Vol. 3(1). – 15 p. 12. *Романов О.Ю.* Аналіз топологій мереж на чіпі / *О.Ю. Романов* // Сучасна інформаційна Україна: інформатика, економіка, філософія: матеріали доповідей конференції, 13-14 травня 2010 р. – Донецьк, 2010. – Т. 1. – С. 407–410.

УДК 004.272

Ресурсоефективний маршрутизатор для багатопроцесорної мережі на чіпі / *Лисенко О.М., Романов О.Ю.* // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 86–92.

Розглянуто різноманітні підходи щодо організації мереж на чипі. Виявлено основний недолік мереж на чипі з комутацією пакетів – надмірно великі обсяги вхідних і вихідних буферів роутерів. Запропоновано нову архітектуру роутера з поліпшеними показниками споживаних ресурсів і високою швидкістю. Іл.: 5. Бібліогр.: 12.

Ключові слова: мережа на чипі, комутація пакетів, буфер, архітектура роутера.

UDC 004.272

Resource efficient router for multiprocessor network on chip / Lisenko O.M., Romanov O.Y. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 86 – 92.

Various approaches to networks on chip organizing are considered. The main drawback of networks on chip packet switching is identified – an excessively large buffers amounts of input and output buffers of routers. The new router architecture with improved resource consumption and high speed action is offered. Figs.: 5. Refs.: 12 titles.

Keywords: network on chip, packet switching, the buffer, router architecture.

Поступила в редакцію 07.02.2011

И.Е. МАРОНЧУК, д.т.н., директор НИЦ НФНМиРТЭ СНУЯЭиП,
Севастополь,

И.И. МАРОНЧУК, к.т.н., зав. НИЛ ПФиНТЭ СНУЯЭиП,
Севастополь,

А.Н. ПЕТРАШ, ст. преп. СНУЯЭиП, Севастополь

NANO-S: ПАКЕТ ДЛЯ РАСЧЕТА ЭНЕРГЕТИЧЕСКИХ СПЕКТРОВ ЭЛЕКТРОНОВ В НАНОГЕТЕРОСТРУКТУРАХ С КВАНТОВЫМИ ТОЧКАМИ

Описан пакет Nano-S, предназначенный для расчета энергетических спектров электронов в наногетероструктурах с ограниченным числом степеней свободы. Кратко рассмотрены математические модели таких структур в приближении эффективной массы. Указываются основные методы численного решения нелинейных алгебраических и дифференциальных уравнений, появляющихся в результате решения уравнения Шредингера с соответствующим модельным потенциалом, которые реализуются в пакете. Библиогр.: 11 назв.

Ключевые слова: наногетероструктуры, модель, уравнение Шредингера.

Постановка проблемы. Полупроводниковые наногетероэпитаксиальные структуры с квантовыми точками [1] привлекают все большее внимание исследователей в связи с перспективами создания на их основе новых поколений существующих приборов, например, солнечных батарей 3-го поколения с эффективностью свыше 50 % [2], а также приборов наноэлектроники. Электронные процессы в таких материалах нельзя описать на основе представлений о газе квазичастиц, которые используются для описания электронных процессов в твердом теле, а технология получения материалов с закономерно расположенными квантоворазмерными объектами требует разработки новых нестандартных, самоорганизующихся процессов, которые не используются в технологии получения обыкновенных кристаллов и эпитаксиальных структур. Большое количество экспериментальных данных, полученных в последнее время, требует теоретического осмысления, что, в свою очередь, выражается в построении соответствующих адекватных математических моделей. Подавляющее большинство таких моделей основывается на численном решении уравнения Шредингера для различных потенциалов, а также трансцендентных уравнений, которые появляются в процессе теоретического анализа. Удобнее всего представлять полученные результаты в графической форме. Возникает необходимость создания пакета, позволяющего рассчитывать энергетические спектры электронов, находящихся в квантовых точках различных размеров и разной геометрией.

Анализ литературы. В настоящее время имеется много пакетов для моделирования процессов, происходящих в наноструктурах. Рассмотрим наиболее популярные из них. Одним из таких пакетов является QUANTUM ESPRESSO (QE) [3]. Его код основывается на теории функционала плотности, расчеты ведутся с использованием теории псевдопотенциала в базисе плоских волн. На аналогичных принципах основан другой, не менее популярный пакет ABINIT [4]. Оба эти пакета позволяют рассчитать общую энергию, плотность заряда и электронную структуру систем электронов и атомов, таких как молекулы и твердые тела. Кроме того, эти пакеты, используя функционал плотности, позволяют оптимизировать геометрию и вычислить силы, действующие в этих системах, провести моделирование методом молекулярной динамики.

Однако, несмотря на большие возможности, оба пакета обладают рядом недостатков, затрудняющих их использование исследователями-экспериментаторами при создании новых материалов с квантоворазмерными структурами. Во-первых, не всегда требуются такие громоздкие вычисления на основе теории псевдопотенциала. Часто бывает достаточно рассмотреть эмпирические и более простые по своей сути модели, например, в приближении эффективной массы. Во-вторых, для работы с такими пакетами необходимы хорошие навыки работы в UNIX-системах, в частности, необходимо уметь работать в командной строке (например, в `bash`), уметь программировать в такой оболочке хотя бы на начальном уровне. Наконец, у всех пакетов практически отсутствует графический интерфейс. Правда, для QE имеется написанный на Tcl интерфейс, который позволяет упростить процесс ввода данных для этого пакета. Тем не менее, для визуализации полученных результатов приходится применять другие пакеты, такие как GNUPLOT или GRACE. Для обучения сотрудников навыкам работы с этими пакетами необходимо определенное время, что негативно сказывается на скорости проведения исследований. Все это подвигло нас разработать визуальный, интуитивно понятный, быстрый пакет для расчетов энергетических спектров электронов на основе простых, но достаточно правдоподобных моделях структур с квантоворазмерными объектами.

Целью статьи является описание разработанного в НИЛ "Прикладной физики и нанотехнологий в энергетике" СНУЯЭиП пакета Nano-S, предназначенного для расчета энергетического спектра носителей зарядов в различных низкоразмерных структурах, таких как квантовые ямы, нити или точки, с различной геометрией поверхности, а также для моделирования физических процессов, происходящих в таких

наноструктурах. Наши вычисления основываются на эмпирических моделях таких квантовых объектов в приближении эффективной массы.

Математическая модель. Движение носителей зарядов (электронов, дырок) в полупроводниковых низкоразмерных структурах в приближении эффективной массы описывается уравнением Шредингера [5, 6]

$$\Delta\psi(\vec{r}) + \frac{2m}{\hbar^2} [E - U(\vec{r})]\psi(\vec{r}) = 0. \quad (1)$$

Здесь $\psi(\vec{r})$ – волновая функция электрона, m – его эффективная масса, E – его полная энергия, а $U(\vec{r})$ – потенциал ямы, в которой находится электрон. Как обычно $\hbar = h/(2\pi)$, а h – постоянная Планка, Δ – оператор Лапласа.

Для разных моделей наногетероструктур используются различные выражения для потенциала $U(\vec{r})$. Кроме того, в зависимости от удерживаемых степеней свободы для соответствующих объектов уравнение (1) будет представлять собой уравнение с одной, двумя или тремя степенями свободы. Ниже мы кратко опишем основные модельные потенциалы и укажем те численные методы решения уравнения Шредингера, которые используются в описываемом пакете.

Наибольший интерес для исследователей, занимающихся созданием материалов с низкоразмерными структурами, представляют такие наноструктуры как квантовые точки. Под квантовыми точками мы будем понимать такие объекты, размеры которых сравнимы с де-бройлевской длиной волны электрона. Сам электрон не имеет вообще степеней свободы. Он находится в соответствующей квантовой яме. Параметры ямы такие, как ее глубина, геометрия, могут быть заданы аналитическим или численным видом потенциала $U(\vec{r})$, а также симметрией всего уравнения (1).

Самым простым является случай квантовой точки в виде прямоугольного параллелепипеда с длинами ребер a , b и c . В случае бесконечно глубокой ямы задача допускает точное решение [5]. Для ямы конечной глубины U_0 уровни энергии находятся из следующего трансцендентного уравнения [6]

$$\arcsin \frac{k\hbar}{\sqrt{2mU_0}} = \frac{ka - \pi n}{2}, \quad (2)$$

где $k = \sqrt{2mE} / \hbar$, $n = 1, 2, \dots$ – номер энергетического уровня. Это уравнение легко может быть решено с помощью метода Ньютона. Аналогично считаются уровни энергии для мод вдоль двух других осей. Сумма энергий, рассчитанных для всех трех удерживаемых степеней свободы, и есть энергетический уровень электрона в такой квантовой точке.

Для цилиндрической квантовой нити бесконечной длины задача нахождения уровней энергии сводится к решению уравнения [7]

$$\frac{kJ_0'(kR)}{J_0(kR)} = \frac{\kappa K_0'(\kappa R)}{K_0(\kappa R)}, \quad (3)$$

где J_0, K_0 – функции Бесселя и Макдональда нулевого порядка соответственно, R – радиус нити, а

$$\kappa^2 = -\frac{2mE}{\hbar^2}, \quad k^2 = \frac{2m(U_0 - |E|)}{\hbar^2}, \quad (4)$$

U_0 – глубина ямы. Для численного решения уравнения (3) необходимо уметь вычислять функции Бесселя и Макдональда и их производные. Такие алгоритмы широко известны.

Если же мы рассматриваем квантовую нить конечной длины и, особенно, квантовую точку, где высота цилиндра сравнима с его радиусом основания, необходимо численно решать соответствующую краевую задачу на собственные значения. Удобным в этом случае оказывается метод пристрелки.

В случае сферической квантовой точки уровни энергии, помимо главного квантового числа n , характеризуются еще угловым l и магнитным m квантовыми числами. Для $m = 0$, $l = 0$ уровни энергии могут быть получены из следующего трансцендентного уравнения [8]

$$E = \frac{\pi^2 \hbar^2}{2md^2} \left(n - \frac{1}{\pi} \operatorname{arctg} \sqrt{\frac{E}{U_0 - E}} \right)^2. \quad (5)$$

Здесь d – диаметр сферической квантовой точки. Аналогичные выражения могут быть найдены для случаев с другими значениями указанных квантовых чисел.

Во всех указанных выше моделях предполагается, что на границе наноструктуры потенциал возрастает скачком с нулевого значения до U_0 . Более реалистические модели требуют изменения поведения соответствующего потенциала. Пусть r – характерный размер квантовой

точки, в зависимости от которого необходимо найти энергетический спектр электрона. Часто используется так называемый параболический потенциал

$$U(r) = \frac{m\omega^2 r^2}{2}, \quad (6)$$

широко известный из теории квантового гармонического осциллятора [5, 6]. Параметр ω здесь выбирается так, чтобы на границе квантовой точки потенциал равнялся бы U_0 . В разное время были предложены другие модельные потенциалы. Среди них:

– потенциал Пешля-Теллера [9]

$$U(r) = -\frac{\hbar^2 \alpha^2}{2m} \frac{\lambda(\lambda-1)}{\text{ch}^2 \alpha r}, \quad (7)$$

где безразмерный параметр $\lambda > 1$, и параметр α , с размерностью обратной длине, определяют вместе глубину и ширину потенциальной ямы;

– гауссов потенциал [10]

$$U(r) = -U_0 \exp\left(-\frac{r^2}{2R^2}\right), \quad (8)$$

где R – область действия потенциала удержания, U_0 – глубина ямы;

– потенциал де Филиппо-Салерно [11]

$$U(r) = -\frac{U_0}{\left(1 + r^2 / R^2\right)^2}, \quad (9)$$

где параметр R является настроечным.

Эти потенциалы на малых расстояниях ведут себя как параболический (6), а на больших расстояниях асимптотически приближаются к нулю. В описываемом нами пакете все эти потенциалы реализованы.

Описание пакета Nano-S. Пакет для моделирования электронных процессов в низкоразмерных квантовых структурах написан на языке C++. Интерфейс был создан с использованием библиотек Qt. Кроме того, для отображения графиков зависимости уровней энергии от размера квантовой точки использовались дополнения для Qt, а именно библиотеки QWT. Все это позволило создать кросс-платформенное приложение, которое с одинаковой легкостью может быть скомпилировано под операционными системами Windows, Linux и MacOS. Пакет должен собираться и под другими операционными

системами, куда портированы библиотеки Qt, однако мы явно это не проверяли.

Необходимость написания такого кросс-платформенного приложения была продиктована разными задачами, решаемыми в нашей лаборатории. Как уже указывалось выше, под ОС Linux имеется ряд пакетов (Quantum Espresso, Abinit), предназначенных для расчета электронных процессов, происходящих в твердых телах методом псевдопотенциала, другими более сложными методами. Так как Nano-S может работать под ОС Linux, исследователь имеет возможность сравнить результаты вычислений в приближении эффективной массы, на котором основан описываемый пакет, с результатами более сложных и громоздких вычислений, предоставляемых сторонними пакетами. В результате такого сравнения можно делать вывод об адекватности соответствующей модели.

С другой стороны, управление лабораторными установками, снятие и обработка результатов измерений происходит с помощью компьютеров, на которых установлена ОС Windows. Имея возможность запустить пакет Nano-S, экспериментатор, работающий в этой системе, может также проверить согласованность различных моделей, использующих различные потенциалы, с экспериментальными данными и в результате этого подобрать наилучшее приближение модели к эксперименту.

Интерфейс пакета разрабатывался с учетом пожеланий сотрудников лаборатории. Мы стремились сделать его интуитивно понятным для любых исследователей, в том числе и для тех, которые имеют низкий уровень компьютерной подготовки.

В верхнем левом углу расположены два выпадающих списка, позволяющих выбрать материал, из которого изготовлена квантовая точка, а также материал матрицы, в которой эта точка была выращена. Мы предусмотрели также идеальные случаи бесконечно глубоких ям и ввели соответствующий бесконечный потенциал. Пользователь может вручную ввести данные для соответствующих объектов. Для этого имеются соответствующие переключатели, блокирующие выпадающие списки и активирующие поля ввода.

В пакете реализованы расчеты энергетического спектра квантовых точек с разной геометрией (кубической, цилиндрической, сферической, пирамидальной, эллиптической). Для выбора соответствующей модели имеется выпадающий список, расположенный ниже элементов управления для выбора материалов структур. Для каждой модели имеются свои параметры, характеризующие энергетический спектр. Соответствующие поля ввода для параметров расчета автоматически

активируются в зависимости от типа модели. Для линейного размера нужно указать параметр изменения и, опционально, шаг.

По нажатию кнопки "Рассчитать спектр" производятся соответствующие вычисления, которые отображаются в таблице, и строится график. Таблицу с данными можно сохранить в виде текстового файла для последующей обработки в других программах. График также можно сохранить в виде графического файла. Стоит отметить, что параметры графика, такие как цвет, толщина, тип линии, являются настраиваемыми.

Выводы. Описан пакет Nano-S, разработанный и широко используемый в НИЛ "Прикладной физики и нанофизики в энергетике" СНУЯЭиП. За все время использования он показал себя как эффективный помощник сотрудникам лаборатории, которые занимаются получением и исследованием наногетероэпитаксиальных структур, на основе которых впоследствии создаются высокоэффективные солнечные элементы 3-го поколения. Проводимые на нем расчеты позволяют упростить подход к выбору необходимых для проведения технологических процессов материалов подложки, квантовых гетероструктур, их размеров и геометрии. Пакет легко может быть скомпилирован под разные платформы, включая такие популярные как Windows и Linux. Для работы с пакетом не требуется подключения к сети Internet, что позволяет установить его на локальный компьютер оператора, управляющего определенной исследовательской установкой. Пакет позволяет легко интегрировать в себя дополнительные программные модули. Это открывает широкие перспективы к его дальнейшему усовершенствованию и наполнению другими, еще не представленными в нем, функциями.

Список литературы: 1. *Knoss R.W.* Quantum dots: research, technology and applications / *R.W. Knoss.* – New York: Nova Science Publ. – 2009. – P. 708. 2. *Cuadra L.* Intermediate band photovoltaic overview / *L. Cuadra, A. Marti, N. Lopez, A. Luque.* – Proceedings of 3rd World Conference on Photovoltaic Energy Conversion, Osaka, Japan. – 2003. – V. 1. – P. 3-8. 3. *Gianozzi P.* QUANTUM ESPRESSO: a modular and open source software project for quantum simulations of materials / *P. Gianozzi* // *J. Phys.: Condes. Matter.* – 2009. – V. 21. – № 39. – 395-502. 4. *Gonze X.* ABINIT: first-principle approach to material and nanosystem properties / *X. Gonze* // *Computer. Phys.Comm.* – 2009. – V. 180. – P. 2582-2615. 5. *Ландау Л.Д.* Квантовая механика / *Л.Д. Ландау, И.М. Лифшиц.* – М.: Наука, 1989. – 768 с. 6. *Давыдов А.С.* Квантовая механика / *А.С. Давыдов.* – М.: Наука, 1973. – 704 с. 7. *Flügge S.* Practical quantum mechanics / *S. Flügge.* – Springer, 1999. – 640 p. 8. *Смирнов С.Б.* Расчет энергетического спектра S-электронов сферической квантовой точки на основе узкозонных полупроводниковых соединений $A^{III}B^V$ в матрице GaP / *С.Б. Смирнов, И.Е. Марончук, И.И. Марончук, А.Н. Петраш* // Сб. науч. трудов СНУЯЭиП. – 2011. – № 1 (37). – С. 95-101. 9. *Pöschl G.* Bemerkungen zur Quantenmechanik des anharmonischen Oszillators / *G. Pöschl, E. Teller* // *Zeitschrift für Physik.* – 1933. – V. 83. – №. 3-4. – P. 143-151. 10. *Adamowski J.* Electron pair in

a Gaussian confining potential / *J. Adamowski, M. Sobkowicz, B. Szafran, S. Bednarek* // Physical Review. – 2000. – V. 62. – № 7. – P. 4234-4237. **11.** *De Filippo S.* Spectral properties of a model potential for quantum dots with smooth boundaries / *S. De Filippo, M. Salerno* // Physical Review. – 2000. – V. 62. – № 7. – P. 4230-4233.

УДК 530.145.61

Nano-S: пакет для розрахунку енергетичних спектрів електронів в наногетероструктурах с квантовими точках / Марончук І.Є., Марончук І.І., Петраш О.М. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 93 – 100.

Описан пакет Nano-S, який призначений для розрахунку енергетичних спектрів електронів в наногетероструктурах с обмеженою чисельністю степенів свободи. Кратко розглянуті математичні моделі таких структур у приближенні ефективної маси. Вказуються основні методи чисельного розв'язування нелінійних алгебраїчних та диференціальних рівнянь, які з'являються в результаті рішення рівняння Шредінгера з відповідним модельним потенціалом, які реалізуються у пакеті. Бібліогр.: 11 назв.

Ключові слова: наногетероструктури, модель, рівняння Шредінгера.

UDC 530.145.61

Nano-S: a package for modeling of electron processes in nanoheterojunctions with quantum dots/ Maronchuk I.E., Maronchuk I.I., Petrash A.N. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 93 – 100.

A package Nano-S intended for calculation of an energy spectrum of electrons in nanoheterojunctions with limited numbers of degrees of freedom are described. Mathematical models of low dimensional structures in effective mass approximation are shortly considered. Main methods of a numerical solution of non-linear algebraic and differential equations that appeared in a solving Schrödinger equation with appropriate model potential that released in the package are pointed. Refs.: 11 titles.

Key words: nanoheterojunctions, model, Schrödinger equation.

Поступила в редакцію 15.02.2011

Д.В. МИШИН, зав. лаб., ассистент каф. ИЗИ ВлГУ, Владимир,
М.М. МОНАХОВА, инженер каф. ИЗИ ВлГУ, Владимир

ОБ ОПТИМИЗАЦИИ АДМИНИСТРИРОВАНИЯ КОРПОРАТИВНЫХ СЕТЕЙ ПЕРЕДАЧИ ДАННЫХ В УСЛОВИЯХ ОГРАНИЧЕННЫХ АДМИНИСТРАТИВНЫХ РЕСУРСОВ АСУП

В работе предлагается модель приоритетов функциональных элементов корпоративной сети передачи данных, предлагаются механизмы ранжирования элементов сети по степени их значимости, решается задача формирования очереди на администрирование. Произведен расчет приоритетов элементов сети предприятия. Ил.: 3. Табл.: 4. Библиогр.: 16 назв.

Ключевые слова: ранжирование, приоритет, функциональный элемент, ресурс администрирования.

Постановка проблемы. Основной интегративной платформой современных АСУП является корпоративная сеть передачи данных (КСПД) [1, 2] – распределенная инфраструктура, представляющая собой организованную совокупность компонентов – функциональных элементов (оконечных узлов, телекоммуникационного оборудования, протоколов и служб передачи данных) – ФЭ и каналов электросвязи [3, 4]. Под администрированием КСПД будем понимать целенаправленные управляющие воздействия на ФЭ, осуществляемые администраторами (человеко-машинными системами, реализующими определенные функции управления сетью) [5] в рамках обеспечения надежной, производительной и безопасной работы КСПД (целевой задачи администрирования) [6]. Оптимизация очереди на администрирование из множества всех проблемных (не соответствующих требованиям) ФЭ КСПД, в условиях ограниченного количества административных ресурсов (администраторов), является одной из актуальных задач в обеспечении требуемого качества функционирования АСУП. Подход к решению поставленной задачи авторы видят в ранжировании (системе приоритетов) ФЭ КСПД по степени их участия в обеспечении транспорта заданному (конечному) множеству информационных процессов (ИП) АСУП. Под приоритетом понимаем общий для всех ФЭ КСПД объективный показатель значимости [7].

Анализ существующих решений. Решение задач администрирования КСПД представляет сложную научную проблему, связанную с разработкой научно-обоснованных моделей и алгоритмов,

методов создания автоматизированных систем администрирования и адаптивного управления. Анализ актуальных исследований в рассматриваемой области [8, 9, 10], существующих на рынке специализированных программных средств администрирования [11, 12, 13] позволяет констатировать отсутствие эффективных механизмов ранжирования ФЭ КСПД по степени их участия в ИП, что затрудняет решение поставленной задачи. Применяемые экспертные оценки частично могут быть использованы, но, в условиях непрерывной модернизации КСПД и изменений прикладных задач АСУП, становятся малоэффективными вследствие их низкой динамичности. Недостаток теоретических и практических разработок в рассматриваемой области определяет актуальность создания научно-обоснованной методики количественного расчета приоритетов ФЭ КСПД.

Цель работы. В работе вводится понятие приоритета ФЭ КСПД, предлагается модель приоритетов и метод расчета, основанный на степени участия ФЭ в реализации заданного множества ИП АСУП.

Математическая модель. Обозначим множество ФЭ КСПД как $S_{\text{кспд}} = \{s_1, s_2, \dots, s_j\}$, $s_r \in S_{\text{кспд}}$ – ФЭ. Приоритетом (R) будем называть показатель значимости ФЭ для реализации ИП АСУП. Обозначим как $R(s_r)$ количественное значение приоритета s_r . Обозначим множество всех ИП КСПД через $P_{\text{кспд}} = \{p_1, \dots, p_m\}$, $p_i \in P_{\text{кспд}}$ – ИП. Информационный процесс p_i будем трактовать как информационное взаимодействие двух и более субъектов (пользователей структурных подразделений корпорации), целью которого является изменение имеющейся хотя бы у одного из них информации, реализация ИП связана с использованием ФЭ. В общем случае ИП представим четверкой:

$$p_i = \langle H_i, A_i, B_i, W_i \rangle, \quad (1)$$

где H_i – количественная оценка ранга p_i , определяемая экспертной группой в соответствии со степенью важности и срочности; $A_i = \{a^i_1, \dots, a^i_k\}$ – множество ФЭ-отправителей p_i , $A_i \subset S_{\text{кспд}}$; $B_i = \{b^i_1, \dots, b^i_e\}$ – множество ФЭ-получателей p_i , $B_i \subset S_{\text{кспд}}$; $W_i = \{w^i_1, \dots, w^i_p\}$, – множество элементарных ориентированных путей (все потенциальные пути прохождения сетевого трафика между абонентами, т.е. множество альтернативных способов реализации p_i), где $w^i_q \in W_i$ – путь от $a^i_j \in A_i$ к $b^i_j \in B_i$, $w^i_q \subset S_{\text{кспд}}$.

Численное значение приоритета ФЭ для ИП, согласно предлагаемой модели, пропорционально коэффициенту "участия" (γ) ФЭ в реализации ИП и его (процесса) рангу (H):

$$R_i(s_r) = \gamma_i(s_r) \times H_i, \quad (2)$$

где $R_i(s_r)$ – численное значение приоритета элемента s_r для процесса p_i , $\gamma_i(s_r)$ – коэффициент будем трактовать как "относительное участие" элемента s_r в реализации p_i .

Для расчета $\gamma_i(s_r)$ представим множество $P_{\text{кспд}}$ как связный неориентированный граф $N_{\text{кспд}}(S, L)$, где множество вершин графа – элементы $S_{\text{кспд}}$, множество дуг графа – каналы электросвязи.

Каждый ИП КСПД, как упорядоченное множество ФЭ, будет представлять собой ориентированный подграф искомого графа $N_{\text{кспд}}$, N_i – орграф p_i (рис. 1).

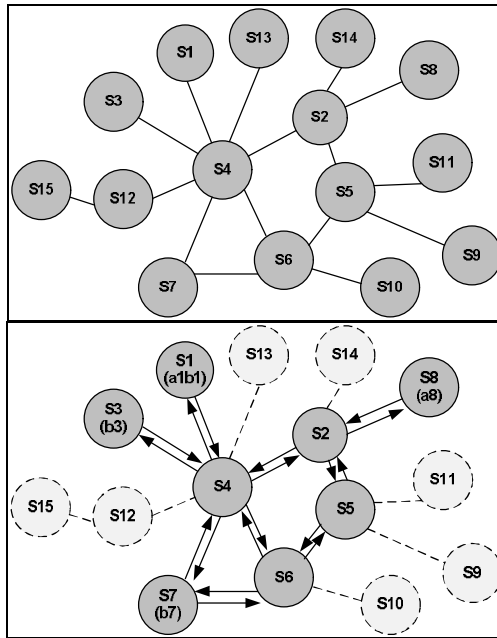


Рис. 1. Неограф $N_{\text{кспд}}$ (вверху), орграф N_i (внизу)

Множество всех альтернативных сочетаний между ФЭ-отправителями и ФЭ-получателями в рамках процесса p_i назовем множеством упорядоченных пар $\{a_j^i, b_j^i\}$ и обозначим как M_i , $a_j^i \in A_i$, $b_j^i \in B_i$. Найдем множество M_i как декартово произведение [14] множеств абонентов (A_i, B_i) по формуле:

$$M_i = A_i \times B_i. \quad (3)$$

Обозначим как $|M_i|$ количество пар $\{a_j^i, b_j^i\}$ (мощность множества M_i). Применяя математический аппарат теории графов [15] для $|M_i|$ пар "отправитель-получатель" процесса p_i найдем все пути взаимодействия $W_i = \{w_1^i, \dots, w_p^i\}$ в виде последовательностей ФЭ – узлов графа N_i , включая соответствующие ФЭ-отправители и ФЭ-получатели. Через $|W_i|$ – обозначим количество найденных путей. Множество путей, проходящих через s_r , обозначим как $W_i^r, W_i^r \subset W_i$. Обозначим количество таких путей $|W_i^r|$.

Параметр $\gamma_i(s_r)$ будем определять числом появлений $|W_i^r|$ на множестве путей W_i по формуле:

$$\gamma_i(s_r) = \frac{|W_i^r|}{|W_i|}. \quad (4)$$

Подставим результат (4) в формулу (2), рассчитаем численное значение $R_i(s_r)$:

$$R_i(s_r) = \frac{|W_i^r|}{|W_i|} \times H_i. \quad (5)$$

Выполнив расчет приоритетов ФЭ множества $S_{\text{кспд}}$ по всему множеству $P_{\text{кспд}}$, получим матрицу приоритетов ФЭ (Табл. 1).

Таблица 1. Матрица проритетов

$\begin{matrix} P_{\text{кспд}} \\ S_{\text{кспд}} \end{matrix}$	p_1	p_2	p_3	...	...	...	...	p_m
s_1	$R_1(s_1)$	$R_2(s_1)$	$R_3(s_1)$	...	...	...	...	$R_m(s_1)$
...	...	...	...	...	...	...	...	...
s_j	$R_1(s_j)$	$R_2(s_j)$	$R_3(s_j)$	...	...	...	...	$R_m(s_j)$

В соответствии с табл. 1 мы приходим к следующему. В графе $N_{\text{кспд}}$ каждому элементу s_r сопоставим m -вектор $\{R_1(s_r), R_2(s_r), \dots, R_m(s_r)\}$ приоритетов, где m – количество ИП АСУП. С использованием евклидовой метрики [16] в пространстве R^m состояния загрузки можно определить с применением весовой нормы:

$$R(s_r) = \sqrt{R_1^2(s_r) + \dots + R_m^2(s_r)} = \sqrt{\sum_{v=1}^m R_v^2(s_r)}, \quad (6)$$

где $R(s_r)$ – значение приоритета для элемента s_r , 0 – нулевой вес $R(s_r) = 0$ соответствует элементу, не принадлежащему множеству $P_{\text{кспд}}$.

Практическая часть. Рассмотрим предлагаемую технику расстановки приоритетов на примере КСПД предприятия ООО "Западно-

Малобалыкское" (ХМАО, Нефтеюганский район), схема рассматриваемой КСПД и ее граф представлены на рисунке (рис. 2).

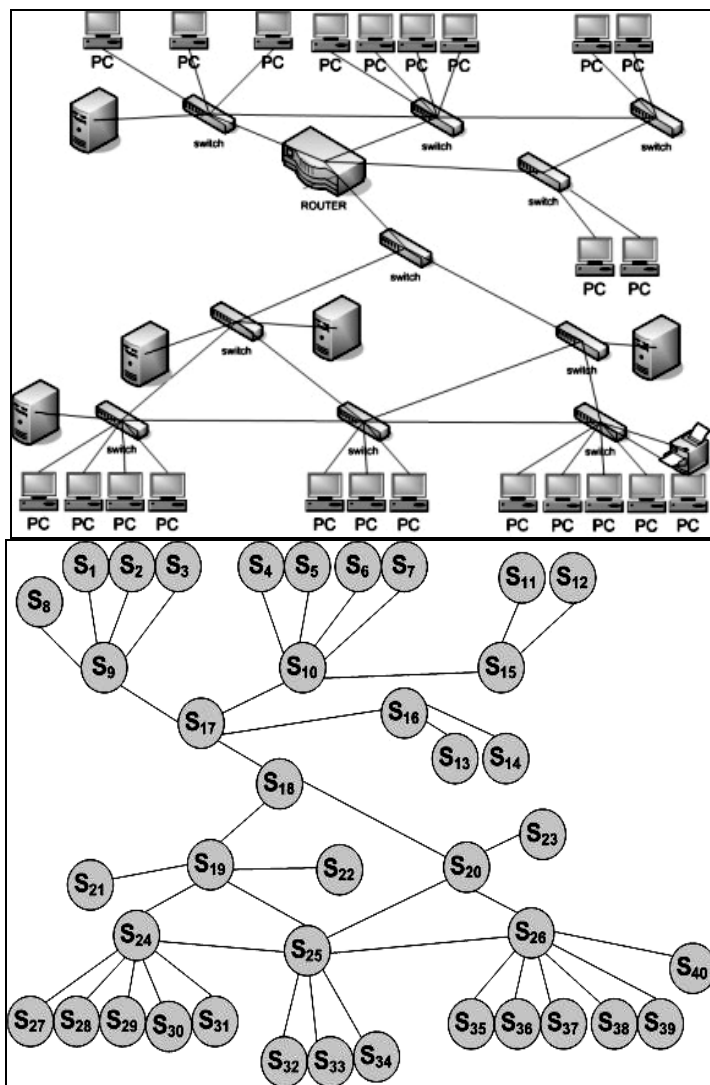


Рис. 2. Схема сети "Западно-Малобалыкское" (вверху) и ее неограф

В качестве исходных данных рассматриваются три основных информационных процесса КСПД подразделения ООО "Западно-Малобалыкское": $H_1(p_1) = 10$, $A_1 = \{s_{32}, s_{33}, s_{34}, s_{35}\}$, $B_1 = \{s_{11}, s_{12}\}$; $H_2(p_2) = 12$, $A_2 = \{s_{27}, s_{32}, s_{35}\}$, $B_2 = \{s_1, s_2\}$; $H_3(p_3) = 3$, $A_3 = \{s_{27}, s_{28}\}$, $B_3 = \{s_{35}, s_{36}\}$.

Выполним расчет количества пар отправитель-получатель для каждого из процессов по формуле (3): $|M_1| = 8$; $|M_2| = 6$; $|M_3| = 4$.

Найдем количество всех путей, проходящих от отправителя к получателю, для каждой из пар: $|W_1| = 108$; $|W_2| = 96$; $|W_3| = 28$.

Найдем количество путей, проходящих через каждый ФЭ (табл. 2):

Таблица 2. Пути, проходящие через каждый ФЭ

$i \backslash r$	1	2	9	10	11	12	15	16	17	18	19	20	24	25	26	27	28	32	33	34	35	36
1	0	0	36	72	54	54	108	36	108	108	60	60	30	102	54	0	0	24	24	24	36	0
2	48	48	96	64	0	0	32	32	96	96	60	60	54	84	54	36	0	36	0	0	36	0
3	0	0	0	0	0	0	0	0	0	12	20	20	28	24	28	14	14	0	0	0	14	14

Рассчитаем для ФЭ коэффициент "участия" по каждому ИП (формула 4) (табл. 3):

Таблица 3. Коэффициент участия по каждому ИП

$i \backslash r$	1	2	9	10	11	12	15	16	17	18	19	20	24	25	26	27	28	32	33	34	35	36
1	0	0	0,3	0,6	0,5	0,5	1	0,3	1	1	0,5	0,5	0,2	0,9	0,5	0	0	0,2	0,2	0,2	0,3	0
2	0,5	0,5	1	0,6	0	0	0,3	0,3	1	1	0,6	0,6	0,5	0,8	0,5	0,3	0	0,3	0	0	0,3	0
3	0	0	0	0	0	0	0	0	0	0,4	0,7	0,7	1	0,8	1	0,5	0,5	0	0	0	0,5	0,5

Выполнив расчет приоритетов ФЭ по всему множеству ИП (формула 5), получим матрицу приоритетов ФЭ (табл. 4):

Таблица 4. Матрица приоритетов ФЭ

$i \backslash r$	1	2	9	10	11	12	15	16	17	18	19	20	24	25	26	27	28	32	33	34	35	36
1	0	0	3,3	6,7	5	5	10	3,3	10	10	5,6	5,6	2,8	9,4	5	0	0	2,2	2,2	2,2	3,3	0
2	6	6	12	8,0	0	0	3,9	3,9	12	12	7,5	7,5	6,7	10,6	6,7	4,5	0	4,5	0	0	4,5	0
3	0	0	0	0	0	0	0	0	0	1,3	2,1	2,13	3	2,58	3	1,5	1,5	0	0	0	1,5	1,5

Найдем итоговое значение приоритета для каждого ФЭ по всему множеству ИП (формула 6) и получим взвешенный граф КСПД ООО "Западно-Малобалыкское", в котором каждому ФЭ сопоставлено количественное значение его приоритета (рис. 3) .

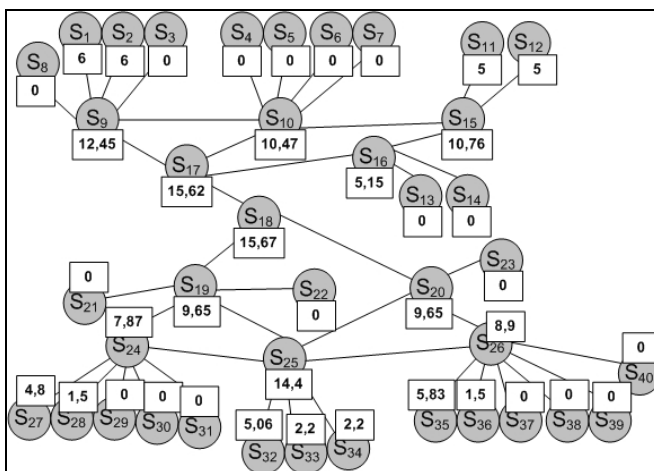


Рис. 3. Взвешенный граф КСПД ООО "Западно-Малобалыкское"

Выводы. Предложенный метод расчета приоритетов позволяет ранжировать множество всех ФЭ КСПД (в каждый момент времени), на основании приоритета как общего для всех ФЭ КСПД критерия. Таким образом, при возникновении множественных неисправностей в корпоративной сети на основе предлагаемого подхода может быть сформирована очередь на обслуживание (администрирование) ФЭ КСПД. Кроме того, приведенное ранжирование ФЭ может помочь выявить "узкие" места в КСПД и принять соответствующие меры по повышению надежности и живучести корпоративной сети.

Список литературы: 1. Каток А.Б. Введение в современную теорию динамических систем / А.Б. Каток, Б. Хасселблат. – М.: Факториал, 1999. – 768 с. 2. Гайфуллин Б.Н. Автоматизированные системы управления предприятиями стандарта ERP/MRP/II / Б.Н. Гайфуллин, И.А. Обухов. – М.: Богородский печатник, 2000. – 237 с. 3. Кульгин М. Технологии корпоративных сетей / М. Кульгин. – Энциклопедия. – СПб.: Питер, 1999. – 704 с. 4. Олифер В.Г. Стратегическое планирование сетей масштаба предприятия / В.Г. Олифер, Н.А. Олифер. – М.: Центр Информационных Технологий, 2000. – 680 с. 5. Мишин Д.В. Модель администратора корпоративной сети передачи данных / Д.В. Мишин, М.М. Монахова / Региональная информатика (РИ-2010). – XII Санкт-Петербургская межд. конф. "Региональная информатика (РИ-2010)" / Санкт-Петербург: Труды конф. / СПОИСУ. – СПб, 2010. – С. 55–56. 6. Мишин Д.В. Модель автоматизированной системы администрирования корпоративной сети передачи данных / Д.В. Мишин, М.М. Монахова / Труды Девятого международного симпозиума "Интеллектуальные системы" (Intels2010). – Россия, ВлГУ, 2010. – С. 268–271. 7. Мишин Д.В. Проблемы оптимизации распределения работ администраторов как основных исполнительных субъектов в рамках решения целевой задачи администрирования КСПД / Д.В. Мишин, М.М. Монахова / Материалы III Межд. научно-практ. конф. – Шуя-Иваново-Владимир: Изд-во ГОУ ВПО "ШГПУ". – С. 165–170.

8. *Леохин Ю.Л.* Многоуровневый подход к управлению мультисервисными корпоративными сетями // Телематика'2009. Труды XVI Всероссийской научно-методической конф. – Санкт-Петербург. 2009. – Том 2. – С. 277–278. 9. *Chen G.* Inte-grated TMN Service Management / *G. Chen, Q. Kong, J. Etheridge, P. Foster* // Journal of Network and system Management. – 1999. – Vol. 7. – № 4. – P. 469–485. 10. *Гребешков А.Ю.* Управление сетями электросвязи по стандарту TMN: Учеб. пособие / *А.Ю. Гребешков.* – М.: Радио и связь, 2004. – 155 с. 11. *Куриц А.Л.* Принципы построения средств управления ИТ-инфраструктурой на примере модели ITSM компании HP / *А.Л. Куриц, А.Л. Фридман, Б.Н. Андерс, Н.А. Фандюшина, Л.Я. Чумаков* // Системы и средства информатики. – 2008. – Т. 18. – № 2. – С. 69-85. 12. IBM Tivoli Software library for technical resources [Electronic resource] / IBM Corp. – 2006. Режим доступа: <http://www-306.ibm.com/software/tivoli/sw-library> 13. *Tammy Zitello.* HP OpenView System Administration Handbook / *Zitello Tammy, Weber Paul, Williams Deborah* / Network Node Manager, Customer Views, Service Information Portal, HP OpenView Operations / Pearson Education, Inc., Upper Saddle River, New Jersey, 2004. – 688 p. 14. *Ершов Ю.Л.* Математическая логика / *Ю.Л. Ершов, Е.А. Палютин.* – СПб.: Лань, 2005. – 320 с. 15. *Салий В.Н.* Алгебраические основы теории дискретных систем / *В.Н. Салий, А.М. Богомолов.* – М.: Физ.-мат. лит., 1997. – 368 с. 16. *Herbert Busemann.* Projective Geometry and Projective Metrics (Dover Books on Mathematics) / *Busemann Herbert, Kelly Paul J.* // New York, Academic Press inc., pub., 2005. – 352 p.

Статья представлена д.т.н. проф. Владимирского государственного университета им. А.Г. и Н.Г. Столетовых, Монаховым М.Ю.

УДК 004.021

Про оптимізації адміністрування корпоративних мереж передачі даних в умовах обмежених адміністративних ресурсів АСУП // Мішин Д.В., Монахова М.М. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 101 – 108.

У роботі пропонується модель пріоритетів функціональних елементів корпоративних мереж передачі даних, пропонуються механізми ранжирування елементів мережі за ступенем їх значущості, вирішується завдання формування черги на адміністрування. Зроблено розрахунок пріоритетів елементів мережі підприємства. Іл.: 3. Табл.: 4. Бібліогр.: 16 назв.

Ключові слова: ранжування, пріоритет, функціональний елемент, ресурс адміністрування.

UDC 004.021

About the optimization of the administration corporate area networks of data transmission under scarce administrative resources // *Mishin D.V., Monakhova M.M.* // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 101 – 108.

The aim of this paper is to development the model of the estimate of priorities functional elements for the corporate area networks of data transmission. In this paper we propose effective mechanisms for ranking (assigning priorities) the functional elements of the network by the degree of their importance. Also we solve the question of creation the queue from the set of the all problem elements. Figs.: 3. Tabl.: 4. Refs.: 16 titles.

Keywords: priority, ranking the functional elements, administrative resources.

Поступила в редакцию 15.02.2011

Б.І. МОРОЗ, д.т.н., проф., зав. каф. АМСУ, Дніпропетровськ,
С.М. КОНОВАЛЕНКО, асп., нач. лаб. АМСУ, Дніпропетровськ

ДЕЯКІ АСПЕКТИ РОЗВИТКУ СИСТЕМИ АНАЛІЗУ РИЗИКІВ ПОРУШЕННЯ МИТНОГО ЗАКОНОДАВСТВА

Розглянуто деякі аспекти розвитку автоматизованої системи аналізу ризиків порушення митного законодавства, а саме можливість застосування методів та засобів штучного інтелекту у вигляді нейромережевого моделювання. Описана модель нейронної мережі типу багатошаровий перцептрон та використано ітеративний метод навчання з можливістю обходу локальних мінімумів. Іл.: 1. Табл.: 1. Бібліогр.: 11 назв.

Ключові слова: аналіз ризиків, нейронна мережа, перцептрон, ітеративний метод.

Постановка проблеми. Одним з основних завдань у діяльності митних органів є здійснення митного контролю. До останнього часу у більшості випадків домінував метод безпосереднього митного контролю, його форми та глибину обирав інспектор [1]. Очевидно, що такий підхід збільшує часові обмеження пропуску та оформлення товарів, транспортних засобів. Окрім того, непоодинокі випадки подання суб'єктами зовнішньоекономічної діяльності недостовірних відомостей про характеристики товарів з метою заниження (завищення) митної вартості, перевезення контрабанди. Концепція створення, впровадження та розвитку автоматизованої системи аналізу та керування ризиками передбачає використання сучасних інформаційних технологій та методів і засобів обробки інформації митного контролю [2, 3].

Аналіз літератури. Публікації в галузі впровадження інформаційних систем аналітики в рамках розвитку проекту *E-customs* носять концептуальний характер [1 – 3], виділяючи проблематику та формалізуючи загальні аспекти функціонування. Як видно із [4 – 6], великий інтерес представляє використання нейромережевих моделей в системах ідентифікації та розпізнавання, проте застосування цих парадигм штучного інтелекту в митній службі України досі предметно не розглядалась, тому розробка та застосування методів та засобів ідентифікації, яким були б притаманні такі риси як здатність до адаптації, навчання є задачею остаточно не вирішеною та потребує подальшого розвитку та дослідження.

Ціль статті. Метою статті є спроба формалізації та розробки моделі ідентифікації об'єкта, явища, що супроводжує митний контроль на основі математичного апарату нейронних мереж. Оскільки в деяких випадках вхідні данні можуть мати невизначений характер та з часом змінюватися,

то моделі ідентифікації необхідно надати такі властивості інтелектуальних систем як адаптивність, можливість оперувати розмитими даними, прийняття рішень в умовах невизначеності.

Митний контроль, як процес прийняття рішень. Під прийняттям рішення будемо розуміти складний процес в якому можна виділити наступні етапи:

1) Аналіз предметної галузі, проблеми що розглядається, тобто виділення факторів та ознак, що мають найбільш важливе значення.

2) Побудова математичної або алгоритмічної моделі процесу ідентифікації, щоб встановити або наблизити співвідношення між вхідним та вихідним вектором. Цей етап включає в себе опису та можливості побудови цільової функції змінних такої, значенням якої відповідала би найкраща ситуація з погляду прийняття рішення.

Оскільки класи об'єктів митного контролю можуть мати досить великий і різноманітний обсяг ознак, то для практичних цілей виділимо деякі з них, які згодом і використаємо:

- товари, об'єми ввезення яких за даними митної статистики України значно перевищують об'єми їх вивезення за даними митних статистик країн-контрагентів;

- товари, щодо яких є інформація про відсутність виробництва або товари, виробництво яких є нехарактерним для певної країни;

- товари, заявлені в одній декларації, але доставлені в декількох автотранспортних засобах, вагонах або контейнерах;

- пред'явлення товарів для митного оформлення в митний орган, відмінний від митного органу призначення, зазначеного в документі контролю за доставкою товарів;

- різниця між брутто і нетто вагою товарів, що перевозяться, відмінна від загальноприйнятої;

- вага одиниці товару не є характерною для даного товару або ідентичних чи подібних (аналогічних) товарів;

- заявлена митна вартість товарів значно відрізняється від ціни ідентичних чи подібних (аналогічних) товарів при їх увезенні на митну територію України;

- товари переміщуються через митний кордон України за зовнішньоекономічними договорами (контрактами), відмінними від договорів купівлі-продажу;

- одна зі сторін зовнішньоекономічного договору зареєстрована в офшорній зоні;

- "білі" та "чорні" списки учасників зовнішньоекономічної діяльності.

Представимо ідентифікаційні ознаки, як вхідну множину $x_i \in X$, а вихідний результат процедури розпізнавання, в контексті системи аналізу ризиків, опишемо як множину зі ступенями ризику $Y = \{\text{високий, помірний, низький}\}$. Як бачимо вхідні та вихідні вектори мають різношкальні характеристики, включаючи навіть лінгвістичні значення, тому відповідні вектори потребуватимуть додаткового нормування та кодування.

Система класифікації на основі нейронних мереж. За останні 15 років дуже корисними з погляду розпізнавання образів виявились моделі нейронних мереж, чисельність комерційних програм що їх використовують постійно збільшується. Одними з найкорисніших переваг з погляду класифікації та розпізнавання образів є здатність до навчання та узагальнення накопичених знань [6].

Для побудови нейромережевої моделі необхідно:

- визначити архітектуру нейронної мережі;
- визначити процедуру навчання мережі;
- визначити критерій якості та шлях його оптимізації.

З відомих архітектур нейронних мереж універсальним апроксиматором є багатошаровий перцептрон прямого поширення сигналу та зворотного розповсюдження помилки [7 – 10]. Як зазначалось вище ми виділили 10 вхідних змінних та одну вихідну з трьома станами. Для того щоб навчити багатошаровий перцептрон, треба привести вхідний та асоційований з ним вихідний вектори у відповідність до того вигляду, з яким може працювати наша нейронна мережа.

Вихідний вектор має три категоричні значення {високий, помірний, низький}, тож враховуючі властивості передаточної функції логістичного сигмоїда, логічно було представити його в межах від 0 до 1, а саме відповідними значеннями {1; 0.5; 0}. Для зручності вхідні значення x_i приводимо до того ж діапазону що і вихідний вектор, наприклад: ознака x_6 – вага одиниці товару, має типовий статичний діапазон значень, відхилення від якого і буде вказувати на ступінь ризику. Таким чином згруповуючи діапазони значень або категорій в три групи кодуємо їх відповідно до ступеня ризику – {1; 0.5; 0}.

Отже в нас є множина вхідної інформації X – ознак об'єкта митного контролю, вихідної інформації, що фактично видає мережа на вихідному шарі Y – рекомендації системи щодо можливого ризику порушення, і, власне, множина бажаного виходу тобто правильно розпізнаного образу D [8, 9]. Із комбінації множин X та D сформуємо навчальну вибірку $\langle x, d \rangle$, де $x = [x_0, x_1, \dots, x_n]^T$, а $d = [d_1, d_2, \dots, d_m]^T$. Для прикладу використаємо

таку товарну позицію, як “монітори кольорові” (номер 8528599000 згідно товарної номенклатури). Було сформовано вибірку із 275 одиниць, із них 193 – навчальна вибірка, 41 – контрольна вибірка та 41 – тестова вибірка.

Проектуючи нейронну мережу потрібно було вирішити скільки буде шарів та кількість нейронів в них. За основу було взято тришаровий персептрон із двома прихованими шарами, де кількість нейронів 1-го шару дорівнює кількості ознак вхідного вектору – 10 нейронів, 2-го шару – 8, і нарешті вихідний шар складається з одного нейрона. Модель, що була розроблена в середовищі *MATLAB 7.11.0 (Neural network toolbox)*, зображена на рисунку.

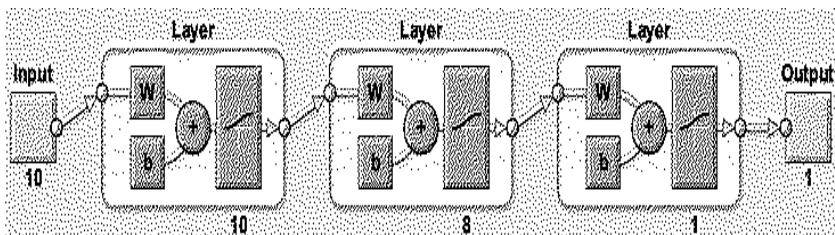


Рис. Схема тришарового персептрона.

Вхідний сигнал S_i^k i -го нейрона, що знаходиться в k -му шарі, обчислюється наступним чином:

$$S_i^k = \sum_{j=0}^t x_j^{k-1} w_{ij}, \quad (1)$$

де x_j^{k-1} – вихідний сигнал нейрона j ($j = \overline{1, t}$) у шарі номер $k-1$; x_0^{k-1} – сигнал зміщення i -го нейрона; w_{ij} – ваговий коефіцієнт зв'язку між нейронами i та j ; t – кількість нейронів в шарі $k-1$. Кожен i -й нейрон в свою чергу формує, за допомогою функції активації $f(S_i^k)$, вихідний сигнал y_i^k :

$$y_i^k = f\left(\sum_{j=0}^t x_j^{k-1} w_{ij}\right). \quad (2)$$

В якості функції активації, як вже зазначалось вище, використаємо сигмоїдну, функцію. Запишемо формулу (2) з урахуванням формули (1) в наступному вигляді:

$$y_i^k = \frac{1}{1 + e^{-S_i^k}}. \quad (3)$$

Для того щоб наша нейромережева модель набула здатності до класифікації та узагальнення потрібно провести процедуру навчання використовуючи вибірку $\langle x, d \rangle$. Для навчання використаємо традиційний метод зворотного розповсюдження помилки (*back propagation*) [10–11].

Було проведено низку експериментів використовуючи різні функції мінімізації з критерієм середньої квадратичної похибки *net.trainParam.goal*=0.01 (таблиця), де було встановлено, що найкраще впорався з поставленою задачею метод Левенберга-Маркардта [6] та метод шкальних зв'язаних градієнтів.

Таблиця. Порівняння методів навчання розробленої нейронної мережі

Функція мінімізації цільової функції	Середня квадратична похибка (0,01)	Коефіцієнт кореляції	Кількість епох (max 1500)
Метод Левенберга-Маркардта	0,0081	0,9664	10
Метод градієнтного спуска	0,6749	0,3682	1500
Метод градієнтного спуска з урахуванням моментів	0,0558	0,5206	1500
Метод шкальних зв'язаних градієнтів	0,0099	0,9576	50

Необхідно зауважити що використаний багатошаровий персептрон (рис.) отриманий емпіричним шляхом. Міняючи кількість прихованих шарів, нейронів в шарах було виявлено лише негативну динаміку у зменшенні коефіцієнта кореляції та неможливості досягти значення похибки хоча б СКП = 0,6, принаймні на тій вибірці, яку було сформовано.

Висновки. В результаті роботи зроблені перші кроки по опису застосування методів штучного інтелекту для потреб митної служби України з подальшою перспективою впровадження.

В подальшому перспективним для дослідження є:

1. Приділення більшої уваги якості вихідних даних, їх кодуванню та нормалізації.

2. Застосування методів визначення структури нейронної мережі.
3. Використання ефективних методів глобальної оптимізації.

Список літератури: 1. Кунев Ю.Д. Управління в митній службі: підручник / Ю.Д. Кунев, І.М. Коросташова, А.В. Мазур, С.П. Шапошник / За ред. Ю.Д. Кунева. – К.: Центр навчальної літератури, 2006. – 408 с. 2. Основи митної справи в Україні: підручник / За ред. П.В. Пашика. – К.: Знання, 2008. – 652 с. 3. Регулювання митної справи: підручник / За ред. А.Д. Войцешука. – Хмельницький: Інтрада, 2007. – 312 с. 4. Ripley B.D. Neural Networks and Related Methods for Classification / B.D. Ripley // Journal of the Royal Statistical Society, Series B, Methodological. – 1994. P. 409-456. 5. Hagan M.T. Neural Networks for Control / M.T. Hagan, H.B. Demuth // Proceedings of the 1999 American Control Conference, San Diego, CA, 1999. – P. 1642-1656. 6. Осовский С. Нейронные сети для обработки информации: Пер. с польского И.Д. Рудинского / С. Осовский. – М.: Финансы и статистика, 2002. – 344 с. 7. Саймон Хайкин Нейронные сети: полный курс, 2-е издание: Пер. с англ. / Саймон Хайкин. – М.: Издательский дом "Вильямс", 2006. – 1104 с. 8. Каллан Роберт Основные концепции нейронных сетей / Роберт Каллан. – М.: Издательский дом "Вильямс", 2001. – 287 с. 9. Лю Б. Теория и практика неопределенного программирования / Б. Лю. – М.: БИНОМ. Лаборатория знаний, 2005. – 416 с. 10. Борисов Е.С. Классификатор на основе многослойной нейронной сети. [Електронний ресурс] – Режим доступу: <http://www.mechanoid.kiev.ua>. 11. Swingler K. Applying Neural Networks. A Practical Guide / K. Swingler. – Academic Press, 1996.

УДК 004.93'1:656.073.5

Некоторые аспекты развития системы анализа рисков нарушения таможенного законодательства / Мороз Б.И., Коноваленко С.Н. // Вестник НТУ "ХПИ". Тематический выпуск: Информатика и моделирование. — Харьков: НТУ "ХПИ". — 2011. — № 17. — С. 109 – 114.

Рассмотрено основные аспекты развития автоматизированной системы анализа рисков нарушения таможенного законодательства, а именно возможность применения методов и средств искусственного интеллекта в виде нейросетевого моделирования. Описана модель нейронной сети типа многослойного персептрона и использован итеративный метод обучения с возможностью обхода локальных минимумов. Ил.: 1. Табл.: 1. Библиогр.: 11 назв.

Ключевые слова: анализ рисков, нейронная сеть, персептрон, итеративный метод.

UDC 004.93'1:656.073.5

Some aspects of development system of the analysis risks infringement of the customs legislation / Moroz B.I., Konovalenko S.N. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – № 17. – P. 109 – 114.

It is considered the basic aspects of development of the automated system of the analysis of risks of infringement of the customs legislation, namely possibility of application of methods and artificial intelligence techniques in a kind neuro-network modelling. The model of a neural network of type multilayered perceptron is described and the iterative method of training with possibility of detour of local minima is used. Figs.: 1. Tabl.: 1. Refs.: 11 titles.

Key words: the analysis of risks, neural network, perceptron, an iterative method.

Поступила в редакцию 14.02.2011

О.Я. НИКОНОВ, д.т.н., проф., зав. каф. ХНАДУ, Харьков,

О.В. МНУШКА, ассистент ХНАДУ, Харьков,

В.Н. САВЧЕНКО, ст. преп. УИПА, Харьков

ОЦЕНКА ТОЧНОСТИ ВЫЧИСЛЕНИЙ СПЕЦИАЛЬНЫХ ФУНКЦИЙ ПРИ РАЗРАБОТКЕ КОМПЬЮТЕРНЫХ ПРОГРАММ МАТЕМАТИЧЕСКОГО МОДЕЛИРОВАНИЯ

Проведен анализ формата вещественных чисел IEEE-754 как составляющей дополнительной погрешности при вычислении специальных функций. Проведено сравнение и показано, что применение арифметики с произвольной точностью (*mpfr_erf()*, *MPFR*) и стандартной функции (*erf()*, *C99*) в программах на C/C++ снижает быстродействие на порядок при той же точности. Ил.: 2. Библиогр.: 8 назв.

Ключевые слова: формат IEEE-754, произвольная точность, точность вычислений специальных функций.

Постановка проблемы. Применение статистических методов в математическом моделировании сложных технических систем, например, спутниковых систем связи [1], позволяет значительно снизить временные и материальные затраты на их разработку, оценить воздействие различных влияющих факторов на поведение системы в реальных условиях эксплуатации. Многократное повторение вычислений в процессе моделирования может приводить к значительным погрешностям результатов, обусловленных ограниченным числом разрядов представления вещественных чисел в памяти компьютера. Специфика многих задач не позволяет пользоваться универсальными пакетами моделирования, требует реализации на одном из универсальных языков программирования и оценки точности вычислений, и адекватности полученных численных результатов.

Анализ литературы. Представление вещественных данных в памяти компьютера регламентируется стандартом IEEE754-2008 [2], в котором определены четыре формата двоичных чисел – *binary16*, *binary32*, *binary64*, *binary128*, и три формата десятичных чисел – *decimal32*, *decimal64*, *decimal128*, а также способы расширения базовых форматов для повышения точности вычислений. Точность представления чисел в определяемых стандартом форматах составляет соответственно 7, 16, 34 десятичных разрядов после запятой. Существенным отличием данной редакции стандарта от предыдущей [3] является отсутствие формата *extended* с размером данных 80 бит, который жестко привязан к размерности регистров сопроцессора x87. Способы реализации стандарта

для языка C++ предлагаются в [4], а сама поддержка стандарта в разной степени реализована в разных компиляторах языка C/C++. Таким образом, требуется оценка кода, получаемого при помощи различных компиляторов, для учета дополнительной погрешности результатов вычислений, обусловленной реализацией алгоритмов вычислений.

Цель статьи – оценка дополнительной погрешности результатов моделирования численными и статистическими методами, возникающей за счет приближенного представления данных в памяти компьютера и форматов вещественных данных языка программирования C/C++.

Основная часть. Основными преимуществами языка C/C++ для научных вычислений является его универсальность, кроссплатформенность и переносимость кода. При этом, если для вычислений с целыми типами данных существует более или менее неплохая совместимость между различными компиляторами и вычислительными платформами, то в случае с вещественными типами данных существует ряд нестыковок, которые могут приводить к неоднозначной трактовке результата и даже потере данных.

В статистических методах математического моделирования часто возникает задача вычисления функции ошибок, в частности, при построении информационных карт вероятности ошибки систем спутниковой мобильной связи. Функция ошибки определена как [4]

$$\operatorname{erf} x = \frac{2}{\sqrt{\pi}} \int_0^x e^{-t^2} dt \quad \text{или} \quad \operatorname{erfc} x = \frac{2}{\sqrt{\pi}} \int_x^{\infty} e^{-t^2} dt. \quad (1)$$

Для функций (1) не существует аналитического решения, поэтому для ее вычисления с высокой точностью используют различные приближения, от выбора которых зависит не только точность вычислений, но и быстродействие программы моделирования.

Наиболее часто функция $\operatorname{erf}()$ аппроксимируется рядами [5] или различными минимаксными полиномиальными приближениями, например, рациональной чебышевской аппроксимацией [6], для которой существует стандартная библиотечная функция на фортране [7]. Недостатком такого подхода является фиксированная точность получаемого результата, что не всегда оправдано с точки зрения физической формулировки задачи и дополнительных накладных вычислительных расходов. Разложение в ряд может быть использовано для построения алгоритмов вычисления с произвольной точностью результата. На таком принципе основаны широко распространенные библиотеки численных методов, такие как GMP (GNU библиотека

арифметики с произвольной точностью) и ее производные: MPFR (библиотека языка Си для вычислений с произвольной точностью и корректным округлением результата), MPC (библиотека для вычислений с комплексными числами с произвольной точностью) и т.д.

Снижение погрешности результатов моделирования статистическими методами требует большого числа испытаний, что может приводить к значительным затратам времени и накоплению ошибки, обусловленной недостаточной точностью промежуточных результатов. С другой стороны, необходимо удерживать большое количество значащих цифр, что требует оценки целевой вычислительной платформы на предмет способности хранить эти знаки.

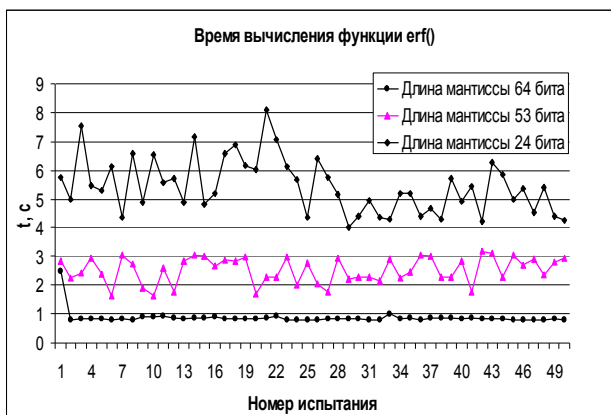
В исследовании использовались компиляторы: Microsoft Visual C++ из состава Microsoft Visual Studio 2008(2010); gcc 4.x (Debian 4.4.5-8) и 4.5.0 (mingw); Intel C Compiler 11.1.054 (Windows), 11.1.064 (Debian amd64), 11.1.074 (Debian x86). Авторами была составлена программа на языке C++ для сравнения эффективности работы различных алгоритмов вычисления интеграла ошибки. Программа позволяет провести указанные вычисления с использованием функции `mpfr_erf()` из пакета MPFR, `erf()` и `erfc()` из стандартной библиотеки C/C++, а также функций на основе аппроксимации, предложенной Cody [6, 7] и аппроксимации рядами по формуле [(7.1.26), 5].

На рис. 1, а показано время, затраченное функцией `mpfr_erf()`, которая вычисляла интеграл ошибки с точностью 32 (7-8 десятичных разрядов), 53 (15-16 десятичных разрядов) и 64 (19-20 десятичных разрядов) бита, что соответствует базовым числовым типам данных языка C/C++. Было проведено 50 испытаний, в каждом из которых вычислялось значение функции ошибки в цикле 100000 раз. Увеличение длины мантиссы от 23 до 64 разрядов увеличивает среднее время вычислений приблизительно в 5 раз (рис. 1, а).

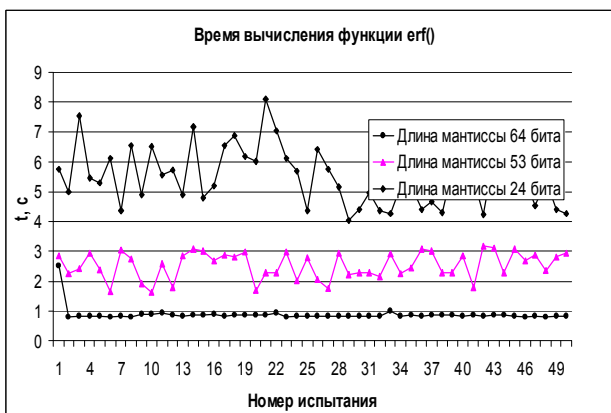
Проведено сравнение работы функции на основе аппроксимации, предложенной Cody [6, 7] и стандартной, определенной стандартом языка C99 (рис. 1, б). Для оценки времени выполнения количество повторений цикла было увеличено в 10 раз, время необходимое стандартному алгоритму оказалось пренебрежимо малым, функция на основе аппроксимации [6, 7] также оказалась значительно быстрее функции с произвольной точностью. Следует отметить, что при компиляции программы не производилась оптимизация под используемую вычислительную платформу.

Учитывая высокие скоростные показатели работы стандартных библиотечных функций, возникает необходимость оценки точности

представления вещественных чисел в памяти компьютера и их интерпретации компиляторами C/C++.



а)



б)

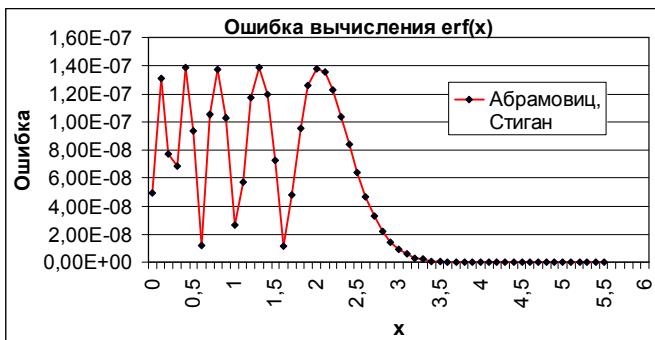
Рис.1. Время вычисления функции ошибок: а) с произвольным числом разрядов; б) на основе реализации алгоритма Cody и стандартного (C99)

Авторами проведен анализ поддержки форматов вещественных чисел наиболее распространенными компиляторами языка C/C++. В зависимости от реализации компилятора и операционной системы вещественные типы бывают: float (binary32), double (binary64), long

double и занимают 10 (extended [3]), 12 или 16 (binary128) байт в памяти. Анализ полученных данных показывает, что: для данных в формате long double компилятором выделяется различное количество байт для хранения одного экземпляра переменной в памяти; неоднозначная интерпретация компиляторами C/C++ этого формата требует аккуратности и обоснованности его применения; большее число байт должно обеспечивать большую точность представления чисел.

Для оценки точности вычислений в различных форматах данных была составлена программа вычисления $\pi = 4 \cdot \arctg(1.0)$. Анализ полученных данных позволяет сделать вывод о том, что: 1) переход от формата double к long double позволяет получить дополнительно 3 точных цифры результата; 2) переменные формата long double компиляторами `icc` и `gcc/g++` преобразуются в формат extended (по IEEE754-1985), а MS VC++ в формат double; результат может содержать "информационный мусор", который отображается на экране и может быть неверно истолкован; 3) заголовочный файл `cmath` для компилятора MS VC++ отличается от аналогичного для компилятора `gcc/g++` отсутствием ряда часто используемых функций, а также не определяет способов работы с форматами long double размером больше 8 байт; 4) только арифметические операции для переменных типа long double гарантированно выполняются; переменные типа long double (больше 8 байт) могут некорректно выводиться как стандартным оператором (`printf()`), так и потоковым (`cout`) в ОС Windows. Результат, который будет отображаться на экране, в данном случае непредсказуем; 5) работа с форматом long double в ОС Linux достаточно корректна. Определен набор функций для переменных этого формата, при этом ввод-вывод также работает корректно, единственным недостатком следует считать "информационный мусор", который отображается после 18 десятичного разряда (а для формата double – после 15).

Была проверена погрешность вычисления `erf()` по формуле [(7.1.26), 5] с аппроксимацией Cody. Было установлено, что первая дает погрешность по сравнению с аппроксимацией Cody $1e-7$ (рис. 2, а) в диапазоне $x = [0, \dots, 3]$, погрешность реализации аппроксимации Cody и стандартной C99 дает погрешность не более чем $1e-17$ (рис. 2, б) по сравнению с "точным" значением. Уменьшение погрешности вычислений (рис. 2, а) объясняется недостатком разрядов для хранения числа. Дальнейшее повышение точности вычислений возможно только с использованием специальных приемов [8].



а)



б)

Рис. 2. Погрешность вычислений функции ошибок:
а) для типа float; б) для типа long double

Выводы. 1. Достижимая погрешность вычислений, которая обеспечивается средствами стандартной библиотеки языка C++ не превышает 18 десятичных разрядов, для достижения более высокой точности следует применять библиотеки вычислений с произвольной точностью – MPFR, интервальную арифметику – MPFI и др., при этом значительно увеличиваются объем требуемой памяти и время вычислений (на порядок и более по сравнению с функциями стандартной библиотеки). Компиляторы gcc / g++ и icc поддерживают представление чисел в формате Binary128 (Decimal128), но в настоящий момент времени не обеспечивают средствами работы с ними. 2. Увеличение быстродействия библиотек типа MPFR должно базироваться на особенности формул, лежащих в их основе, которые должны поддаваться распараллеливанию и использованию SIMD – расширений процессора

Intel. При этом требуется как разработка соответствующих алгоритмов, так и поддержка со стороны компилятора языка программирования.

Перспективами дальнейших исследований является разработка более эффективных алгоритмов вычисления значения специальных функций, таких как $\text{erf}()$, с произвольной точностью и увеличенным быстродействием, на основе использования алгоритмов параллельных и многопоточных вычислений, с целью повышения точности и ускорения вычислений при моделировании сложных технических систем.

Список литературы: 1. Мазманишвили А.С. Визуализация информационных характеристик электромагнитной обстановки в системах спутниковой связи / А.С. Мазманишвили, О.Я. Никонов // Вісник СумДУ. Серія "Технічні науки". – 2008. – № 4. – С. 30-37. 2. IEEE Standard for Floating-Point Arithmetic. – New York, 2008. – 70 P. 3. IEEE Standard for Binary Floating-Point Arithmetic. – New York, 1985. – 23 P. 4. Information Technology – Extension for the programming language C++ to support decimal floating-point arithmetic [Электронный ресурс]. – Режим доступа: <http://www.open-std.org/jtc1/sc22/wg21/docs/papers/2009/n2849.pdf>. 5. Справочник по специальным функциям с формулами, графиками и математическими таблицами / [М. Абрамовиц и др.]. – М.: Наука, 1979. – 832 с. 6. Cody W.J. Rational Chebyshev approximations for the error function / W.J. Cody // Math. Comp. – 1969. – No.23. – P. 631-637. 7. SUBROUTINE CALERF(ARG,RESULT,JINT) [Электронный ресурс] / Режим доступа: <http://www.netlib.org/specfun/erf>. 8. Chevillard S. Computation of the error function erf in arbitrary precision with correct rounding / S. Chevillard, N. Revol. – Proc. 8th Conference on Real Numbers and Computers, July 2008. – P. 27-36.

УДК 519.651:004.222.3

Оцінка точності обчислень спеціальних функцій при розробці комп'ютерних програм математичного моделювання / Никонов О.Я., Мнушка О.В., Савченко В.М. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 115 – 121.

Проведений аналіз формату дійсних чисел IEEE-754 як складової додаткової похибки при обчисленні спеціальних функцій. Проведено порівняння й показано, що застосування арифметики з довільною точністю ($\text{mpfr_erf}()$, MPFR) і стандартної функції ($\text{erf}()$, C99) у програмах на C/C++ знижує швидкодію на порядок при тій же точності. Іл.: 2. Бібліогр.: 8 назв.

Ключові слова: формат IEEE-754, довільна точність, точність обчислень спеціальних функцій.

UDC 519.651:004.222.3

Estimation of accuracy of special functions computation in the development of mathematical simulation software / Nikonov O.Ya., Mnushka O., Savchenko V.N. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 115 – 121.

Format reals in IEEE-754 as a component of the additional error in the calculation of special functions have been analysed. A comparison is shown that the use of arithmetic with arbitrary precision ($\text{mpfr_erf}()$, MPFR) and the standard function ($\text{erf}()$, C99) in programs for C / C++ reduces the speed of the order with the same accuracy. Figs.: 2. Refs.: 8 titles.

Keywords: IEEE-754 format, an arbitrary precision, accuracy of special functions computation.

Поступила в редакцію 04.02.2011

Н.В. ПЛЮТА, асп. кафедри математичного моделювання ДВНЗ
"Запорізький національний університет", Запоріжжя,
С.І. ГОМЕНЮК, д.т.н., проф., декан математичного факультету
ДВНЗ "Запорізький національний університет", Запоріжжя

МОДЕЛЬ КООРДИНАЦІЙНОЇ ВЗАЄМОДІЇ В СКЛАДНІЙ ІЄРАРХІЧНО ВПОРЯДКОВАНІЙ СИСТЕМІ

Обґрунтовано актуальність проблеми розробки моделей та методів, що дозволять комплексно вирішити задачу координації. Побудовано модель координаційної взаємодії, що відображає зв'язок між системою управління та процесами в складній ієрархічно впорядкованій системі. Запропоновано постановку задачі координації для систем даного класу. Бібліогр.: 10 назв.

Ключові слова: задача координації, ієрархічно впорядкована система, модель координаційної взаємодії.

Постановка проблеми. Поява та стрімкий розвиток розподілених систем різної природи, зокрема, віртуальних підприємств, розподілених інформаційних систем і т.п. викликає необхідність підтримки вирішення задачі координації сучасними програмними засобами. Проте математичний апарат теорії координації, як базис такого програмного забезпечення, не в повній мірі відповідає потребам сучасних ієрархічно впорядкованих систем через те, що не існує єдиної методології та відповідних моделей та методів для комплексного вирішення на їх основі задачі координації. Тому актуальною є наукова проблема розробки математичного апарату теорії координації, що дозволить не лише обрати "оптимальний координуючий сигнал" в рамках сталої структури системи управління, але й визначити шляхи її вдосконалення, тобто вибору "оптимальної схеми взаємозв'язку" між центрами прийняття рішень для знаходження найбільш прийнятної, з точки зору глобальної мети, рішення для кожної задачі, що вирішується в рамках системи асинхронних послідовно-паралельних процесів.

Аналіз літератури. Досить широке коло наукових праць вітчизняних та зарубіжних вчених присвячене проблемам математичної теорії координації [1 – 10], тому доцільним є виділення двох основних напрямків досліджень: 1. Побудова координаційного механізму у відповідності до конкретної задачі. Базова концепція даного підходу закладена Дж. Данцигом та П. Вульфом і набула подальшого розвитку в роботах [5 – 10]. Методи даного підходу дозволяють визначити "оптимальну схему взаємозв'язку" при розв'язанні конкретної задачі, проте не враховують багатозадачність ієрархічно впорядкованих систем

та складність структури асинхронної системи послідовно-паралельних процесів. 2. Побудова координаційного механізму у відповідності до конкретної системи управління. Методологічне підґрунтя даного підходу закладене в роботі [4] та активно розвивається по наш час, зокрема, в роботах [1 – 3]. Методи даного підходу дозволяють визначити "оптимальний координуючий сигнал", проте не пропонують засоби вдосконалення структури системи управління.

Мета статті – побудова моделі координаційної взаємодії в ієрархічно впорядкованій системі та постановка задачі координації в рамках процесного підходу до побудови координаційного механізму.

Основний розділ. Для комплексного вирішення задачі координації доцільним є застосування процесного підходу до побудови координаційного механізму ієрархічно впорядкованої системи, оскільки аналіз системи управління у її взаємозв'язку з сукупністю асинхронних послідовно-паралельних процесів дозволяє визначати як "оптимальний координуючий сигнал", так і шляхи вдосконалення структури керуючої системи. Першим етапом реалізації даного підходу є побудова моделі координаційної взаємодії.

Визначимо складну ієрархічно впорядковану систему як трійку:

$$C = (Ps, Cs, Q), \quad (1)$$

де $Ps = (P, Pe)$ – орієнтований граф, що представляє систему послідовно-паралельних процесів, функціонування яких забезпечує система C ; $P = \{P_1, P_2, \dots, P_i, \dots, P_K\}$ – множина вершин графа Ps , тобто процесів P_i , $i \in [1, K]$; $Pe = \{Pe_1, Pe_2, \dots, Pe_j, \dots, Pe_{Kj}\}$ – множина спрямованих дуг графа Ps , які встановлюють зв'язки між процесами множини P ; $Cs = (S, Se)$ – орієнтований граф, що представляє систему управління складною системою C ; $S = \{S_1, S_2, \dots, S_l, \dots, S_{Kl}\}$ – множина вершин графа Cs , тобто центрів прийняття рішень S_l , $l \in [1, Kl]$; $Se = \{Se_1, Se_2, \dots, Se_y, \dots, Se_{Ky}\}$ – множина спрямованих дуг, що встановлюють управляючі та зворотні зв'язки між центрами прийняття рішень множини S ; $Q = \{Q_1, Q_2, \dots, Q_v, \dots, Q_{Kv}\}$ – множина впливів налаштування та сигналів зворотного зв'язку, які встановлюють зв'язки між центрами прийняття рішень множини S та процесами множини P .

Оскільки об'єктом дослідження виступає складна ієрархічно впорядкована система, то в залежності від рівня деталізації центр прийняття рішень S_l , $l \in [1, Kl]$ може виступати як підсистема

управління $Cs_l, l \in [1, KI]$, а процес $P_i, i \in [1, K]$ – як система послідовно-паралельних процесів нижчого рівня $Ps_i, i \in [1, K]$. Виходячи з цього твердження, при розгляді складної системи на найвищому рівні доцільним є визначення процесу P_0 як умовного загального для системи C процесу, результатом декомпозиції якого є Ps . Налаштування даного процесу є задачею центру прийняття рішень найвищого рівня S_0 .

Розглянемо процес P_0 , який представляє собою цілеспрямовану діяльність по перетворенню деякого входу на вихід та має наступне формальне представлення:

$$P_0 : X_0 \rightarrow Y_0, \quad (2)$$

де X_0 – вектор входів процесу P_0 ; Y_0 – вектор виходів процесу P_0 .

При декомпозиції процесу P_0 на підпроцеси $P_i, i \in [1, K]$ визначаються вектори їх входів та виходів X_i та Y_i відповідно. Вектори X_0 та Y_0 є або компонентами векторів $X_i, i \in [1, K]$ та $Y_i, i \in [1, K]$, або розраховуються за формулою $X_0 = F_x(X_u)$ та $Y_0 = F_y(Y_u)$, де $X_u \in X$, $Y_u \in Y$, де $F_x(\cdot)$, $F_y(\cdot)$ – деякі функції агрегації, $X = \{X_0, X_i\}, i = [1, K]$, $Y = \{Y_0, Y_i\}, i = [1, K]$.

Ефективність процесу P_0 визначається функцією $\Phi(X_0, Y_0)$, де $\Phi(\cdot)$ – векторний критерій. При цьому внутрішнє та зовнішнє середовище накладають деякі обмеження на входи та виходи процесу P_0 : $X_0 \in X^S$, $Y_0 \in Y^S$. Обмеженість ресурсів в системі C та її зовнішньому середовищі на глобальному рівні представлено наступним чином: $H(X_0, Y_0) \geq b_0$.

Таким чином, задача центру прийняття рішень S_0 по забезпеченню ефективності процесу P_0 приймає наступний вигляд:

$$\begin{cases} \Phi(X_0, Y_0) \rightarrow \text{extr}, \\ H(X_0, Y_0) \geq b_0, \\ X_0 \in X^S, \\ Y_0 \in Y^S. \end{cases} \quad (3)$$

Оскільки, межі цільових значень функції $\Phi(X_0, Y_0)$ є, як правило, відомими, то перейдемо до нечіткого представлення цілі центру прийняття рішень, тоді (3) прийме наступний вигляд:

$$\begin{cases} \mu_D(\Phi(X_0, Y_0)) \rightarrow \text{extr}, \\ H(X_0, Y_0) \geq b_0, \\ X_0 \in X^S, \\ Y_0 \in Y^S, \end{cases} \quad (4)$$

де $\mu_D(\cdot)$ – функція приналежності значення цільової функції глобального центру управління S_0 до нечіткої множини ефективних значень.

Обмін керуючими сигналами та сигналами зворотного зв'язку в межах C_S відбувається шляхом передачі компоненти w_l , $l \in [1, K]$ вектора w від S_0 до S_l в тому разі, якщо в графі C_S існує направлена дуга Se_y з вершини S_0 до S_l . Розглянемо задачу центру прийняття рішень S_l :

$$\begin{cases} \mu_{Dl}(\Phi_l(X_l, Y_l, w_l)) \rightarrow \text{extr}, \\ H_l(X_l, Y_l) \geq b_l, \\ X_l \in X_l^S, \\ Y_l \in Y_l^S, \end{cases} \quad (5)$$

де $\mu_{Dl}(\cdot)$ – функція приналежності значення цільової функції центру прийняття рішень S_l до нечіткої множини ефективних значень.

Рішення задачі (5) – (X_l^*, Y_l^*) є відкликом центру прийняття рішень S_l на отриманий координуючий сигнал w_l та передається до глобального центру управління S_0 .

Задача координації S_0 приймає вигляд:

$$\begin{cases} \mu_D(\Phi(X_0, Y_0, w)) \rightarrow \text{extr}, \\ H(X_0, Y_0) \geq b_0, \\ X_0 \in X^S, \\ Y_0 \in Y^S, \end{cases} \quad (6)$$

де w – вектор координуючого сигналу елементу S_0 .

Рішенням задачі (6) є (X_0^*, Y_0^*) та значення вектору w , компоненти якого S_0 направляє до локальних центрів прийняття рішень $S_l, l \in [1, K]$. Ітеративний обмін інформацією між рівнями управління відбувається до тих пір, поки функція $\mu_D(\cdot)$ не досягне оптимуму при (X^*, Y^*) , що задовольняють обмеженням задачі (5) – (6).

Вирішення задачі координації в рамках процесного підходу до побудови координаційного механізму в постановці (5) – (6) повинно відбуватися в 3 етапи:

1. Для складної системи $C = (Ps, Cs, Q)$ зі сталою структурою процесів Ps , системи управління Cs та системи зв'язків між ними Q знайти рішення (X^*, Y^*) задачі (5) – (6), таке що $\mu_D(\Phi_0(X_0^*, Y_0^*, w))$ досягає оптимального значення.

2. Для складної системи $C = (Ps, Cs, Q)$ зі сталою структурою процесів Ps та системи управління Cs знайти таку систему зв'язків між ними Q' , при якій існує рішення (X^{**}, Y^{**}) задачі (5) – (6), таке що $\mu_D(\Phi(X^{**}, Y^{**})) > \mu_D(\Phi(X^*, Y^*))$.

3. Для складної системи $C = (Pr, Cs, Q')$ зі сталою структурою процесів Ps знайти таку структуру системи управління Cs'' та системи зв'язків Q'' при якій існує рішення (X^{***}, Y^{***}) задачі (5) – (6), таке що $\mu_D(\Phi(X^{***}, Y^{***})) > \mu_D(\Phi(X^{**}, Y^{**}))$.

Перший етап задачі координації вирішується застосуванням ітеративного алгоритму координації до задачі (5) – (6), проте вирішення другого та третього етапів вимагає застосування нових для математичної теорії методів.

Висновки. Побудовано модель координаційної взаємодії в складній ієрархічно впорядкованій системі. Запропоновано постановку задачі координації в рамках процесного підходу до побудови координаційного механізму. Подальшим напрямком досліджень буде розробка методів вирішення другого та третього етапу задачі координації.

Список літератури: 1. Алиев Р.А. Методы и алгоритмы координации в промышленных системах управления / Р.А. Алиев, М.И. Либерзон. – М.: Радио и связь, 1987. – 208 с. 2. Алтунин А.Е. Модели и алгоритмы принятия решений в нечётких условиях: монография / А.Е. Алтунин, М.В. Семухин. – Тюмень: Издательство Тюменского государственного университета, 2000. – 352 с. 3. Катренко А.В. Механізми координації у складних

ієрархічних системах / *А.В. Катренко, І.В. Савка* // Вісник Національного університету "Львівська політехніка". Серія: Інформаційні системи та мережі. – Львів : Видавництво Національного університету "Львівська політехніка", 2008. – С. 156-166. **4.** *Месарович М.* Теория иерархических многоуровневых систем / *М. Месарович, Д. Мако, И. Такахага*; пер. с англ. И.Ф. Шахнова. – М.: Мир, 1973. – 344 с. **5.** *Kim H.M.* Target cascading in optimal system design / *H.M. Kim, N.F. Michelena, P.Y. Papalambros, T. Jiang* // Proceedings of DETC 2000 26th Design Automation Conference. – 2000. – P. 14-18. **6.** *Kroo I.* Distributed multidisciplinary design and collaborative optimization [Електронний ресурс] / *I. Kroo* // VKI lecture series on Optimization Methods & Tools for Multidisciplinary Design, 2004. – Режим доступу: http://aero.stanford.edu/Reports/VKI_CO_Kroo_A.pdf. **7.** *Roth B.D.* Enhanced collaborative optimization: a decomposition-based method for multidisciplinary design / *B.D. Roth, I.M. Kroo* // Proceedings of the ASME design engineering technical conferences. – 2008. – P. 12-16. **8.** *Tosserams S.* Multi-modality in augmented Lagrangian coordination for distributed optimal design / *S. Tosserams, L.F.P. Etman, J.E. Rooda* // "Structural and Multidisciplinary Optimization". – 2010. – № 40. – P. 329-352. **9.** *Sobieszczanski-Sobieski J.* Bilevel integrated system synthesis for concurrent and distributed processing / *J. Sobieszczanski-Sobieski, T.D. Altus, M. Phillip, J.R. Sandusky* // AIAA Journal. – 2003. – № 41. – P. 1996-2003. **10.** *Лэсдон Л.С.* Оптимизация больших систем / *Л.С. Лэсдон*. – М.: Наука, 1975. – 432 с.

УДК 004.021

Модель координационного взаимодействия в сложной иерархической системе / **Плюта Н.В., Гоменюк С.И.** // Вестник НТУ "ХПИ". Тематический выпуск: Информатика и моделирование. — Харьков: НТУ "ХПИ". – 2011. – № 17. – С. 122 – 127.

Обоснована актуальность проблемы разработки моделей и методов, позволяющих комплексно решить задачу координации. Разработана модель координационного взаимодействия, которая отражает взаимосвязь между системой управления и процессами в сложной иерархической системе. Предложена постановка задачи координации для систем данного класса. Библиогр.: 10 назв.

Ключевые слова: задача координации, иерархическая система, модель координационного взаимодействия.

UDC 004.021

Model of coordination interaction in complex hierarchical system / **Plyuta N.V., Gomenyuk S.I.** // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 122 – 127.

The urgency of the problem of developing the models and methods to solve the complex coordination task is justified. The model of coordination interaction, which reflects the relationship between the control system and processes in a complex hierarchical system, is developed. Formulation of the coordination task for systems of this class is proposed. Refs.: 10 titles.

Key words: task of coordination, hierarchical system, model of coordination interaction.

Поступила в редакцию 14.02.2011

М.Ю. ПОЛЯКОВА, студентка НТУ "ХПИ", Харьков,
Б.Н. СУДАКОВ, к.т.н., проф. НТУ "ХПИ", Харьков

РАЗРАБОТКА ПОДХОДА К СОЗДАНИЮ АЛГОРИТМА СИНТАКСИЧЕСКОГО АНАЛИЗА ЕСТЕСТВЕННО- ЯЗЫКОВОГО ТЕКСТА ИНФОРМАЦИОННО-ПОИСКОВЫХ СИСТЕМ

Рассмотрены существующие методы синтаксического анализа естественно-языкового текста и выделены основные преимущества и недостатки. Разработан усовершенствованный алгоритм синтаксического анализа. Показано, что параллельное использование синтаксического и семантического анализа позволяет сократить временные затраты на обработку естественно-языкового текста. Ил.: 1. Библиогр.: 10 назв.

Ключевые слова: синтаксический анализ, семантический анализ, естественно-языковый текст.

Постановка проблемы. В настоящее время существует множество методик информационного поиска, условно разделенных на три группы. К первой группе относятся статистические методы, которые являются наиболее распространенными методами информационного поиска. Основной их особенностью является качественная математическая модель, позволяющая получать хорошие оценки релевантности файлов поиска. Поисковые машины, основанные на данных методах, отличаются простотой интерфейса. Основным минусом данного метода является тот факт, что не учитывается смысловая нагрузка текста документов коллекции и текста запроса. Отсутствие учета смысловой нагрузки текстов зачастую приводит к нерелевантным результатам. Примерами поисковых машин такого типа являются поисковые машины Google, Yandex, Rambler, Yahoo и т.д.

Основная идея второй группы методов информационного поиска заключается в том, что все исходные данные представлены в виде онтологий, а поиск ведется путем указания свойств искомого объекта. Данные методы, в отличие от статистических, учитывают смысловую нагрузку информации, поскольку информация представлена в виде онтологий. Однако данные методы имеют ряд недостатков:

- сложность пользовательского интерфейса, требующая от пользователя дополнительных затрат на конкретизацию объектов и свойств;
- большинство информации в Интернет представлено в виде простых HTML-страниц и не содержат семантического описания

контента. В качестве примера подобной системы можно рассматривать систему АСНИ (Автоматизированная система научных исследований).

К третьей группе методов информационного поиска относятся методы, которые помимо статистических методов поиска используют методы семантического анализа текстов. Данная группа методов развивается в настоящее время наиболее интенсивно. Основным плюсом систем комбинированного типа является комбинация качественной статистической модели поиска и учета семантических конструкций. Основные минусы подобных систем, существующих в настоящее время:

- большое время отклика;
- мало где используются механизмы логического вывода;
- ограничения на структуру запроса (при использовании простого пользовательского интерфейса);
- необходимость установки дополнительных параметров поиска (при использовании сложных пользовательских интерфейсов);
- большинство систем подобного типа используют в качестве исходной информации стандартные тексты, проводя семантический анализ на конечном этапе задачи поиска, что приводит к медлительности данных систем.

Несмотря на то, что третья группа методов наиболее полно отвечает требованиям, предъявляемым к системам информационного поиска на основе семантики, все системы данного типа имеют недостатки.

Анализ литературы. В [1 – 3] изложены основные принципы проектирования экспертных систем. Рассматриваются базовые концепции технологии экспертных систем. Освещаются основные схемы представления проблемно-ориентированных знаний в программах и методы применения этих знаний к построению искусственного интеллекта. В [4, 5] освещены теоретические концепции и практические методы автоматической обработки естественно-языкового текста на всех уровнях лингвистического анализа. В [6, 7] представлены методы синтаксического анализа. В [8] рассматривается решение по синтаксическому разбору структуры вводимого текста, применяемых в средах разработки для повышения эффективности работы программиста. В электронных источниках приведены базовые понятия синтаксического компьютерного анализа.

Цель статьи. Разработать подход к построению алгоритма синтаксического анализа естественно-языкового текста.

Актуальность вопроса. В последние годы большое распространение получили различного рода интеллектуальные системы, выполняющие обработку текстов на естественном языке (ЕЯ). В

простейшем случае они используются в информационно-поисковых системах (ИПС), ориентированных на естественно-языковое общение с пользователем.

Одним из основных элементов ИПС является лингвистический процессор (ЛП), выполняющий роль посредника между пользователем и базой данных, в которой хранится интересующая его информация. Классическая структура лингвистического процессора содержит три последовательных блока – для морфологического, синтаксического и семантического анализа текста. Синтаксический анализ, является главным блоком, определяющим качество работы ЛП в целом.

Задача синтаксического анализа. Используя морфологическую информацию о словоформах, необходимо построить синтаксическую структуру входного предложения (осуществить разбор предложения). Морфологический анализ – обработка отдельных словоформ, в результате которой каждой словоформе относится в соответствие ее информация – характеристика, которая отображает те свойства словоформы, которые необходимы для следующего синтаксического анализа.

К началу синтаксического анализа весь текст представляется в виде последовательности характеристик к словоформам, так что алгоритм синтаксического анализа имеет дело не со словоформами, а лишь с соответствующими характеристиками.

Программа синтаксического анализа, как правило, состоит из двух компонентов: сегментации предложения и установления связей между словами. Компоненты работают параллельно или последовательно, в зависимости от архитектуры синтаксического модуля.

Сегментация соединяет простые предложения в составе сложного. В любое простое предложение могут быть вставлены причастные или деепричастные обороты, придаточные предложения, которые, в свою очередь, тоже могут быть "разбиты" другими оборотами. Существуют примеры, когда части цельного высказывания находятся на значительном расстоянии друг от друга, а глубина вложения небольших предложений теоретически не ограничена.

На следующем этапе происходит установление связей между словами в построенных сегментах. На этом этапе появляется проблема морфологической омонимии, то есть неоднозначности. Морфологическая омонимия возникает, когда одна и та же форма может выражать разные морфологические значения. Пример: форма "весла" может быть родительным падежом единственного числа, именительным падежом множественного числа.

Явление морфологической омонимии весьма негативно отражается на скорости работы программы синтаксического анализа. На "длинных" предложениях количество комбинаторных вариантов иногда достигает нескольких сотен, поэтому используются разного рода математические и лингвистические ухищрения, позволяющие избежать анализа всех комбинаторно возможных вариантов.

Разработка основных принципов построения алгоритма синтаксического анализа. Синтаксический анализ может рассматриваться как процесс поиска дерева синтаксического анализа. Существуют два противоположных способа задания соответствующего пространства поиска. Во-первых, можно начать с вершины и искать дерево, листьями которого являются соответствующие слова. Такой способ называется нисходящим синтаксическим анализом (поскольку символ вершины изображается в верхней части рисунка, на котором показано дерево, повернутое корнем вверх). Во-вторых, можно начать со слов и выполнять поиск дерева начиная с корня. Такой метод называется восходящим синтаксическим анализом.

Традиционным методом построения синтаксической структуры является метод фильтров. В данном методе построение дерева зависимостей начинается с построения наборов всевозможных связей (синтаксических отношений) между словами. При этом для большинства слов устанавливается несколько потенциально возможных связей с различными управляющими словами. Применение фильтров позволяет отбросить многие из первоначально установленных связей и, в идеале, получить количество связей, равное числу слов (при условии, что между всеми словами установлены синтаксические отношения). В чистом виде метод фильтров для практической реализации неприменим, так как число всевозможных связей между словами весьма велико, а число всевозможных способов выбора из них конкретного дерева зависимостей огромно. На практике для получения эффективных алгоритмов необходимо применять методы, направляющие и ускоряющие выбор правильных вариантов анализа. Одним из таких методов является метод параллельного синтаксического и семантического анализов, в котором последний выступает в качестве фильтра тупиковых вариантов.

Создание алгоритма синтаксического анализа естественно-языкового текста. Алгоритм представляет собой структуру построения синтаксических групп (СГ), которая состоит из существительного с зависимыми от него словами. Присоединяя к существительному словоформы, находящиеся в предложении слева от него, отыскиваются все СГ (B1 – B2) (см. рис.).

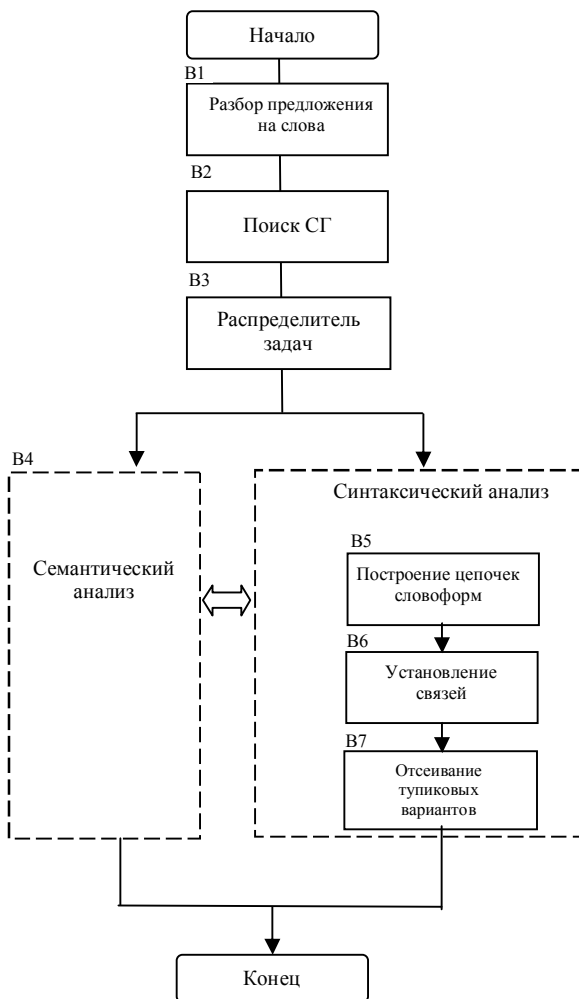


Рис. Алгоритм синтаксического анализа

Для установления поверхностно-синтаксических связей в группах анализ начинается с конца СГ. При этом анализируются пары слов с целью установления связей между ними. Если связь установлена, то переходят к следующей паре. Если для соотношения слов найдено

главное слово, то записывается его номер. Анализ продолжается до тех пор, пока в анализируемой группе все соотношения слов не будут иметь главные слова, кроме одной словоформы. Если этого сделать не удастся, то уточняется вариант разбиения на группы. Строятся цепочки словоформ и устанавливается связь между ними в предложении (B5 – B6).

Для отсеивания тупиковых вариантов параллельно используется семантический анализ (B4), направленный на решение задач, связанных с возможностью понимания смысла фразы и выдачи запроса поисковой системе в необходимой форме. Используя модуль семантического анализа текстов, повышают эффективность лингвистических систем (программ автоматического перевода, информационно-поисковых систем, рубрикаторов и рефераторов текстов) на основе реализации извлечения и обработки смысловой информации. Таким образом, время, затрачиваемое на синтаксический анализ, сокращается за счет отброса заранее не существующих соотношений между словоформами (B7). Описанная схема алгоритма представлена на рис.

Выводы. Предложен новый метод синтаксического анализа для информационно-поисковой системы, суть которого заключается в более быстром нахождении связей между словоформами и подготовке качественного материала для семантического уровня. Это позволяет в целом улучшить качество работы лингвистического процессора поисковой системы, поскольку семантический и синтаксический анализы проходят параллельно. За счет этого заранее отбрасываются тупиковые варианты.

Список литературы: 1. Джексон П. Введение в экспертные системы / П. Джексон. – М.: Основа, 2001. – 624 с. 2. Карпова Г.Д. Компьютерный синтаксический анализ: описание моделей и направлений разработок / Г.Д. Карпова, Ю.К. Пирогова, Т.Ю. Кобзарева, Е.В. Микаэлян. – М.: ВИНТИ, 1991. – 304 с. 3. Поспелова Д.А. Искусственный интеллект / Д.А. Поспелова. – М.: Наука, 1990. – 246 с. 4. Романенко В.Н. Сетевой информационный поиск / В.Н. Романенко. – М.: Профессия, 2005. – 290 с. 5. Фостер Дж. Автоматический синтаксический анализ / Дж. Фостер. – М.: Мир, 1975. – 71 с. 6. Попов Э.В. Общение с ЭВМ на естественном языке / Э.В. Попов. – К.: Наука, 1992. – 254 с. 7. Анно Е.И. Типологии алгоритмов синтаксического анализа (для формальных моделей естественного языка) / Е.И. Анно. – СПб.: Питер, 1989. – 152 с. 8. Омар А.Х. Авадала Подход к синтезу естественно-языковых сообщений по формализованному представлению базы знаний: автореф. дис. на соискание ученой степени канд. техн. наук / Омар А.Х. Авадала. – Х., 2001. – 20 с. 9. Руских И.В. Инкрементальный синтаксический анализ в средах разработки и текстовых редакторах // Нижегородский университет. – 2007. – С. 277 10. Автоматическая обработка текста [Электронный ресурс] / Леонтьев Н.Н. – 2003. – С. 5. – Режим доступа к статье: <http://www.aot.ru/technology/html>.

Статья представлена д.т.н., проф. НТУ "ХПИ" Серковым А.А.

УДК 004.031.42

Розробка підходу до створення алгоритму синтаксичного аналізу природно-мовного тексту інформаційно-пошукових систем/ Судаков Б.М., Полякова М.Ю. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 128 – 134.

Розглянуто існуючі методи синтаксичного аналізу природно-мовного тексту та виділено основні переваги та недоліки. Розроблено удосконалений алгоритм синтаксичного аналізу. Показано, що паралельне використання синтаксичного та семантичного аналізу дозволяє скоротити часові витрати на обробку природно-мовного тексту. Іл.: 1. Бібліогр.: 10 назв.

Ключові слова: синтаксичний аналіз, семантичний аналіз, природно-мовний текст.

UDC 004.031.42

Developing an approach to creating an algorithm parsing natural language text of information retrieval systems / Sudakov B.N., Poliakova M.U. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 128 – 134.

The article deals with the existing methods of parsing natural language text. The main advantages and disadvantages were identified. Developed an improved algorithm for parsing. It is shown that the parallel use of syntactic and semantic analysis can reduce the time required for processing natural language text. Figs.: 1. Refs.: 10 titles.

Key words: syntactic analysis, semantic analysis, natural language text.

Поступила в редакцію 26.01.11

Н.О. РИЗУН, к.т.н., доц. ДУЭП им. А. Нобеля, Днепропетровск

КОНЦЕПЦИЯ ПОСТРОЕНИЯ ЭКСПЕРТНОЙ СИСТЕМЫ ПОДДЕРЖКИ ПРИНЯТИЯ РЕШЕНИЙ ПО УПРАВЛЕНИЮ УЧЕБНЫМ ПРОЦЕССОМ В ВУЗЕ

Предложена концепция построения модульной структуры экспертной системы поддержки принятия решений по управлению учебным процессом в ВУЗе на базе развития методологий совершенствования теории тестирования и распределения учебной нагрузки с использованием инструментов системного анализа и проектирования SADT. Ил.: 5. Библиогр.: 11 назв.

Ключевые слова: экспертная система, принятие решений, теория тестирования, системный анализ, учебная нагрузка, проектирование, SADT.

Постановка проблемы. Главной целью функционирования системы управления учебным процессом в ВУЗе, по-прежнему, остается повышение уровня образования за счет эффективной и скоординированной организации работы с предоставлением возможности контроля, анализа и корректировки принятых управленческих решений на базе объективных результатов непрерывного мониторинга количественного и качественного уровня знаний студентов. Поскольку одним из эффективных инструментов оценки уровня знаний в настоящее время являются системы автоматизированного тестового контроля, актуальной является научная проблема интеллектуализации методологических и технологических инструментов тестирования путем разработки и создания на их основе экспертных систем поддержки принятия решений по управлению учебным процессом в ВУЗе.

Анализ литературы. В настоящее время ведутся исследования в области разработки методик совершенствования теории тестового контроля и теории расписаний. Основным недостатком рассматриваемых методов является отсутствие системного подхода и учета взаимосвязи разрабатываемых методов с глобальной целью образовательного процесса – повышением качества управления и организации учебного процесса. Работы характеризуются узкой направленностью на автоматизацию технологических процессов тестирования [1, 2]; малое внимание уделяется методикам формирования качественного тестового материала как одной из главных предпосылок обеспечения адекватности и объективности оценки [3, 4]; методология составления расписания строится на принципах полной автоматизации, не позволяющей учитывать специфику каждого отдельного учебного заведения [5, 6].

Целью статьи является разработка концепции построения модульной структуры экспертной системы поддержки принятия решений по управлению учебным процессом в ВУЗе на базе развития методологий совершенствования теории тестирования и распределения учебной нагрузки с использованием инструментов системного анализа и проектирования SADT [7].

Результаты исследований. Поставленная в работе научная проблема предопределила необходимость использования для проектирования языка IDEF, предоставляющего возможности функционального и процессного моделирования экспертных систем.

Согласно выбранной методологии на нулевом уровне декомпозиции была сформулирована цель проектируемой экспертной системы поддержки принятия решений (ЭСППР), а также входы (*Учебный материал*), выходы (*Качественные знания студентов*), управляющие воздействия (*Стандарты Министерства образования и науки Украины*) и ресурсы (*Методы организации учебного процесса*).

Дальнейшая декомпозиция процесса повышения качества учебного процесса на основе процессных *IDEF0* и функциональных *IDEF3* моделей позволила на первом уровне выделить основные четыре взаимосвязанных модуля принятия решений по управлению качеством учебного процесса (УКУП) (рис. 1).

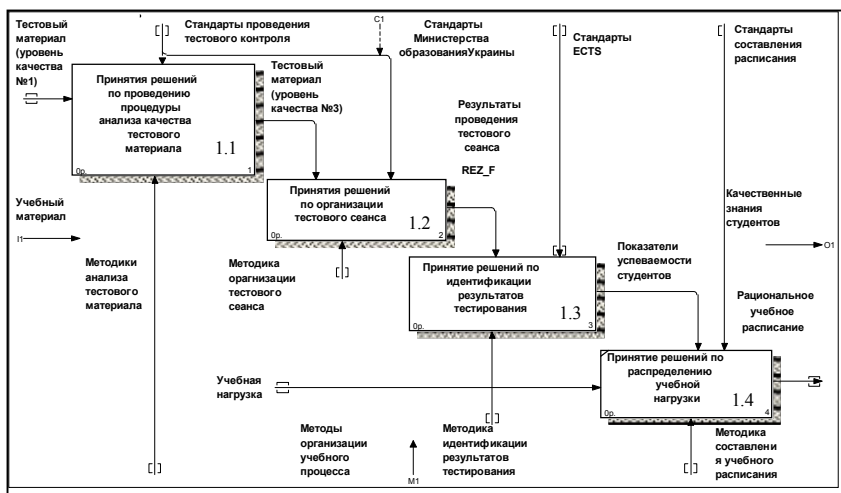


Рис 1. Модель *IDEF0* – уровень декомпозиции A1 процесса 1.1.

Новизна полученного научного результата состоит в обосновании необходимости выделения отдельных подпроцессов модуля тестирования знаний студентов: анализа качества тестового материала (процесс 1.1), организации тестового сеанса (процесс 1.2) и идентификации результатов тестирования (процесс 1.3) [8] с целью повышения уровня объективности средств измерения и самой процедуры тестирования.

Декомпозиция *второго* уровня, представленная на рис. 2, детализирует функциональную модель процесса принятия решений по проведению "Процедуры расширенного анализа качества тестового материала" (процесс 1.1).

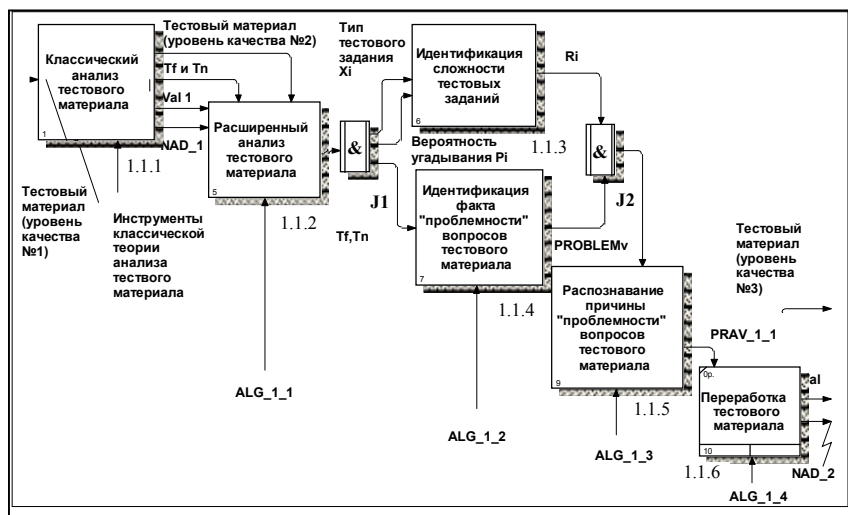


Рис 2. Модель *IDEF3* – уровень декомпозиции *A1.1* процесса 1.1

Управляющим параметром предлагаемой методологии принятия решений по процессу 1.2 является введенное авторами соотношение между нормативным *TN* и фактическим *TF* временем на выполнение тестового задания, позволившее выделить дополнительные функциональные подпроцессы 1.1.2 – 1.1.5, реализующие алгоритмы принятия решений по расширенному анализу качества тестового материала *ALG_1_1*, идентификации *ALG_1_2* и распознаванию причины *ALG_1_3* "проблемности" вопроса, характеризующегося низкими показателями надежности *NAD_1* и валидности теста *VAL_1*, а также процедуры повторной переработки тестового материала *ALG_1_4*, позволяющей, по результатам проведенных исследований и анализа полученной статистики, повысить значения показателей надёжности

NAD_2 и валидности VAL_2 тестового материала на 15 – 20% [9].

Декомпозиция процесса принятия решений по процессу 1.2. "Организации тестового сеанса" позволила выделить дополнительные функциональные подпроцессы 1.2.1, 1.2.3 – 1.2.5 с целью реализации предлагаемого авторами [10] адаптивного алгоритма последовательной подачи ALG_2_1 тестовых заданий разных типов в порядке уменьшения уровня сложности ALG_2_2 и идентификации показателя "достаточности" прохождения студентом индивидуального количества типов тестовых заданий, которое зависит от фактически установленного в процессе тестирования уровня знаний студента ALG_2_3 , а также алгоритма проверки истинности гипотезы угадывания и/или неустойчивости знаний студента ALG_2_4 , состоящего в формализации условий необходимости повторения тестового задания, вызвавшего подозрение с точки зрения значений динамического коэффициента $D_k = \{\text{Очень быстро, Очень медленно}\}$ (рис. 3).

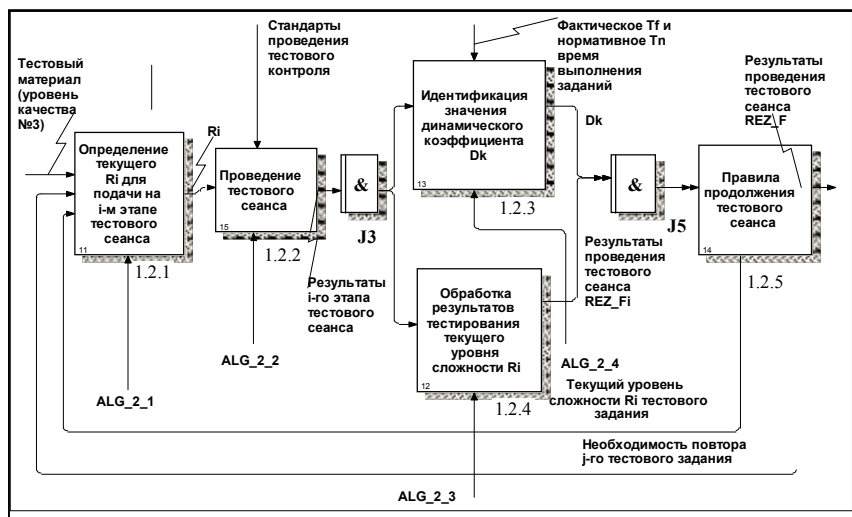


Рис 3. Модель IDEF3 – уровень декомпозиции A1.2 процесса 1.2

Управляющим параметром предлагаемой методологии принятия решений по процессу 1.3 "Идентификации результатов тестирования" (рис. 4) является предложенный авторами [11] показатель коэффициента корреляции K_k между рядами фактического и нормативного времени, потраченного на правильный ответ как инструмента оценки степени устойчивости знаний и управления методами снижения влияния фактора

"угадывания" на объективность результатов тестирования, позволивший выделить дополнительные подпроцессы 1.3.1 – 1.3.4 реализации алгоритмов корректировки фактического набранного количества баллов в зависимости от величины K_k , уровня сложности вопроса R_i и значения динамического коэффициента D_k как показателя степени несоответствия эталону фактического времени (D_k), потраченного на правильный ответ: $ALG_3_1 - ALG_3_3$.

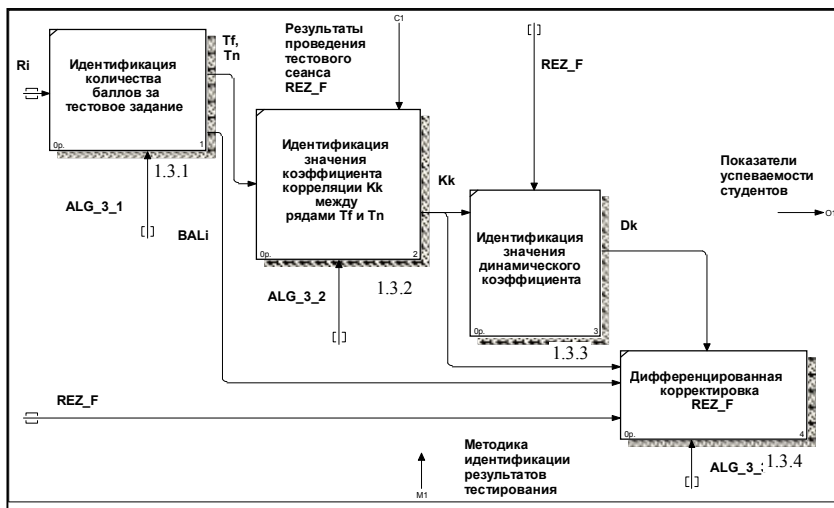


Рис 4. Модель *IDEF0* – уровень декомпозиции *A1.3* процесса 1.3.

Декомпозиция процесса принятия решений по процессу 1.4. "Распределение учебной нагрузки" (рис. 5) позволила выделить дополнительные функциональные подпроцессы 1.4.3, 1.4.5, 1.4.6, призванные обеспечить повышение точности и качества полученного расписания учебных занятий для всех курсов и направлений учебы за счет реализации алгоритмов минимизации синтаксических и логических ошибок в полученном расписании Alg_4_3 , а также предоставления фактической информации о качестве полученного расписания Alg_4_5 с возможностью дальнейшей его корректировки Alg_4_6 с учетом биологических ритмов и уровня сложности учебных дисциплин, в свою очередь, зависящую от показателей количества междисциплинарных связей, среднего балла по дисциплине и коэффициента степени наличия логической последовательности расположения в полученном расписании лекционных и практических занятий в течение недели, а также

чередования разных видов учебных занятий в пределах дня.

Формализация рекомендации *RECOM* призваны обеспечить равномерность загрузки студентов и способствовать повышению уровня усвоения материала и эффективности учебного процесса в целом [12].

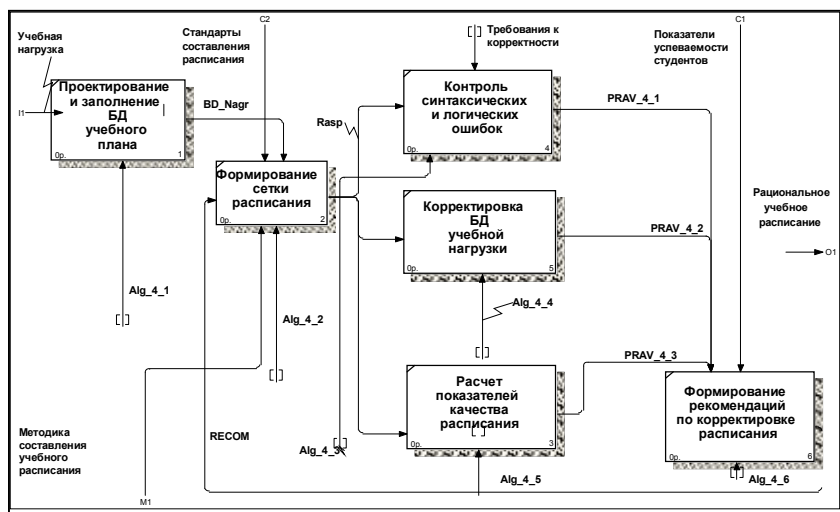


Рис 5. Модель *IDEF0* – уровень декомпозиции *A1.4* процесса 1.4.

Выводы. Таким образом, предложенная авторами концепция построения экспертной системы позволяет решить поставленную научную задачу путем разработки декомпозиционной методики поэтапной реализации взаимосвязанных модулей принятия решений по совершенствованию процесса управления качеством учебного процесса в ВУЗе.

Список литературы: 1. Мелехин В.Б. Автоматизированная система контроля знаний студентов в вузе / В.Б. Мелехин, Е.И. Павлюченко // Транспортное дело России. – 2009. – №1. – С. 23-25. 2. Ціделко В.Д., Яремчук Н.А., Шведова В.В. Автоматизована система тестування, навчання та моніторингу. Пат. 43616 України: МПК G09B 7/00 / Замовник та патентовласник: Національний технічний університет України "Київський політехнічний інститут". – № 200902620, заявл. 23.03.2009, опубл. 25.08.2009, Бюл. № 16, 2009. 3. Белоус Н. Методика определения качества тестовых заданий, оцениваемых по непрерывной шкале / Н. Белоус, И. Куцевич, И. Белоус // International Book Series "Information Science and Computing". The paper is selected from XVth International Conference "Knowledge-Dialogue-Solution". – Kyiv, 2009. – С. 127-133. 4. Максимова О.А. Технология комплексной экспертизы качества тестовых материалов для контроля учебных достижений обучающихся / О.А. Максимова // Журнал научных публикаций аспирантов и докторантов. – 2008. – № 10. – С. 140-146. 5. Дуденгефер А.А. Методика автоматизации составления

расписания средствами табличного процессора Excel. – пст. Комсомольск-на-Печере, Троицко-Печерский район, Республика Коми. – 2009. – 26 с. **6. Ревчук И.Н.** Анализ расписания занятий с использованием электронных таблиц. – 2009. [Электронный ресурс] / Режим доступа: URL: http://www.rusedu.ru/detail_2733.html. **7. Куликов Г.Г.** Автоматизированное проектирование информационно-управляющих систем. Проектирование экспертных систем на основе системного моделирования: монография / Г.Г. Куликов, А.В. Речкалов, Л.Р. Черняховская, А.Н. Набатов, Е.Б. Старцева, Н.О. Никулина. – Уфа: Изд. Уфимск. гос. авиац. техн. ун-та, 1999. – 224 с. **8. Ризун Н.О.** Эвристический алгоритм совершенствования технологии оценки качества тестовых заданий / Н.О. Ризун. – "Східно-Європейський журнал передових технологій". – № 3/11 (45). – 2010. – С. 40-49. **9. Ризун Н.О.** Методика адаптивного інтелектуального тестування знань студентів // Н.О. Ризун, Ю.К. Тараненко, Р.С. Моргун / II Всеукраїнська науково-практична конференція "Інформатика та системні науки". – 2011. – Полтава. – С. 98–101. **10. Тараненко Ю.К., Ризун Н.О.** Спосіб виміру рівня знань учнів при комп'ютерному тестуванні. Патент на корисну модель № 51559 України: МПК G06F 7/00. Замовник та патентовласник: Тараненко Ю.К., Ризун Н.О. – № u200913726, заявл. 28.12.2009, опубл. 26.07.2010. Бюл. № 14, 2010. **11. Тараненко Ю.К., Ризун Н.О.** Спосіб складання потижневого розкладу навчальних занять у вищому навчальному закладі з використанням електронної таблиці. Патент на корисну модель 51682 Україна: МПК G06F 7/00. Замовник та патентовласник: Тараненко Ю.К., Ризун Н.О. – № u201001399, заявл. 11.02.2010, опубл. 26.07.2010. Бюл. № 14, 2010.

Стаття представлена д.т.н., с.н.с. ДУЕП Тараненко Ю.К.

УДК 681.3:378.146

Концепція будування експертної системи підтримки прийняття рішень по керуванню навчальним процесом у ВНЗ / Ризун Н.О. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 135 – 141.

Запропонована концепція будування модульної структури експертної системи підтримки прийняття рішень по керуванню навчальним процесом у ВНЗ на базі розвитку методології вдосконалення теорії тестування та розподілу навчального навантаження із використанням засобів системного аналізу та проектування SADT. Ил.: 5. Бібліогр.: 11 назв.

Ключові слова: експертна система, прийняття рішень, теорія тестування, системний аналіз, навчальне навантаження, проектування, SADT.

UDC 681.3:378.146

The concept of building the expert system of supporting the decision making in the study process controlling in universities / Rizun N.O. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 135 – 141.

The concept of building the module structure of the expert system of supporting the decision making in the study process controlling in universities on the basis of development of the methodology of testing theory improvement and study workload distribution with the usage of system analysis and projecting means SADT is suggested. Figs.: 5. Refs.: 11 titles.

Key words: expert system, decision making, testing theory, system analyse, study workload, projecting, SADT.

Поступила в редколлегию 15.02.2011

Г.А. САМИГУЛИНА, д.т.н., зав. лаб. Института проблем информатики и управления Министерства образования и науки Республики Казахстан, Алматы,

С.В. ЧЕБЕЙКО, к.х.н., с.н.с. Института проблем информатики и управления Министерства образования и науки Республики Казахстан, Алматы

РАЗРАБОТКА ТЕХНОЛОГИИ ИММУННОСЕТЕВОГО МОДЕЛИРОВАНИЯ ДЛЯ КОМПЬЮТЕРНОГО МОЛЕКУЛЯРНОГО ДИЗАЙНА ЛЕКАРСТВЕННЫХ ПРЕПАРАТОВ

Разработан иммунносетевой подход к моделированию зависимостей "структура-свойство" лекарственных препаратов. Предложенная интеллектуальная технология на основе искусственных иммунных систем позволяет уменьшить погрешности энергетических оценок и повысить достоверность прогноза зависимости "структура-свойство" химических соединений. Библиограф.: 12 назв.

Ключевые слова: технология иммунносетевого моделирования, интеллектуальная технология, погрешности энергетических оценок.

Постановка проблемы и анализ литературы. Создание методов прогнозирования свойств новых химических соединений и направленный компьютерный молекулярный дизайн соединений с заданным набором свойств являются важнейшими и актуальными задачами биоинформатики. Применение последних достижений вычислительной техники и новейших информационных технологий открывает широкие возможности для решения одной из главных проблем современной науки – целенаправленного поиска новых веществ и материалов с заранее заданными свойствами, в том числе проектирование новых лекарственных средств.

История дизайна с помощью компьютеров началась более 25 лет назад, когда стало возможным изображение и вращение молекул на экране компьютера [2]. Компьютерный молекулярный дизайн основан на концепции взаимосвязи молекулярной структуры и биологической активности химических соединений. Данное направление предполагает создание принципиально новых компьютерных алгоритмов и программ поиска и отбора активных веществ целевого назначения.

Количественное описание молекулярной структуры химических соединений в компьютерном молекулярном дизайне осуществляется с помощью дескрипторов [3]. Дескриптор – это математический параметр, который характеризует структуру органического соединения, отмечая

наиболее важные черты этой структуры. Существует проблема создания дескрипторов, наиболее полно характеризующих рассматриваемое соединение и позволяющих в удобной форме использовать их в вычислительном процессе. Построение адекватной компьютерной молекулярной модели позволяет в дальнейшем прогнозировать различные терапевтические и физико-химические свойства синтезируемых молекул, что определяет актуальность и перспективность развития данного научного направления.

Среди методов прогнозирования зависимости "структура – свойство" следует отметить рост исследований по искусственным нейронным сетям [4]. В рамках поиска зависимостей между структурами органических соединений и их биологической активностью наиболее популярна многослойная нейронная сеть прямого распространения, обучающаяся по методу обратного распространения ошибки.

Моделирование биологической активности органических соединений также возможно с помощью нового биологического направления искусственного интеллекта – искусственных иммунных систем (ИИС).

Процессы, происходящие при обработке информации естественными системами и принципы их функционирования, поражают своей эффективностью, экономичностью и быстроедействием [5, 6]. Прежде всего, вызывает интерес способность данных систем решать многомерные задачи огромной вычислительной сложности в реальном времени. ИИС – это адаптивные системы для обработки и анализа данных, которые представляют собой математическую структуру, имитирующую некоторые функции иммунной системы человека и обладающие способностью к обучению, к прогнозированию на основе уже имеющихся временных рядов и принятию решения в незнакомой ситуации. ИИС в принципе не нуждаются в заранее известной модели, а строят ее сами на основе полученной информации в виде временных рядов. Данные системы применяются при решении плохо алгоритмизируемых задач, таких как прогнозирование, классификация и управление.

Математическая основа подхода ИИС заключается во введении понятия формального пептида как математической абстракции свободной энергии белковой молекулы от ее пространственной формы, описанной в алгебре кватернионов. В работах [7, 8] предложена математическая модель формального пептида.

При реализации интеллектуальных систем, основанных на выше приведенных принципах, существует ряд проблем [9, 10]. Основная трудность заключается в создании алгоритмов безошибочного

распознавания образов, так как ошибки энергетических оценок не позволяют добиться сто процентного распознавания. Особенно эта проблема актуальна для схожих по структуре формальных пептидов, которые находятся на границах различных классов и разделение между классами нелинейно. Как и в искусственных нейронных сетях, существует проблема создания эффективных и простых методик обучения иммунной сети за минимально короткое время. Необходимо из множества факторов выделить главные, которые оказывают наибольшее влияние на процесс обработки информации, построить оптимальную структуру иммунной сети на основе информативных дескрипторов, обучать ИИС и оценить процесс обучения. Проблема значительно усложняется при увеличении размерности системы.

Кроме того, очень важным является способность иммунной сети обобщать результат на новые данные, которые не были использованы в обучающем множестве. Таким образом, решение задачи минимизации ошибки обобщения позволяет повысить прогностическую способность модели и является наиболее трудной при построении данных систем.

Цель статьи. Разработать эффективную интеллектуальную информационную технологию для компьютерного молекулярного дизайна (моделирования и предсказания свойств новых лекарственных препаратов с заданными параметрами) на основе биологического подхода искусственных иммунных систем.

Технология иммунносетевого моделирования. Разработана интеллектуальная информационная технология, позволяющая моделировать зависимость "структура – свойство" на основе искусственных иммунных сетей [11].

Используется следующий алгоритм:

- описываются структуры исследуемых соединений числовыми параметрами (дескрипторами), создаются базы данных (БД);
- осуществляется предварительная обработка данных: нормирование, центрирование, заполнение пропущенных данных;
- выбирается оптимальный набор дескрипторов, строится оптимальная структура иммунной сети;
- весь массив данных разбивается на обучающую и контролирующую выборки;
- экспертами осуществляется классификация решений;
- производится обучение иммунной сети с учителем;
- решается задача распознавания образов и нахождения минимальной энергии связывания между формальными пептидами (антителами и антигенами);

– осуществляется оценка решения задачи распознавания образов на основе гомологов и расчет коэффициентов риска прогнозирования на основе ИИС;

– осуществляется прогноз свойств неизвестных соединений.

Рассмотрим подробнее реализацию данного алгоритма. Разработанная интеллектуальная технология состоит из трех основных этапов:

Этап 1. Предварительная обработка данных

Пусть исходная совокупность данных записана в виде матрицы $A = (a_{ij})$ ($i = 1, \dots, m, j = 1, \dots, n$). Так как дескрипторы, характеризующие вещества, измеряются в разных единицах, то результат может существенно зависеть от выбора масштаба измерения. Поэтому необходим переход к безразмерным величинам с помощью нормирования и центрирования дескрипторов. Для этого элементы каждого вектора преобразуем таким образом, чтобы математическое ожидание было равно нулю, а дисперсия – единице.

Основной целью нормирования данных является приведение их к сопоставимому виду. Новая матрица стандартизированных переменных

X записывается из элементов: $x'_{ij} = \frac{x_{ij} - m_j}{s_j}$ где m_j – среднее значение

исходных элементов j -го вектора; s_j – стандартное отклонение исходных элементов j -го вектора, которое вычисляется по формуле:

$$s_j = \left(\frac{1}{N-1} \sum_{i=1}^n (x_{ij} - m_j)^2 \right)^{\frac{1}{2}}.$$

Задача снижения размерности анализируемого признакового пространства и отбора наиболее информативных дескрипторов решена с помощью факторного анализа и метода главных компонент на основе вращения собственного вектора [12].

Определим базисное пространство R и проекции векторов данных на каждую из n ортогональных осей. Тогда исходную матрицу данных A размерности $(m \times n)$ можно представить в виде:

$$A = CV^T,$$

где V – матрица, столбцы которой ортогональны оси; C – матрица, строками которой являются координаты проекций каждого вектора данных в базисном пространстве R . Новую матрицу B получим следующим образом:

$$B = R^T A.$$

Матрица преобразования R^T в двумерном пространстве имеет вид:

$$R^T = \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix}.$$

Рассчитывается корреляционная матрица:

$$C = \frac{1}{N-1}(X^T X),$$

где N – число столбцов в матрице X .

Пусть $Y = B^T$, $X = A^T$, тогда получим: $Y = XR$, $Y^T = R^T X^T$.

Необходимо найти матрицу преобразования R^T такую, чтобы, применив ее к матрице X , получить новую матрицу Y , которая удовлетворяет выражению:

$$Y^T Y = R^T X^T X R = R^T C R = \Lambda,$$

где Λ – диагональная матрица.

Необходимо, чтобы выполнялось условие $CR = \lambda R$, тогда получим:

$$(C - \lambda I)R = 0, \quad (1)$$

где λ – вектор диагональных элементов в матрице Λ .

Задача будет иметь решение при выполнении условия:

$$|C - \lambda I| = 0.$$

После нахождения вектора λ подставим его в (1) и найдем матрицу преобразования R .

На основе проведенных преобразований исходные данные можно изобразить в новой системе, где координатные оси являются собственными векторами. После анализа дескрипторов (для построения оптимальной структуры иммунной сети) необходимо отбросить те, которые лежат ближе к началу координат и являются наименее информативными.

Этап 2. Распознавание образов

Ключевым моментом в разработанной интеллектуальной технологии на основе ИИС является решение задачи распознавания образов [8]. Для каждого класса, выделенного экспертами, формируются матрицы эталонов $A_1, A_2, A_3, \dots, A_n$ (n – количество классов). Выполнив

сингулярное разложение данных матриц, получаем правые и левые сингулярные векторы $\{x_1, y_1\}$, $\{x_2, y_2\}$ и т. д. эталонных матриц. Затем формируется множество матриц, рассматриваемых в качестве образов: $B_1, B_2, B_3, \dots, B_m$ (m – количество образов).

Согласно подходу ИИС энергию связи между формальными пептидами можно представить в виде:

$$W_1 = -x_1^T B y_1, W_2 = -x_2^T B y_2, W_3 = -x_3^T B y_3, \dots, W_n = -x_n^T B y_n,$$

где t – символ транспонирования.

Нативная (функциональная) укладка белковой цепи соответствует минимуму энергии связи, поэтому минимальное значение энергии связи определяет класс n , которому принадлежит данный образ: $W_n = \min \{W_1, W_2, W_3, \dots, W_n\}$.

Этап 3. Оценка энергетических погрешностей

Обработка многомерной совокупности данных на основе технологии ИИС неизбежно приводит к увеличению энергетических погрешностей, зависящих от ряда факторов, и существенно влияет на достоверность прогноза. Разработана процедура оценки энергетических погрешностей на основе гомологичных белков [10].

Выводы. Достоинством предложенной интеллектуальной технологии на основе иммунносетевого моделирования является:

- способность системы глубоко анализировать скрытые (латентные) взаимодействия между дескрипторами и основополагающие факторы, влияющие на них;
- распознавать пептиды, находящиеся на границе нелинейно разделенных классов (имеющие схожие структуры);
- сокращение времени на обучение иммунной сети за счет построения оптимальной структуры и редукции дескрипторов, несущих существенные погрешности;
- уменьшение погрешностей энергетических оценок, так называемых ошибок обобщения; повышение достоверности прогноза зависимостей "структура – свойство" химических соединений.

На разработанное программное обеспечение получены два авторских свидетельства о государственной регистрации объекта интеллектуальной собственности.

Список литературы: 1. Кубины Г. В поисках новых соединений – лидеров для создания лекарств / Г. Кубины // Российский химический журнал. – 2006. – № 2. – С. 5-17. 2. Иванов А.С. Интегральная платформа "От гена до прототипа лекарства" in silico и in vitro / А.С. Иванов, А.В. Веселовский, А.В. Дубанов, В.С. Скворцов, А.И. Арчаков // Российский химический журнал, 2006. – № 2. – С. 18-35. 3. Раевский О.А. Дескрипторы водородной

связи в компьютерном молекулярном дизайне / *О.А. Раевский* // Российский химический журнал, 2006. – № 2. – С. 97-108. **4.** *Гальберштам Н.М.* Нейронные сети как метод поиска зависимостей структура – свойство органических соединений / *Н.М. Гальберштам, И.И. Баскин, В.А. Палиулин, Н.С. Зефирова* // Успехи химии, 2003. – № 72 (7). – С. 706-727. **5.** *Альбертс Б.* Молекулярная биология клетки / *Б. Альбертс, Д. Брей, Дж. Льюис* – М.: Мир, 1994. – Т. 2. – С. 287-301. **6.** Искусственные иммунные системы и их применение / *Под ред. Д. Дасгупта*. – М.: Физматлит, 2006. – 344 с. **7.** *Тараканов А.О.* Математические модели биомолекулярной обработки информации: формальный пептид вместо формального нейрона / *А.О. Тараканов* // Проблемы информатизации. – 1998. – С. 65-70. **8.** *Tarakanov A.O.* Formal peptide as a basic of agent of immune networks: from natural prototype to mathematical theory and applications / *A.O. Tarakanov* // Proceedings of the I Int. workshop of central and Eastern Europe on Multi-Agent Systems (CEEMAS'99). – St. Petersburg, Russia, June 1-4, 1999. – P.281-292. **9.** *Самигулина Г.А.* Разработка интеллектуальных экспертных систем управления на основе искусственных иммунных систем / *Г.А. Самигулина*. – Алматы: ИПИУ МОН РК, 2010. – 252 с. **10.** *Самигулиной Г.А.* Разработка интеллектуальных экспертных систем прогнозирования и управления на основе искусственных иммунных систем / *Г.А. Самигулиной* // Проблемы информатики. – Новосибирск, 2010. – № 1. – С. 15-22. **11.** *Самигулина Г.А.* Прогнозирование зависимости структура-свойство органических соединений на основе иммунносетевого моделирования / *Г.А. Самигулина, С.В. Чебейко* // Химический журнал Казахстана. – Алматы, 2010. – № 3. – С. 164-172. **12.** *Иберла К.* Факторный анализ / *К. Иберла*. – М.: Статистика, 1980. – 304 с.

УДК 004.89:004.41.

Розробка технології імунносетевого моделювання для комп'ютерного молекулярного дизайну лікарських препаратів / Самигуліна Г.А., Чебейко С.В. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 142 – 148.

Розроблений імунносетевої підхід до моделювання залежностей "структура-властивість" лікарських препаратів. Запропонована інтелектуальна технологія на основі штучних імунних систем дозволяє зменшити погрішності енергетичних оцінок і підвищити достовірність прогнозу залежності "структура-властивість" хімічних сполук. Бібліогр.: 12 назв.

Ключові слова: технології імунносетевого моделювання, інтелектуальна технологія, погрішності енергетичних оцінок.

УДК 004.89:004.41.

Development of immune-networks modeling technology for computers molecular design of medical products / Samigulina G.A., Chebeiko C.B. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 142 – 148.

It is developed immune nets the approach to modeling dependences "structure – property" of medical drugs. The offered intellectual technology allows to reduce errors of power estimations and to raise reliability of the forecast of dependence "structure – property" of chemical compounds. Refs.: 12 titles.

Keywords: immune-networks modeling technology, intellectual technology, errors of power estimations.

Поступила в редакцію 14.02.2011

С.А. СУББОТИН, к.т.н., доц. ЗНТУ, Запорожье

МЕТОДЫ ФОРМИРОВАНИЯ ВЫБОРОК ДЛЯ ПОСТРОЕНИЯ ДИАГНОСТИЧЕСКИХ МОДЕЛЕЙ ПО ПРЕЦЕДЕНТАМ

Решена актуальная задача разработки математического обеспечения для формирования обучающих выборок. Получили дальнейшее развитие переборные и эволюционные методы комбинаторного поиска, которые модифицированы для формирования выборок путем введения разработанных критериев для отбора, цензурирования и псевдокластеризации экземпляров, что позволяет ускорить процесс формирования выборок и обеспечить их соответствие заданным критериям при ограниченном объеме. Определены оценки сложности разработанных методов. Библиогр.: 9 назв.

Ключевые слова: выборка, эволюционные методы комбинаторного поиска, диагностическая модель.

Постановка проблемы и анализ литературы. Для обеспечения конкурентоспособности и высокого качества выпускаемой продукции, её безотказности в процессе эксплуатации возникает необходимость в своевременном выполнении диагностических процедур, что, в свою очередь, требует наличия диагностической модели [1]. Поскольку в подавляющем большинстве практических задач экспертный опыт является весьма ограниченным, он оказывается недостаточным для построения диагностической модели и возникает необходимость построения модели по набору прецедентов – обучающей выборке, извлекаемой из доступной исследователю исходной выборки.

Традиционным подходом для выделения обучающей выборки из исходной совокупности прецедентов является использование методов формирования случайных выборок [2 – 6], которые обладают такими недостатками, как необходимость задания объема формируемой выборки человеком (неопределенность объема выборки), возможность невключения важных и включения малозначимых экземпляров в формируемую выборку при малом объеме формируемой выборки.

Другим подходом к формированию выборки является использование процедур кластер-анализа [1, 5], позволяющих выделить все основные типы наблюдений. Однако недостатком данных методов является то, что количество кластеров (типов прецедентов) априори неизвестно, а статистические свойства (частоты экземпляров разных кластеров) в сформированной выборке могут не соответствовать исходной. Кроме того, методы [1, 5] могут выделить чрезмерно большое число кластеров, что приведет к избыточности сформированной выборки.

Наиболее точным методом формирования обучающих выборок, способным гарантированно обеспечить наилучшее решение для заданного

критерия качества, является метод полного перебора [1] всех возможных подвыборок исходной выборки. Однако метод является чрезвычайно трудоемким и для выборок большого объема не применим.

Поэтому для формирования выборок необходимо разработать методы, способные в автоматическом режиме выделять из исходной выборки подмножество экземпляров минимального объема, содержащее наиболее важные экземпляры для построения диагностической модели.

Целью данной работы являлось создание методов, позволяющих автоматизировать процесс формирования обучающих выборок для синтеза диагностических моделей по прецедентам.

Постановка задачи. Пусть задана исходная выборка $\langle X, Y \rangle$ объемом S экземпляров, характеризуемых набором значений N входных (описательных) признаков X и одного выходного (целевого) признака Y . Тогда задачу формирования обучающей выборки $\langle x, y \rangle$ из исходной выборки $\langle X, Y \rangle$ можно представить как поиск такого минимального подмножества $\langle X, Y \rangle$, для которого значение заданного функционала качества $\bar{I}(\langle x, y \rangle)$ будет иметь максимальное значение. При этом функционал качества $\bar{I}(\langle x, y \rangle)$ должен отражать требования относительно топологической и статистической репрезентативности формируемой подвыборки $\langle x, y \rangle$ относительно $\langle X, Y \rangle$.

Метод формирования выборки на основе сокращенного перебора с цензурированием и псевдокластеризацией. Для устранения недостатков метода полного перебора предлагается исходную выборку разделить на подмножества, расположенные в компактных областях пространства признаков – кластерах, и из каждого такого подмножества извлекать только те экземпляры, которые наиболее перспективны для формирования множества решений. По сути, заменить исходную выборку выборкой меньшего размера, содержащей потенциально наиболее ценные экземпляры. Далее из полученной выборки сформировать множество решений, среди которых отобрать наилучшие путем перебора с цензурированием. Для ускорения поиска предлагается кластерный анализ исходной выборки в многомерном пространстве признаков заменить на псевдокластеризацию, представляемую как объединение результатов частных кластеризаций выборки в одномерных проекциях на оси признаков. Разработанный метод включает следующие этапы.

1. Инициализация: задать исходную выборку $\langle X, Y \rangle$ объемом S экземпляров, а также максимально допустимый объем S_{ϕ} формируемой выборки $\langle x, y \rangle$. Рассчитать значение критерия качества исходной выборки $\bar{I}$.

2. Псевдокластеризация. Для каждого i -го признака ($i = 1, 2, \dots, N$) выполнить пп. 2.1 – 2.2.

2.1. Упорядочить экземпляры исходной выборки в порядке неубывания значений i -го признака.

2.2. Просматривая упорядоченное множество экземпляров по оси i -го признака слева направо (от меньших значений к большим) попарно для каждых двух соседних экземпляров, включить оба экземпляра в выборку $\langle X', Y' \rangle$, если они принадлежат к разным классам. Также включить в выборку $\langle X', Y' \rangle$ крайние левый и правый экземпляры по оси значений i -го признака. При включении экземпляров в выборку $\langle X', Y' \rangle$ исключить дуближ, для чего перед добавлением нового экземпляра найти расстояния от него до каждого из уже имеющихся в выборке экземпляров и включать экземпляр только тогда, когда минимальное из расстояний больше нуля.

3. Для сформированной выборки $\langle X', Y' \rangle$ объемом S' сгенерировать все возможные подвыборки $\langle x(k), y(k) \rangle$, содержащие комбинации не более, чем S_Φ экземпляров, $S_\Phi \leq S'$. Здесь k – номер подвыборки, $x(k), y(k)$ – соответственно, экземпляры k -ой выборки и сопоставленные им значения выходного признака.

4. Цензурирование и отбор решений.

4.1. Для $\forall k$ исключить из рассмотрения подвыборки, удовлетворяющие условию: $\exists q, q = 1, 2, \dots, K: S^q(k) = 0$, где K – количество классов, $S^q(k)$ – количество экземпляров q -го класса в k -ой подвыборке.

4.2. Для оставшихся в рассмотрении подвыборок, исключить из рассмотрения те, которые удовлетворяют условию: $\exists q, q = 1, 2, \dots, K: |S^q(k) - S^q| / S > \delta_K$, где δ_K – некоторая заранее заданная константа, $0 < \delta_K < 1$.

4.3. Для всех оставшихся подвыборок рассчитать значения критерия качества $\bar{I}(k)$ и исключить из рассмотрения те из оставшихся подвыборок, для которых: $\bar{I}(k) < \bar{I}^*$, где $\bar{I}^*$ – среднее значение критерия качества для оставшихся подвыборок. Критерий качества выборки можно определить на основе показателей качества, предложенных в [3 – 5].

5. Среди оставшихся подвыборок $\{\langle x(k), y(k) \rangle\}$ в качестве решения $\langle x, y \rangle$ выбрать ту подвыборку $\langle x(p), y(p) \rangle$, которая наилучшим образом соответствует заранее заданному критерию выбора решения. Предлагается использовать один из следующих критериев:

– максимум качества формируемой выборки: $p = \arg \max_k \{\bar{I}(k)\}$;

– максимум соответствия качества формируемой выборки качеству исходной выборки: $p = \arg \min_k \{|\bar{I}(k) - \bar{I}|\}$;

– минимум объема выборки: $p = \arg \min_k \{S(k)\}$;

– максимум ограниченного объема выборки: $p = \arg \max_k \{S(k)\}$;

– комбинированные критерии: $p = \arg \min_k \{|\bar{I}(k) - \bar{I}| / S(k)\}$;

$$p = \arg \min_k \{\bar{I}(k) / S(k)\}; \quad p = \arg \min_k \{S(k) |\bar{I}(k) - \bar{I}|\}; \quad p = \arg \min_k \{S(k) \bar{I}(k)\}.$$

Достоинством данного метода является то, что он перед формированием подвыборок существенно сокращает размерность исходной выборки за счет исключения малоинформативных экземпляров, сохраняя при этом экземпляры, расположенные на границе разделения классов. Таким образом, с одной стороны, существенно сокращает время поиска решений, а, с другой стороны, сохраняет наиболее важные для построения моделей прецеденты. Недостатком метода является то, что он из-за потери информации вследствие сокращения исходной выборки, может существенно изменить частоты классов в извлекаемых выборках. Другим недостатком метода является то, что информация о качестве уже сгенерированных выборок не учитывается при формировании новых выборок.

Эволюционный метод формирования выборок. Для сокращения числа перебираемых комбинаций рационально обеспечить использование информации об уже проанализированных решениях для перехода к рассмотрению новых решений, похожих на рассмотренные ранее. Также необходимо обеспечить шансы для каждого из возможных решений быть рассмотренным. Для этого предлагается использовать эволюционный подход, представляющий собой разновидность случайного поиска [7].

Метод формирования выборки на основе эволюционного поиска с псевдокластеризацией будет включать следующие этапы.

1. Инициализация: задать исходную выборку $\langle X, Y \rangle$ объемом S экземпляров, а также максимально допустимый объем $S_{\text{ф}}$ формируемой выборки $\langle x, y \rangle$. Рассчитать значение критерия качества исходной выборки $\bar{I}$ (критерий качества выборки можно определить на основе показателей качества, предложенных в [7 – 9]). Задать размер популяции решений N , максимальное число итераций T , вероятность мутации $P_{\text{м}}$, а также приемлемое значение критерия качества результата $\bar{I}^* \leq \bar{I}$.

2. Псевдокластеризация. Для каждого i -го признака ($i = 1, 2, \dots, N$) выполнить пп. 2.1 – 2.3.

2.1. Упорядочить экземпляры исходной выборки в порядке неубывания значений i -го признака.

2.2. Просматривая упорядоченное множество экземпляров по оси i -го признаков слева направо (от меньших значений к большим) попарно для каждых двух соседних экземпляров, включить оба экземпляра в выборку $\langle X', Y' \rangle$, если они принадлежат к разным классам. Также включить в выборку $\langle X', Y' \rangle$ крайние левый и правый экземпляры по оси значений i -го признака. При включении экземпляров в выборку $\langle X', Y' \rangle$ исключить дублиаж, для чего перед добавлением нового экземпляра найти расстояния от него до каждого из уже имеющихся в выборке экземпляров и включать экземпляр только тогда, когда минимальное из расстояний больше нуля.

3. Формирование начальной популяции решений. Представим k -ое решение h^k как бинарную комбинацию из S разрядов, s -й разряд которой h_s^k определяет включение в решение s -го экземпляра исходной выборки (если $h_s^k = 0$, то s -й экземпляр не входит в k -ое решение, в противном случае, когда $h_s^k = 1$, s -й экземпляр входит в k -ое решение). Сформируем случайным образом H бинарных комбинаций путем выполнения п. 3.1 – 3.2 для $k = 1, 2, \dots, H$; $s = 1, 2, \dots, S$.

3.1. Задать вероятности включения экземпляров в k -ое решение:

$$P(h_s^k) = \begin{cases} 0,5(\lambda + \text{rand}), & X^s \in \langle X', Y' \rangle; \\ 0,5\text{rand}, & X^s \notin \langle X', Y' \rangle, \end{cases} \quad \lambda = \begin{cases} 1, & S' \leq S_\phi; \\ \min\{0,5; S_\phi / S'\}, & S' > S_\phi, \end{cases}$$

где rand – функция, возвращающая случайное число в диапазоне $[0, 1]$.

3.2. Двигаясь от разрядов с большими вероятностями включения экземпляров в k -е решение к разрядам с меньшими вероятностями, установить равными единице не более S_ϕ разрядов с наибольшими вероятностями, но не меньшими 0,5, остальные разряды установить равными нулю.

4. Проверка на окончание поиска. Для каждого k -го решения популяции сформировать соответствующую выборку, для которой оценить $\bar{I}(k)$. Если выполнено более чем T итераций или среди множества решений имеется такое решение с номером k , для которого $\bar{I}(k) \geq \bar{I}^*$, то прекратить поиск и вернуть в качестве результата выборку с наибольшим значением критерия качества.

5. Отбор решений для скрещивания. Рассматривая $\bar{I}(k)$ в качестве максимизируемой фитнес-функции, сформировать родительские пары для производства решений-потомков на основе правила "колеса рулетки"

[7], обеспечивая тем самым учет $\bar{I}(k)$ для оценивания вероятности решения быть допущенным к скрещиванию.

6. Скрещивание. Реализовать скрещивание отобранных решений для производства новых решений на основе одноточечного кроссингвера.

7. Мутация. Для каждого из имеющихся решений инвертировать случайным образом не более $\text{round}(P_m S)$ разрядов, где round – функция округления. Для тех решений, в которых число битов, равных единице, превышает S_ϕ , инвертировать случайным образом лишние единичные биты. Исключить из текущей популяции решения, встречавшиеся ранее на предыдущих циклах работы метода. Перейти к этапу 4.

Данный метод совмещает идеи случайного формирования выборки и детерминированного поиска лучших решений. Он начинает работу с выделения наиболее перспективных для включения в решения экземпляров, однако сохраняет шансы остальных экземпляров войти в формируемые выборки, и в процессе своей работы целенаправленно улучшает рассматриваемые решения. При этом метод гарантирует, что каждая из рассматриваемых выборок будет иметь объем не более S_ϕ .

Анализ сложности методов формирования выборок. Для разработанных методов формирования выборок представляется целесообразным оценить условия их практической применимости. Очевидно, что основными показателями, определяющими сложность методов формирования выборок, являются число генерируемых подвыборок-комбинаций экземпляров исходной выборки, а также количество вычислений интегрального критерия качества. Для сравнения методов формирования выборок оценим их временную сложность.

Метод формирования выборки на основе полного перебора будет иметь сложность порядка $O(G_F G_g)$, где G_F – сложность расчета интегрального показателя качества для выборки (для выборки из S экземпляров завышенная оценка сложности составит $O(S^2)$ при эффективной реализации вычислений), G_g – количество генерируемых выборок. Для данного метода $G_g = 2^S - 1$. Поскольку генерируемые выборки будут содержать разное число экземпляров, для простоты в среднем примем: $G_F = (0,5S)^2$, таким образом, получим оценку сложности $O((0,5S)^2 (2^S - 1))$. Эта оценка свидетельствует о практической пригодности полного перебора только для исходных выборок небольшого объема.

Метод формирования выборки на основе перебора с цензурированием и псевдокластеризацией имеет сложность порядка $O(G_F G_c + G_g)$, где $G_c \ll G_g$. В худшем случае $G_c = G_g$, а, поскольку, $G_g = 2^S - 1$, то сложность метода можно оценить как $O((2^S - 1)(G_F + 1))$.

Примем приближенно $G_F = S_\Phi^2 \approx (0,5S)^2$, тогда сложность метода может быть оценена как $O((2^S - 1)((0,5S)^2 + 1))$. В наихудшем случае ($S_\Phi = S$) он будет в $(2^S - 1)(0,5S)^2 / ((2^{S'} - 1)((0,5S)^2 + 1)) \approx 2^{(S-S')}$ раз быстрее работать по сравнению с методом полного перебора. Данный метод может применяться для больших выборок, поскольку содержит процедуры устранения избыточных прецедентов перед выполнением поиска, который дополнительно еще и цензурируется.

Для эволюционного метода формирования выборки сложность будет зависеть от количества итераций, которое в худшем случае составит T , размера популяции решений H и сложности расчета показателя качества выборки G_F , который примем приближенно $G_F = (0,5S)^2$. В результате сложность метода приближенно может быть оценена как $O(THG_F) = O(0,25THS^2)$. Таким образом, для данного метода можно определить требования к значениям параметров, обеспечивающим его эффективность относительно:

- полного перебора: $TH \ll 2^S - 1$;

- перебора с цензурированием и псевдокластеризацией $TH \ll 2^{S'} - 1$. Только при соблюдении данных условий, предложенный эволюционный метод будет эффективнее этих переборных методов.

Выводы. С целью автоматизации построения диагностических моделей решена актуальная задача разработки математического обеспечения для формирования обучающих выборок.

Научная новизна работы заключается в том, что получили дальнейшее развитие переборные и эволюционные методы комбинаторного поиска, которые модифицированы для формирования выборок путем введения разработанных критериев для отбора, цензурирования и псевдокластеризации экземпляров, что позволяет ускорить процесс формирования выборок и обеспечить их соответствие заданным критериям при ограниченном объеме.

Практическая ценность результатов работы состоит в определении оценок сложности разработанных методов формирования выборок, позволяющих определить условия их применимости на практике. Использование предложенных оценок делает возможным задание критериев качества, учитывающих предпочтения пользователя при формировании выборок и синтезе моделей с учетом имеющихся ресурсов.

Дальнейшие исследования могут быть направлены на разработку процедур цензурирования решений, обеспечивающих большее сокращение пространства поиска в методах формирования выборок.

Список литературы: 1. Интеллектуальные средства диагностики и прогнозирования надежности авиадвигателей: монография / В.И. Дубровин, С.А. Субботин, А.В. Богуслаев, В.К. Яценко. – Запорожье: ОАО "Мотор-Сич", 2003. – 279 с. 2. Кокрен У. Методы выборочного исследования / У. Кокрен. – М.: Статистика, 1976. – 440 с. 3. Bernard H.R. Social research methods: qualitative and quantitative approaches / H.R. Bernard. – Thousand Oaks: Sage Publications, 2006. – 784 p. 4. Chaudhuri A. Survey sampling theory and methods / A. Chaudhuri, H. Stenger. – New York: Chapman & Hall, 2005. – 416 p. 5. Multivariate analysis, design of experiments, and survey sampling / [ed. S. Ghosh]. – New York: Marcel Dekker Inc., 1999. – 698 p. 6. Прогрессивные технологии моделирования, оптимизации и интеллектуальной автоматизации этапов жизненного цикла авиационных двигателей: монография / А.В. Богуслаев, Ал.А. Олейник, Ан.А. Олейник, Д.В. Павленко, С.А. Субботин. Под ред. Д.В. Павленко, С.А. Субботина. – Запорожье: ОАО "Мотор Сич", 2009. – 468 с. 7. Subbotin S.A. The Training Set Quality Measures for Neural Network Learning / S.A. Subbotin // Optical Memory and Neural Networks (Information Optics). – 2010. – Vol. 19. – № 2. – P. 126–139. 8. Субботин С.А. Комплекс характеристик и критериев сравнения обучающих выборок для решения задач диагностики и распознавания образов / С.А. Субботин // Математичні машини і системи. – 2010. – № 1. – С. 25–39.

Статья представлена д.т.н., проф. НТУ "ХПИ" Дмитриенко В.Д.

УДК 681.518.5:004.93

Методи формування вибірок для побудови діагностичних моделей за прецедентами / Субботін С.О. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 149 – 156.

Вирішено актуальне завдання розробки математичного забезпечення для формування навчальних вибірок. Дістали подальшого розвитку переборні й еволюційні методи комбінаторного пошуку, які модифіковані для формування вибірок шляхом уведення розроблених критеріїв для відбору, цензурування та псевдокластеризації екземплярів, що дозволяє прискорити процес формування вибірок і забезпечити їхню відповідність заданим критеріям при обмеженому обсязі. Визначено оцінки складності розроблених методів. Бібліогр.: 9 назв.

Ключові слова: вибірка, еволюційні методи комбінаторного пошуку, діагностична модель.

UDC 681.518.5:004.93

The sampling methods for diagnostic model construction on precedents / Subbotin S.A. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 149 – 156.

The actual problem of mathematical support development for training sample forming is solved. The exhaustive search and evolutionary methods of combinatorial search were further developed. They are modified for sampling by the introduction of the developed criteria for exemplar selection, censoring and pseudo-clustering. This allows to speed up the process of sampling and to ensure its compliance with specific criteria in limited volume. The complexity of the developed methods is estimated. Refs.: 9 titles.

Key words: sample, evolutionary methods of combinatorial search, diagnostic model.

Поступила в редакцію 08.09.2010

Б.М. СУДАКОВ, к.т.н., проф. НТУ "ХПІ", Харків,

Д.Е. ДВУХГЛАВОВ, к.т.н., доц. ХУПС, Харків,

І.М. ВОЛОДИНА, магістр НТУ "ХПІ", Харків

МЕТОД СТРУКТУРИЗАЦІЇ ПРЕДМЕТНОЇ ГАЛУЗІ В СИСТЕМАХ ПІДТРИМКИ ПРИЙНЯТТЯ РІШЕНЬ

Обґрунтовано доцільність забезпечення можливості розширення меж предметної галузі в системах підтримки прийняття рішень. Представлено переваги використання методу для користувачів, які не мають навиків розробки програмного забезпечення. Розроблено основні принципи структуризації предметної галузі у вигляді системи взаємопов'язаних об'єктів. Іл.: 2. Бібліогр.: 10 назв.

Ключові слова: предметна галузь, система підтримки прийняття рішень, система взаємопов'язаних об'єктів.

Аналіз літератури та постановка проблеми. Результат діяльності людей у більшості галузей на сучасному етапі визначається ефективністю рішень, що приймаються. Це стосується задач аналізу поточної ситуації, в яких необхідно точно визначити тип ситуації для вирішення подальших дій. Також це відноситься до задач прогнозування результатів дій, що дозволить оцінити власну вигоду та затрати; і задач оперативного управління, коли від швидкості реакції на зміни в обстановці залежить загальний успіх [1, 2]. В сучасних умовах прийняття ефективних рішень у встановлені терміни в більшості галузей діяльності людини вимагає використання систем підтримки прийняття рішень (СППР) [3, 4].

В роботі [5] показано, що розробка програмного забезпечення СППР здійснюється з використанням інформаційного або когнітивного підходу. Простота та відносно невелика вартість створення систем обробки даних робить інформаційний підхід більш популярним на теперішній час. Питання розробки баз даних розглядаються у великій кількості наукових та навчальних публікацій [6 – 9]. Але створення системи обробки даних з часом вимагає зміни структури бази даних, що, в свою чергу, потребує модифікації програмних модулів наповнення. Таким чином, використання існуючої технології розробки систем обробки даних передбачає залучення тільки фахівців, які володіють спеціальними знаннями та навичками. Можливість розширення переліку інформації, що зберігається у базах даних, користувачами системи виключена. Але з огляду на те, що достатньо часто розробник і користувач знаходяться на відстані один від одного, тривалість модифікації структури бази даних та програмних модулів маніпулювання даними робить актуальним забезпечення можливості внесення змін користувачами системи.

Мета статті. Розробити метод структуризації предметної галузі, який дозволить забезпечити можливість розширення складу даних, що накопичуються в системі, користувачами без участі програмістів.

Основна частина. Сучасні системи баз даних мають в своїй основі реляційну модель даних – систему взаємопов’язаних таблиць відношень [9]. У випадку зміни структури бази даних буде або додане поле (поля) у певну таблицю, або додана нова таблиця. У першому випадку на форму для наповнення модифікованої таблиці має бути додано новий компонент форми для вводу та відображення значень цього поля, а також перероблені процедура вводу даних та модифікації даних. У другому випадку у проекті з’явиться 2-3 нових форми. І якщо змінювати структуру користувач ще може навчитись, то для розробки та переробки модулів потрібне залучення відповідних спеціалістів.

Таким чином, щоб забезпечити можливість внесення змін до структури бази даних без залучення програмістів та без розробки програмних модулів наповнення слід уніфікувати форму для представлення інформації та забезпечити її динамічне формування в залежності від наповнення. Для динамічного конструювання форми в сучасних системах програмування є відповідні компоненти [10].

Сутність підходу до структуризації базується на тому, що предметна галузь являє собою сукупність об’єктів. Кожен об’єкт описує набір характеристик, значення яких формують чітку уяву про нього. Експерти мають визначити типи класів для структуризації предметної галузі. Такий розподіл дозволить організувати характеристики у групованому вигляді та надати у певній послідовності за вимогами користувача. Принцип представлення типів об’єктів демонструє рис. 1. Формально тип об’єктів TOb_i можна описати так

$$TOb_i = \langle NTOb_i, RATO b^i \rangle, \quad (1)$$

де $NTOb_i$ – найменування типу; а $RATO b^i$ – множина описів характеристик, що пов’язуються із даним типом

$$RATO b^i = \{RA_m^{ij}\}. \quad (2)$$

Кожен елемент множини $RATO b^i$ відповідає встановленню зв’язку між типом об’єкту та відповідною характеристикою:

$$RA_m^{ij} \Leftrightarrow A_m \mapsto TOb_i. \quad (3)$$

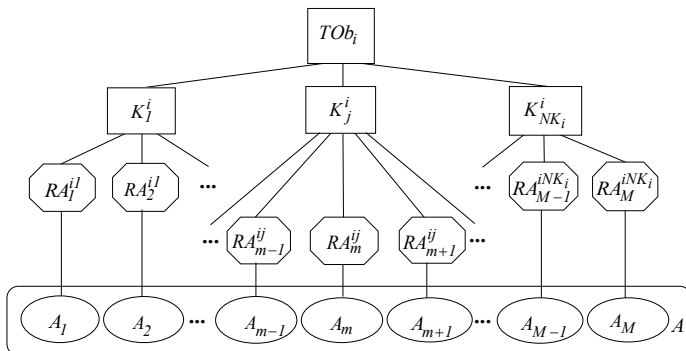


Рис. 1. Структура представлення типового об'єкту предметної галузі

При цьому передбачається, що існує набір характеристик $A = \{A_1, A_2, \dots, A_{m-1}, A_m, A_{m+1}, \dots, A_M\}$, які можуть бути використані для опису будь-якого типу об'єктів. Також, при встановленні зв'язку "тип об'єкту – характеристика" визначається клас K_j^i ($K_j^i \in K^i$, K^i – сукупність класів i -го типу, визначених для змістовної структуризації характеристик), до якого буде включена характеристика, а також визначаються обмеження на значення характеристики PA_m^{ij} . Отже:

$$RA_m^{ij} = \langle A_m, K_j^i, PA_m^{ij} \rangle. \quad (4)$$

Таким чином, характеристика включається тільки до одного класу K_j^i . Перелік обмежень PA_m^{ij} залежить від типів значень характеристики і відрізняється при використанні характеристики для опису об'єктів різних типів. Сфера використання способу дозволяє визначити, що характеристика може зберігати значення одного з шести типів: цілі значення (не допускається дрібна частина); числові (допускається дрібна частина); символічні; логічні (істина або хибне); дата та час. Для кожної характеристики встановлюються обмеження на припустимі значення. Якщо тип характеристики встановлений "перерахована", то для неї попередньо визначається набір можливих значень, з яких можна обрати потрібний варіант. Для не перерахованих характеристик значення є довільними. В залежності від типу значень можна встановити додаткові параметри. Так, спосіб передбачає, що логічна величина – це перерахована величина із двома можливими значеннями – "так" та "ні". Для символічних не перерахованих значень треба встановити максимальну довжину значення. Для інших типів, за умов прийняття довільного значення, слід встановити максимальне та мінімальне

значення, які будуть визначати діапазон можливих значень. На основі визначених типів об'єктів далі будуть створювати екземпляри об'єктів, які відповідатимуть певним об'єктам предметної галузі (див. рис. 2).

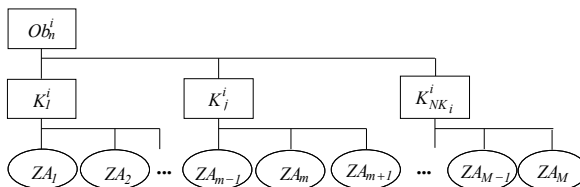


Рис. 2. Структура екземпляру об'єкту

Це характеризуватиме перехід від загального опису структури предметної галузі до детального. Визначена структура конкретного об'єкту Ob_n^i вже містить значення характеристик ZA_m . Рис. 2 ілюструє, як буде представлена інформація користувачу.

Висновки. Сукупність типів об'єктів та їх характеристик дозволяють представити інформацію про об'єкти предметної галузі будь-якої природи. Використання даного підходу має певні переваги у порівнянні із традиційними методами представлення інформації, а саме:

1. Створення моделі предметної галузі має схожу послідовність із створенням структури таблиць бази даних. Але запропонований спосіб забезпечує створення універсальної форми для введення інформації про об'єкти, налагодження вигляду представлення інформації для кінцевих користувачів. Таким чином, спосіб забезпечує скорочення часу на розробку програмних засобів введення та модифікації інформації про предметну галузь у випадку розширення її меж.

2. Оперування поняттями "об'єкт" та "характеристика" є більш зрозумілими для кінцевого користувача системи. Тому набути навички створення моделі предметної галузі користувачі зможуть швидше, ніж вносити зміни у структуру бази даних. Також слід врахувати, що створення універсальної форми для введення інформації дозволить користувачу відразу після створення типу об'єкту та визначення характеристик вносити до пам'яті СППР відомості про ці об'єкти.

3. В сфері розробки систем підтримки прийняття рішень має місце тенденція нарощування можливостей систем у напрямку інтелектуалізації. Розроблена методика дозволяє забезпечити як звичайне зберігання і узагальнення інформації, так і обробку з використанням інтелектуальних технологій.

Список літератури: 1. Карданская Н.Л. Управленческие решения: Учебник для вузов / Н.Л. Карданская. – М.: ЮНИТИ ДАНА, 2004. – 465 с. 2. Катренко А.В. Теорія прийняття рішень: підручник з грифом МОН / А.В. Катренко, В.В. Пасічник, В.П. Пасько. – К.: Видавнича група ВНУ, 2009. – 448 с. 3. Ситник В.Ф. Системи підтримки прийняття рішень: Навч. посіб. / В.Ф. Ситник. – К.: КНЕУ, 2004. – 614 с. 4. Герасимов Б.М. Системы поддержки принятия решений: проектирование, применение, оценка эффективности / Б.М. Герасимов, М.М. Дивизинюк, И.Ю. Субач. – Севастополь: Изд.центр СНИЯЭИП, 2004. – 320 с. 5. Теоретичні основи автоматизації процесів вироблення рішень в системах управління Повітряних Сил / О.В. Александров, Д.Е. Двухглавов, М.А. Павленко, І.О. Романенко, О.І. Тимочко. – Х.: ХУ ПС, 2010. – 172 с. 6. Карпова Т.С. База данных: модели, разработка, реализация / Т.С. Карпова. – СПб.: Питер, 2002. – 304 с. 7. Коннолли Т. Базы данных. Проектирование, реализация и сопровождение. Теория и практика: Пер. с англ. / Т. Коннолли, К. Бегг. – М.: Вильямс, 2003. 8. Роланд Ф.Д. Основные концепции баз данных: Пер. с англ. / Ф.Д. Роланд. – М.: Издательский дом "Вильямс", 2002. – 256 с. 9. ДСТУ 28764-94. Системи обробки інформації. Бази і банки даних. Основні терміни та визначення. – К.: Держстандарт, 1994. – 32 с. 10. Фленов М.Е. Библия Delphi / М.Е. Фленов. – СПб.: БХВ-Петербург, 2004. – 1026 с.

Стаття представлена д.т.н. проф. НТУ "ХПІ" Серковим О.А.

УДК 004.048

Метод структуризации предметной области в системах поддержки принятия решений / Судаков Б.Н., Двухглавов Д.Э., Володина И.М. // Вестник НТУ "ХПИ". Тематический выпуск: Информатика и моделирование. — Харьков: НТУ "ХПИ". – 2011. – № 17. – С. 157 – 161.

Обоснована целесообразность обеспечения возможности расширения границ предметной области в системах поддержки принятия решений. Представлены преимущества использования метода для пользователей, которые не имеют навыков разработки программного обеспечения. Предлагаются основные принципы структуризации предметной области в виде системы взаимосвязанных объектов. Ил.: 2. Библиогр.: 10 назв.

Ключевые слова: предметная область, система поддержки принятия решений, система взаимосвязанных объектов.

UDC 004.048

Method of problem area structurization in decision support systems / Sudakov B.N., Dvukhglavov D.E., Volodina I.M. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – № 17. – P. 157 – 161.

The expediency of a problem area borders expansion possibility in decision support systems is proved. Advantages of method usage by system users without the working out software skills are shown. The main principles of the problem area composition in the form of the interconnected objects system are suggested. Figs.: 2. Refs.: 10 titles.

Key words: problem area, decision support system, system of interconnected objects.

Поступила в редакцію 04.01.2011

А.В. ТЕЛІШЕВСЬКА, асистент, ЧФ НТУ "ХПІ", Чернівці,
А.І. ПОВОРОЗНЮК, к.т.н., доц. НТУ "ХПІ", Харків

ФОРМАЛІЗАЦІЯ ВХІДНОЇ ІНФОРМАЦІЇ ДЛЯ ДІАГНОСТИКИ НЕВРОЛОГІЧНИХ ЗАХВОРЮВАНЬ

В роботі розглядається проблема постановки діагнозу неврологічного захворювання. В якості вхідних ознак розглянуті лабораторні дані і нейровізуальні дослідження, а також неврологічний статус. В результаті побудована множина, що містить підмножини ознак по лабораторним та клінічним дослідженням, які проведені в Кам'янець-Подільській міській лікарні № 1. Табл.: 1. Бібліогр.: 8 назв.

Ключові слова: діагностика, ознака, формалізація, неврологічні захворювання, клінічні дослідження.

Постановка проблеми та аналіз літератури. За останні десять років відмічається значне зростання неврологічних захворювань (НЗ) в Україні: захворюваність на нервові хвороби, а також поширеність їх збільшились майже вдвічі [1]. Інсульт, або, як його називали раніше, мозковий удар, обумовлений раптовим припиненням кровопостачання частини головного мозку або крововиливом у порожнину черепа та стисненням тканини мозку кров'ю, що вилилась. Обидві ці причини призводять до припинення функціонування та до загибелі клітин мозку в ушкодженій ділянці, що призводить до порушення або втрати функції тих частин тіла, якими вони керують. В Україні захворюваність на інсульт складає приблизно 120 тис. випадків на рік. Серед наших співвітчизників, які перенесли інсульт, відносна кількість загиблих або таких що залишилися інвалідами, суттєво вища, ніж в цивілізованих країнах. На відміну від розвинених країн, в Україні інсульт як причина смерті посідає не третє, а друге місце, поступаючись лише захворюванням серця. Витрати на лікування хворих на інсульт або його наслідки в Україні набагато нижчі, ніж в розвинених країнах, та й навіть в багатьох країнах, що розвиваються. На сьогоднішній день є чітке уявлення про НЗ [2, 3]. Вони зумовлені навколишніми чинниками, такими, як забрудненість навколишнього середовища та шкідливі звички, та внутрішніми чинниками: високим кров'яним тиском, холестерином, тромбозом та ін. Тому для діагностики необхідні дослідження, що дозволяють оцінити збереженість чи порушення кровоплину та визначити місця крововиливів (ультразвукове дослідження сонних артерій, ангіографія – рентгенологічне обстеження судин). Дослідження крові для визначення порушень здатності крові до згортання, що сприяють чи утворенню тромбів, чи кровоточивості. Електрокардіографія (ЕКГ) або ультразвукове дослідження серця (ехокардіографія) для

виявлення серцевих джерел тромбів, які можуть, відірвавшись, пересуватися в судини головного мозку [4]. Таким чином, для діагностики НЗ необхідно виконати аналіз різнорідних діагностичних ознак (дихотомічних, числових та візуальних), тому необхідно формалізувати множину діагностичних ознак. При цьому необхідно вирішити задачі формального опису ознак [5, 6] та відновлення пропущених даних [7, 8].

Ціллю статті є аналіз та формалізація вхідних даних для діагностування неврологічних захворювань.

Формалізація вхідних даних. Вхідні дані для діагностування НЗ представлені Кам'янець-Подільською міською лікарнею № 1. Для аналізу були відібрані 200 пацієнтів з підозрою на НЗ. В ході комплексного обстеження встановлені неврологічні порушення, а саме геморагічний інсульт, ішемічний інсульт, епілепсія, остеохондроз, енцефаліт та мігрень. Особливістю вхідних даних являється наявність великої кількості різнорідної інформації. При цьому багато ознак мають описовий характер.

Нехай кожен пацієнт представляє собою об'єкт ω_i ($i = \overline{1, N}$, N – кількість пацієнтів) в багатовимірному просторі ознак. Простір ознак являє собою множину X . Із елементів простору формується вектор ознак $\vec{x} = (x_1, x_2, \dots, x_m)$. Отже, кожен об'єкт ω_i в просторі ознак описується вектором $\vec{x}_i^\omega = (x_{i1}, x_{i2}, \dots, x_{im})$. Задачею діагностики є віднесення пацієнта ω_i до одного з формалізованих станів НЗ D_i^j ($i = \overline{1, 6}$), де D_1^1 – ішемічний інсульт, D_1^2 – геморагічний інсульт, D_1^3 – епілепсія, D_1^4 – остеохондроз, D_1^5 – енцефаліт, D_1^6 – мігрень.

Далі необхідно розбити вхідну множину ознак X на підмножини X_k , які не перетинаються, так, щоб $\bigcup_{i=1}^k X_i = X$, $X_j \cap X_i = \emptyset$, $j, i = \overline{1, k}$, $j \neq i$. В результаті були виділені наступні підмножини: X_1 – неврологічний статус, X_2 – лабораторні дослідження та X_3 – нейровізуальні дослідження. Кожна підмножина, в свою чергу, розбивається на підмножини X_k^l . Підмножина X_1 розбита на підмножини: X_1^1 – черепномозкові нерви (12 пар), X_1^2 – рухова сфера (5 ознак), X_1^3 – менінгальні контрактури (2 ознаки), X_1^4 – чутливість

(7 ознак), X_1^5 – сухожилкові та періостальні рефлексі (5 ознак), X_1^6 – патологічні ознаки (6 ознак), X_1^7 – координація рухів (8 ознак), X_1^8 – вегетативна нервова система (6 ознак). Підмножина X_2 розбита на підмножини: X_2^1 – аналіз крові на глюкозу, X_2^2 – загальний аналіз сечі, X_2^3 – загальний аналіз крові, X_2^4 – печінкові проби, X_2^5 – копрограма, X_2^6 – аналіз сечі на амілазу, X_2^7 – аналіз сечі на діастазу, X_2^8 – аналіз сечі на глюкозу. Підмножина X_3 , в свою чергу, розбита на підмножини: X_3^1 – електрокардіограма, X_3^2 – ультразвукова діагностика, X_3^3 – ехоенцефалоскопія, X_3^4 – магніто-резонансна томографія головного мозку, X_3^5 – комп'ютерна томографія.

Ознаки, які входять в підмножину X_1^i ($i = \overline{1, 8}$), вимірюються в дихотомічній шкалі, тому 1 – присутність ознаки, 0 – відсутність. Значення показників підмножини X_2^i ($i = \overline{1, 8}$) вимірюються в кількісній шкалі. Значення показників підмножини X_3^i ($i = \overline{1, 5}$) мають візуальне представлення. Таким чином, в якості експериментальних даних для діагностики НЗ необхідно сформуувати таблицю наступної структури:

Таблиця. Експериментальні дані

$X \backslash N$	X_1								X_2			X_3				
	X_1^{1j}				...	X_1^{8j}				X_2^{1j}	...	X_2^{8j}	X_3^{1j}	...	X_3^{5j}	
	$x_1^{1j,1}$	...	$x_1^{1j,12}$	...		$x_1^{8j,1}$	...	$x_1^{8j,6}$								
1	$x_1^{11,1}$	...	$x_1^{11,12}$	...		$x_1^{11,1}$	...	$x_1^{11,1}$		x_2^{11}	...	x_2^{81}	x_3^{11}	...	x_3^{51}	
2	$x_1^{12,1}$	...	$x_1^{12,12}$	...		$x_1^{12,1}$	...	$x_1^{12,1}$		x_2^{12}	...	x_2^{82}	x_3^{12}	...	x_3^{52}	
3	$x_1^{13,1}$	...	$x_1^{13,12}$	...		$x_1^{13,1}$	...	$x_1^{13,1}$		x_2^{13}	...	x_2^{83}	x_3^{13}	...	x_3^{53}	
⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮
j	$x_1^{1j,1}$	...	$x_1^{1j,12}$	...		$x_1^{8j,1}$	...	$x_1^{8j,6}$		x_2^{1j}	...	x_2^{8j}	x_3^{1j}	...	x_3^{5j}	
⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮	⋮
N	$x_1^{1n,1}$	...	$x_1^{1n,12}$	...		$x_1^{8n,1}$	...	$x_1^{8n,6}$		x_2^{1n}	...	x_2^{8n}	x_3^{1n}	...	x_3^{5n}	

Для комп'ютерної обробки експериментальних даних необхідно, щоб всі ознаки x_k^{lj} були виражені в числовому представленні. Отже множина X_1 представлена в дихотомічній шкалі, то x_1^{lj} будуть набувати значень 1 або 0, тобто присутність ознаки або відсутність. Множина X_2 представлена в номінальній шкалі, тобто ознаки x_2^{lj} мають цифрове представлення, яке відповідає нормам аналізів пацієнтів. Оскільки множина X_2 має числові дані, то необхідно виконати центрування і нормування. Першим етапом, як правило є центрування – знаходження точки середнього значення всіх ознак – геометричного центра багатовимірної множини точок даних. Зазвичай зручно зсунути всі точки даних на один і той же вектор таким чином, щоб центр множини опинився на початку координат. Наступний етап – це нормування – тобто ділення всіх значень ознак на певне число таким чином, щоб значення ознак потрапляли в подібні по величині інтервали. В якості такого числа, зазвичай вибирають одну із характерних відстаней. В багатовимірній множині існує декілька відстаней. Перша – це середньоквадратичне відхилення:

$$\sigma = \sqrt{\frac{1}{N} \sum_{l=1}^N (X_2^l - \bar{X})^2},$$

де $\bar{X} = \frac{1}{N} \sum_{l=1}^N X_2^l$ (X_2^l – вектор даних, x_2^{lj} – j -та координата l -го вектора).

У випадку, якщо вибірка може вважатися отриманою із нормального розподілу, то в колі з центром в $\bar{X}$ та радіусом σ знаходиться близько двох третин від числа точок даних. Існує відстань, яка характеризує максимальне розсіювання в множині даних

$$R = \max_{l=1, N} \|X_2^l - \bar{X}\|.$$

Нормування всіх ознак на R призводить до того, що вся множина даних поміщена в коло одиничного радіусу.

Якщо в якості відстані вибрані σ або R , то відповідні формули обробки (нормування на "одиничну дисперсію" і "на одиничне коло") мають вигляд:

$$\tilde{X}_2^l = \frac{X_2^l - \bar{X}}{\sigma}, \quad \tilde{X}_2^l = \frac{X_2^l - \bar{X}}{R},$$

де $\tilde{X}_2^l$ – новий вектор ознак, X_2^l – старий вектор ознак.

Крім того, якщо діапазони значень для різних ознак сильно відрізняються один від одного, то краще для кожної з ознак застосувати власний масштаб. Тобто, для кожної з ознак можна ввести своє середньоквадратичне відхилення та розсіювання:

$$\sigma_j = \sqrt{\frac{1}{N} \sum_{l=1}^N (x_2^{lj} - \bar{x}_j)^2}, \quad R_j = \max_{l=1, N} \|x_2^{lj} - \bar{x}_j\|,$$

де $\bar{x}_j = \frac{1}{N} \sum_{l=1}^N x_2^{lj}$ (x_2^{lj} – значення j -ої ознаки на l -му векторі).

Як результат отримаємо формули для нормування на "одиничну дисперсію для кожної ознаки" і "на одиничний куб":

$$\tilde{x}_2^{lj} = \frac{x_2^{lj} - \bar{x}_j}{\sigma_j}, \quad \tilde{x}_2^{lj} = \frac{x_2^{lj} - \bar{x}_j}{R_j}.$$

Оскільки, множина X_3 має візуальне представлення, тому для подальшої роботи ознаки x_3^{lj} необхідно перевести у числове представлення.

Висновки. В даній роботі виконані етапи збору, формалізації і попереднього аналізу вхідних даних та вихідних діагностуємих станів. Побудовані множини, що містять підмножини ознак неврологічного статусу та клінічних досліджень з метою подальшої розробки інформаційної структури медичної бази даних.

Список літератури. 1. *Волошин П.В.* Епідеміологія мозкового інсульту в Україні / *П.В. Волошин, Т.С. Міценко, І.В. Здесенко* // Матеріали науково-практичної конференції з міжнародною участю "Діагностика, лікування, профілактика гострих та хронічних порушень мозкового кровообігу". – Харків. – 2005. – С. 74. 2. *Міценко Т.С.* Стан та перспективи розвитку неврологічної служби в Україні / *Т.С. Міценко* // Матеріали науково-практичної конференції "Фармакотерапія захворювань нервової системи". – Харків. – 2005. – С. 167. 3. *Віничук С.М.* Нервові хвороби. / *С.М. Віничук, Є.Г. Дубенко, Є.Л. Мачерет*. За ред. С.М. Віничука, Є.Г. Дубенка. – К.: Здоров'я, 2001. – 696 с. 4. *Кареліна Т.І.* Медсестринство в неврології. Підручник / *Т.І. Кареліна, Н.М. Касевич*. — К.: ВСВ "Медицина", 2010. – 296 с. 5. *Орлов А.И.* Прикладная статистика. Учебник / *А.И. Орлов*. – М.: Издательство "Экзамен", 2004. – 656 с. 6. *Джоррратано Джозеф.* Экспертные системы: принципы разработки и программирование / *Джозеф Джоррратано, Гари Райли*. – М.: ООО "И. Д. Вильямс", 2007. – 1152 с. 7. *Айвазян С.А.* Классификация многомерных наблюдений / *С.А. Айвазян, З.И. Бежаева, О.В. Староверов*. – М.: Статистика, 1974. – 240 с. 8. *Дейвисон М.* Многомерное шкалирование: Методы наглядного представления данных / *М. Дейвисон*. – М.: Финансы и статистика, 1988.

Статья представлена д.т.н., с.н.с. ДУЭП Тараненко Ю.К.

УДК 681.513:620.1

Формализация входной информации для диагностики неврологических заболеваний / Телишевская А.В., Поворознюк А.И. // Вестник НТУ "ХПИ". Тематический выпуск: Информатика и моделирование. — Харьков: НТУ "ХПИ". — 2011. — № 17. — С. 162 – 167.

В работе рассматривается проблема постановки диагноза неврологического заболевания. В качестве входных признаков рассмотрены данные лабораторных и нейровизуальных исследований, а также неврологический статус. В результате построено множество, которое содержит подмножества признаков по лабораторным и клиническим исследованиям, которые проведены в Каменец-Подольской городской больнице № 1. Табл.: 1. Библиогр.: 8 назв.

Ключевые слова: диагностика, признак, формализация, неврологические заболевания, клинические исследования.

UDC 681.513:620.1

Formalization of entrance information is for diagnostics of neurological diseases / Telishevskaya A.V., Povoroznyuk A.I. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. — Kharkov: NTU "KhPI". — 2011. — №. 17. — P. 162 – 167.

The problem of raising of diagnosis of neurological disease is considered in work. As source signs were considered the laboratory and neurovisual research, and also neurological status. As a result built plural, which contains subsets of signs for by laboratory and clinical research, what are conducted in Kamenec-Podol'skoy a city hospital № 1. Tabl.: 1. Refs.: 8 titles.

Key words: diagnostics, sign, formalization, neurological diseases, clinical researches.

Поступила в редакцию 29.03.2011

А.Е. ФИЛАТОВА, к.т.н., доц. НТУ "ХПИ", Харьков

НЕЛИНЕЙНАЯ ФИЛЬТРАЦИЯ БИМЕДИЦИНСКИХ СИГНАЛОВ С ЛОКАЛЬНО СОСРЕДОТОЧЕННЫМИ ПРИЗНАКАМИ В ЗАДАЧЕ СТРУКТУРНОЙ ИДЕНТИФИКАЦИИ

В работе рассматривается задача структурной идентификации биомедицинских сигналов (БС) с локально сосредоточенными признаками (ЛСП) с помощью нелинейного фильтра. Рассмотрены модели полезного сигнала с ЛСП, а также методы преобразования БС с ЛСП на основе моделей полезного сигнала. Предложен обобщенный метод нелинейной фильтрации БС с ЛСП. Ил.: 3. Библиогр.: 8 назв.

Ключевые слова: структурная идентификация, биомедицинский сигнал, локально сосредоточенные признаки, нелинейный фильтр, модель полезного сигнала.

Постановка проблемы. Проектирование компьютерных диагностических систем (КДС) оценки состояния сердца и сердечно-сосудистой системы с целью повышения эффективности диагностики сердечно-сосудистых заболеваний является актуальной задачей медицинской кибернетики. В [1, 2] были выделены основные этапы обработки БС с ЛСП в КДС. Рассматриваемые БС с ЛСП – это квазипериодические сигналы сложной формы, состоящие из периодически чередующихся структурных элементов (СЭ). СЭ – это небольшие фрагменты интервала наблюдения БС с ЛСП, несущие информацию о состоянии объекта. В качестве примера такого сигнала можно привести ЭКГ. Традиционно диагностическими признаками ЭКГ являются амплитудно-временные параметры СЭ. Например, основные признаки ишемической болезни сердца сосредоточены на S-T интервале и зубце Т [3 – 6]. Одним из ответственных и трудно формализуемых этапов обработки БС с ЛСП является этап структурной идентификации (СИ), которая заключается в выделении на фоне помех СЭ. Эту задачу можно решить с помощью нелинейного фильтра (НФ), в основу которого положена модель полезного сигнала (МПС). Задача НФ – на основе множества моделей структурных элементов БС с ЛСП найти некоторое преобразование, в результате которого может быть получен сигнал, обладающий заданными характеристиками.

Анализ литературы. В [1, 4 – 8] рассмотрены различные модели полезного сигнала, каждая из которых в соответствии с некоторым критерием наилучшим образом описывает СЭ.

МПС-1: представление БС $x(t)$ в виде решетчатой функции времени

идентификации БС с ЛСП относятся к методам контурного анализа. В основе методов данной группы лежат априорные сведения о структуре обрабатываемого сигнала. Поэтому основными недостатками таких методов являются их сложность и плохая адаптация для сигналов, имеющих нетипичную структуру [2 – 4].

Преобразование Фурье лежит в основе спектрального анализа БС с ЛСП. Преобразование Фурье не позволяет локализовать во времени частотные компоненты, поэтому может использоваться для стационарных сигналов. Однако БС с ЛСП не являются стационарными. Попыткой учесть временную характеристику сигнала является оконное преобразование Фурье. При этом ухудшается разрешающая способность по частоте. При увеличении ширины окна улучшается разрешающая способность по частоте, но теряется разрешение по времени, и наоборот. Для преодоления этого недостатка в последнее время используется вейвлет преобразование. Основным недостатком этого метода является вычислительная сложность алгоритмов преобразования. Кроме того, в МПС-2, которые лежат в основе этих преобразований, не учитываются особенности БС с ЛСП [7, 8].

Особый интерес для решения поставленной задачи имеют преобразования, основанные на МПС-4 и МПС-5.

Цель статьи – разработка обобщенного метода нелинейной фильтрации БС с ЛСП с учетом моделей и методов преобразования полезного сигнала БС с ЛСП для решения задачи структурной идентификации при проектировании компьютерных диагностических кардиологических систем.

Структурная идентификация сигнала с помощью метода преобразования БС с ЛСП в фазовое пространство. Пусть имеется выборка искаженных дискретных реализаций $x_1[k], \dots, x_M[k]$ ($M \geq 2$). Каждая выборка соответствует одному циклу квазипериодического БС с ЛСП. Методами численного дифференцирования получены производные $\dot{x}_m[k]$ в каждой k -й точке m -й реализации ($m = \overline{1, M}$). В результате получаем M последовательностей (фазовых траекторий) $Q_m = \{z_m[k], k = \overline{1, K_m}\}$, $m = \overline{1, M}$, где $z_m[k] = (x_m^*[k], \dot{x}_m^*[k])$ – нормированный вектор, координаты $x_m^*[k]$, $\dot{x}_m^*[k]$ которого вычисляются по следующим выражениям:

$$x_m^*[k] = \frac{x_m[k] - \min_k x_m[k]}{\max_k x_m[k] - \min_k x_m[k]}, \quad \dot{x}_m^*[k] = \frac{\dot{x}_m[k] - \min_k \dot{x}_m[k]}{\max_k \dot{x}_m[k] - \min_k \dot{x}_m[k]}.$$

После преобразования сигнала в фазовое пространство фазовые траектории Q_m по соответствующему алгоритму усредняются, и выделяется опорная траектория. По эталонной траектории строится эталонный цикл во временной области. В результате структурная идентификация выполняется не на исходном сигнале, а на полученном известными эвристическими методами эталонном фрагменте. Если обрабатываемый сигнал имеет нетипичные циклы, то в фазовом пространстве выделяется не одна, а несколько опорных траекторий, по каждой из которых строится эталонный сигнал во временной области.

Структурная идентификация сигнала с помощью метода преобразования в адаптированное пространство параметров аппроксимирующих функций (АППАФ). Пусть каждый объект ω_m описан вектором амплитуд $\vec{x}^m = (x_1^m, \dots, x_j^m, \dots, x_{N_x}^m)$ точек исходного сигнала x_t , принадлежащих этому СЭ, где N_x – длина $\vec{x}^m$. Тогда в АППАФ объект ω_m будет описан вектором $\vec{y}^m = (y_1^m, \dots, y_j^m, \dots, y_{N_y}^m)$, где $y_j^m = f_k(\vec{x}_j^m)$, $f_k(\vec{x}_j^m)$ – значение опорной функции, заданной на участке, ограниченном опорными точками; N_y – длина $\vec{y}^m$, причем $N_y < N_x$ (рис. 1).

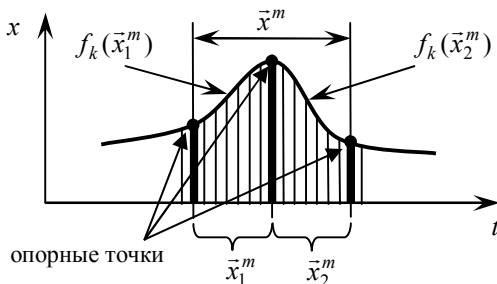


Рис. 1. Пример описания СЭ в АППАФ

В качестве опорных функций могут выступать разделенные разности первого и второго порядков, коэффициенты аппроксимирующих полиномов различных порядков и т. д. Адаптивным пространство является в силу того, что количество координат этого пространства, вид опорной функции и другие параметры преобразования выбираются на основании выбранного класса СЭ. Т.е. параметры преобразования как бы "подстраиваются" под искомые СЭ.

Для СИ сигнала строится функция дифференциации расстояний вида (ФДР) $f_r[t] = d^{(\alpha)}(\omega^3, \omega_t)$, где $d^{(\alpha)} \in [0; 1]$ – расстояние между эталоном ω^3 и объектом ω_t в АППАФ, основанное на идее потенциальных функций. Расстояние вычисляется по следующему выражению

$$d^{(\alpha)}(\omega^3, \omega_t) = \sqrt{\alpha \sum_{j=1}^{N_y} (y_j^3 - y_j^t)^2 \left/ \left(1 + \alpha \sum_{j=1}^{N_y} (y_j^3 - y_j^t)^2 \right) \right.},$$

где y_j^3, y_j^t – координаты в АППАФ объектов ω^3 и ω_t соответственно.

Принадлежность объекта ω_t к классам Ω_1 или Ω_2 определяется из анализа значений локальных минимумов M_ξ ФДР:

$$\omega_t \in \begin{cases} \Omega_1, & \text{если } \theta(\omega_t, P_d) = P_d - M_\xi > 0; \\ \Omega_2, & \text{если } \theta(\omega_t, P_d) = P_d - M_\xi \leq 0, \end{cases}$$

где P_d – адаптивный порог, который может учитывать как статические, так и динамические значения параметров сигнала.

Обобщенный метод нелинейной фильтрации. Анализ рассмотренных методов позволил выделить однотипные этапы при структурной идентификации БС с ЛСП. В результате предложена схема структурной идентификации, изображенная на рис. 2.

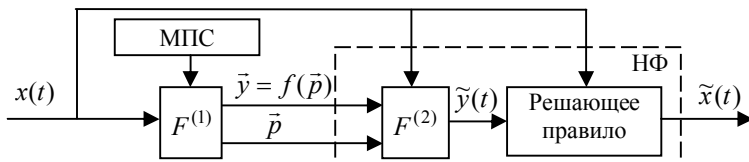


Рис. 2. Схема структурной идентификации БС с ЛСП на основе нелинейного фильтра

Преобразование 1-го уровня $F^{(1)}$ – это один из методов преобразования БС с ЛСП $x(t)$ на основе моделей полезного сигнала с учетом вектора параметров $\vec{p}$ в вектор $\vec{y} = f(\vec{p})$. Например, переход в фазовое пространство или адаптивное пространство параметров аппроксимирующих функций.

Преобразование 2-го уровня $F^{(2)}$ – получение новой функции во временной области $\tilde{y}(t)$, на основании анализа которой с помощью соответствующего решающего правила выполняется структурная идентификация. В первом из рассмотренных методов в качестве преобразования 2-го уровня выступает этап синтеза эталонного фрагмента сигнала; во втором методе – построение функции дифференциации расстояний. Преобразование $F^{(2)}$ вместе с решающим правилом являются основой нелинейного фильтра.

Т.к. адекватность определенной МПС для каждого СЭ различна, то при проектировании НФ предлагается объединять частные решающие правила (ЧРП) в коллектив решающих правил (КРП) (рис. 3).

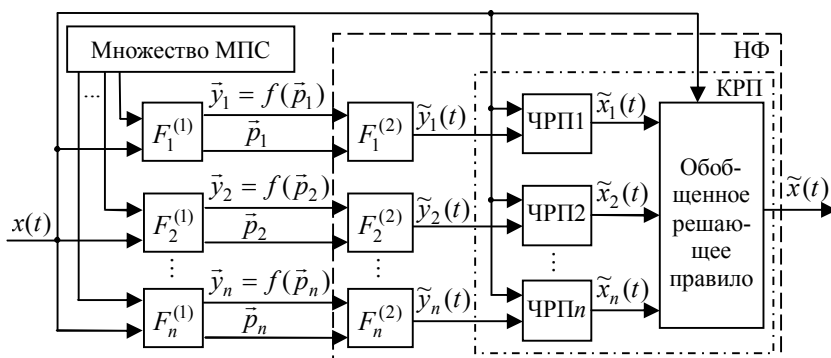


Рис. 3. Обобщенная схема структурной идентификации БС с ЛСП на основе нелинейного фильтра

Такой подход позволяет учитывать все сильные стороны преобразований, основанных на различных МПС. А это, в свою очередь, повышает качество структурной идентификации БС с ЛСП.

Выводы. Анализ моделей полезного сигнала и методов преобразования БС с ЛСП позволил выделить однотипные этапы при структурной идентификации БС с ЛСП, в результате чего на основе нелинейного фильтра предложена обобщенная схема структурной

идентифікації БС с ЛСП, которая позволила учесть достоинства основных методов преобразования БС с ЛСП.

Список литературы: 1. Філатова Г.Є. Структурна ідентифікація сигналів у кардіологічних системах: дис. канд. техн. наук: 05.11.17 / Філатова Ганна Євгенівна. – Харків, 2002. – 177 с. 2. Абакумов В.Г. Біомедичні сигнали. Генезис, обробка, моніторинг. Навчальний посібник / В.Г. Абакумов, О.І. Рибін, Й. Сватош. – К.: Нора-прінт, 2001. – 516 с. 3. Мурашко В.В. Электрокардиография: Учебное пособие / В.В. Мурашко, А.В. Струтинский. – М.: МЕДпресс, Элиста: Джангар, 1998. – 313 с. 4. Вычислительные системы и автоматическая диагностика заболеваний сердца / Под ред. Ц. Касереса, Л. Дрейфуса. – М.: Мир, 1974. – 504 с. 5. Файнзильберг Л.С. Методи та інструментальні засоби оцінювання стану об'єктів за сигналами з локально зосередженими ознаками: автореф. дис. на здобуття наук. ступеня доктора техн. наук: спец. 05.13.06 "Автоматизовані системи управління та прогресивні інформаційні технології" / Л.С. Файнзильберг. – К., 2004. – 35 с. 6. Файнзильберг Л.С. ФАЗАГРАФ® – ефективна інформаційна технологія обробки ЕКГ в задачі скринінга ішемічної хвороби серця / Файнзильберг Л.С. // Клінічна інформатика і телемедицина. – 2010. – Т. 6. – Вип. 7. – С. 22-30. 7. Айфичер Э. Цифровая обработка сигналов: Практический подход / Э. Айфичер, Б. Джервис. – М.: Издательский дом "Вильямс", 2004. – 992 с. 8. Сергиенко А.Б. Цифровая обработка сигналов: Учебное пособие / А.Б. Сергиенко. – СПб.: Питер, 2006. – 752 с.

Статья представлена д.т.н. проф. НТУ "ХПИ" С.М. Порошиным.

УДК 61.007+004.932.72

Нелінійна фільтрація біомедичних сигналів з локально зосередженими ознаками в задачі структурної ідентифікації / Філатова Г.Є. // Вісник НТУ "ХПИ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПИ". – 2011. – № 17. – С. 168 – 174.

У роботі розглядається задача структурної ідентифікації біомедичних сигналів (БС) з локально зосередженими ознаками (ЛЗО) за допомогою нелінійного фільтра. Розглянуті моделі корисного сигналу із ЛЗО, а також методи перетворення БС із ЛЗО на основі моделей корисного сигналу. Запропонований узагальнений метод нелінійної фільтрації БС із ЛЗО. Іл.: 3. Бібліогр.: 8 назв.

Ключові слова: структурна ідентифікація, біомедичний сигнал, локально зосереджені ознаки, нелінійний фільтр, модель корисного сигналу.

UDC 61.007+004.932.72

Nonlinear filtration of biomedical signals with the locally concentrated signs in task of structural identification / Filatova A.E. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 168 – 174.

The task of structural identification of biomedical signals (BS) with the locally concentrated signs (LCS) by means of nonlinear filter is examined in the work. The models of useful signal with LCS and transformation methods of BS with LCS on the basis of models of useful signal are considered. The generalized method of nonlinear filtration of BS with LCS is offered. Figs.: 3. Refs.: 8 titles.

Keywords: structural identification, biomedical signal, locally concentrated signs, nonlinear filter, model of useful signal.

Поступила в редакцію 04.04.2011

А.И. ЦВЕТКОВ, асп. Волжской государственной академии водного транспорта, Нижний Новгород

БИКРИТЕРИАЛЬНАЯ МОДЕЛЬ ОБСЛУЖИВАНИЯ ПОТОКА ОБЪЕКТОВ В РАЗВЕТВЛЕННОЙ РАБОЧЕЙ ЗОНЕ ОБСЛУЖИВАЮЩЕГО ПРОЦЕССОРА

Рассматривается модель обслуживания детерминированного потока объектов, проходящих транзитом узловую рабочую зону обслуживающего mobile-процессора. Эффективность стратегий обслуживания оценивается по значению двух критериев. Предлагается алгоритм синтеза оптимальных по Парето стратегий обслуживания. Приводится результат численного эксперимента. Ил.: 1. Библиогр.: 7 назв.

Ключевые слова: детерминированный поток объектов, стратегия обслуживания, оптимальность по Парето.

Постановка проблемы и анализ литературы. Технология обслуживания судов на ходу при транзитном прохождении ими крупномасштабной зоны ответственности сервисного предприятия получает все большее распространение на внутреннем водном транспорте РФ. В качестве обслуживающего выступает специализированное судно, предназначенное для выполнения одного или некоторого фиксированного набора работ: материально-технического снабжения, сбора подсланевых вод, ремонта и т.п.

Любое транзитное судно при прохождении зоны ответственности сервисного предприятия может запросить обслуживание, получить его или не быть обслуженным в зависимости от складывающейся ситуации.

Одна из задач диспетчера сервисного предприятия (лица принимающего решения – ЛПР) заключается в выработке (и последующем обеспечении реализации) наиболее рациональной в конкретной оперативной обстановке стратегии обслуживания потока судов. В условиях интенсивного судоходства оперативная выработка таких стратегий традиционными средствами, базирующимися на субъективных оценках и ситуационных представлениях ЛПР, весьма затруднена. В то же время существенную помощь в повышении эффективности функционирования сервисного предприятия может оказать автоматизация синтеза проектов стратегий обслуживания в процессе вычислительного анализа адекватной математической модели и решения соответствующих оптимизационных задач. Отметим при этом, что регламенты диспетчерского управления устанавливают, как правило, достаточно жесткие ограничения на длительность формирования стратегий обслуживания. Очевидно, что эти ограничения должны соблюдаться и при автоматизированной технологии. В рамках проектов

по её созданию удалось сформулировать [1 – 2] адекватные реальным транспортно-технологическим процессам базовые модели обслуживания, в которых понятие "mobile-процессор" было введено для обозначения обслуживающего судна; были предложены и алгоритмы синтеза стратегий обслуживания, построенные с соблюдением принципа оптимальности динамического программирования [3, 4].

Цель статьи заключается в дальнейшем развитии модельно-алгоритмического обеспечения автоматизированной технологии синтеза стратегий обслуживания путем решения в комплексе следующих задач:

1. Обобщение линейной модели [1 – 2] на практически значимый случай трехкомпонентной узловой рабочей зоны mobile-процессора.

2. Формализация бикритериального принципа оценки стратегий обслуживания; такой подход, переводя классическую постановку задачи оптимизации в разряд задач принятия решений, позволяет при оперативном планировании учитывать наряду с доходом за обслуживание также штраф за отказы в обслуживании.

3. Разработка алгоритма синтеза Парето-оптимальных [5, 6] стратегий обслуживания, конструктивно реализующего идеологию бикритериального динамического программирования [7], и его экспериментальная оценка.

Математическая модель обслуживания. Имеется n -элементный детерминированный поток объектов $O(n) = \{o(1), o(2), \dots, o(n)\}$, проходящих транзитом трехкомпонентную узловую рабочую зону Ξ процессора P , предназначенного для реализации однофазного обслуживания объектов (см. рис.).

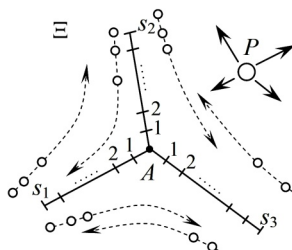


Рис. Трехкомпонентная узловая рабочая зона

В дискретной идеализации представляем зону Ξ как набор упорядоченных последовательностей элементарных участков с номерами $1, 2, \dots, s_q$, каждая из которых соответствует ветви зоны Ξ ($q = \overline{1, 3}$).

Поток $O(n)$ состоит из 6 подпотоков O_j ($j = \overline{1, 6}$) таких, что $\bigcup_{j=1}^6 O_j = O(n)$ и $\bigcap_{j=1}^6 O_j = \emptyset$. Объекты каждого подпотока поступают в зону Ξ по некоторой ветви, проходят через точку A и покидают ее по другой ветви. В пределах любого компонента зоны Ξ перемещение любого объекта осуществляется равномерно.

Обслуживающий объекты потока $O(n)$ процессор P характеризуется следующими целочисленными параметрами: u – номер ветви рабочей зоны, на которой расположен процессор P в начальный момент времени $t = 0$; z – номер элементарного участка на ветви u , на котором расположен процессор P в момент $t = 0$; $T_{-q}(T_{+q})$ – норма времени пребывания процессора P на элементарном участке при его автономном движении по q -й ветви рабочей зоны по направлению к точке A (от точки A), $q = \overline{1, 3}$; $w^+(i)$ – доход за обслуживание объекта $o(i)$, $i = \overline{1, n}$; $w^-(i)$ – штраф за отказ в обслуживании объекта $o(i)$, $i = \overline{1, n}$.

Каждый объект $o(i)$ характеризуется следующими целочисленными параметрами: $t(i)$ – момент поступления объекта в зону Ξ ; $\chi^-(i)$ ($\chi^+(i)$) – норма времени пребывания объекта на элементарном участке при движении по направлению к точке A (от точки A); $d^-(i)$ ($d^+(i)$) – номер ветви, по которой объект начинает движение в зоне Ξ (покидает зону Ξ), $d^-(i) \neq d^+(i)$; $\tau(i)$ – норма длительности обслуживания объекта $o(i)$ процессором P .

Стратегия ρ обслуживания объектов потока $O(n)$ процессором P представляет собой m -элементный кортеж ($m \in [0, n]$)

$$\rho = \begin{cases} (\varphi_1, \psi_1, \theta_1), (\varphi_2, \psi_2, \theta_2), \dots, (\varphi_m, \psi_m, \theta_m), & \text{при } m \geq 1, \\ \emptyset, & \text{при } m = 0, \end{cases}$$

в записи которого использованы следующие обозначения:

φ_j – идентификатор объекта $o(\varphi_j)$, обслуживаемого процессором P в очередь j ($\varphi_j \in [1, n]$); ψ_j – номер участка начала обслуживания объекта $o(\varphi_j)$ в очередь j ($\psi_j \in [1, s_q]$); θ_j – номер ветви зоны Ξ , на которой расположен участок начала обслуживания ψ_j ($\theta_j = [1, k]$); $j = \overline{1, m}$.

Множество допустимых стратегий обслуживания Ω однозначно определяется следующими ограничениями: каждая стратегия содержит только объекты, запросившие обслуживание; каждый объект $o(\varphi_j)$ в допустимой стратегии может быть обслужен процессором P не более одного раза; обслуживание более одного объекта одновременно запрещено; обслуживание объекта начинается и заканчивается в пределах зоны Ξ . Суммарный доход от обслуживания объектов потока $O(n)$ при

реализации стратегии ρ обозначим через $W^+(\rho)$, а суммарный штраф за отказы в обслуживании объектов – через $W^-(\rho)$. Как очевидно, указанные характеристики стратегии определяются выражениями

$$W^+(\rho) = \sum_{j=1}^m w^+(\phi_j), \quad W^-(\rho) = \sum_{j=1}^m w^-(\phi_j).$$

Общий подход к исследованию проблемы принятия решений при наличии нескольких критериев оценки базируется на концепции Парето [1] и в условиях рассматриваемой модели приводит к следующей биокритериальной задаче

$$\{\max_{\rho \in \Omega} (W^+(\rho)), \min_{\rho \in \Omega} (W^-(\rho))\}. \quad (1)$$

Решающий алгоритм. Поток объектов $O(n)$ и обслуживающий их процессор будем рассматривать как дискретную управляемую систему, на каждом этапе j управления которой для свободного процессора формируется вектор $\{\nu_j, \omega_j, \eta_j\}$. Здесь ν_j – номер назначенного на обслуживание на этапе j объекта, ω_j – номер участка начала его обслуживания, η_j – номер ветви рабочей зоны, на которой начинается обслуживание объекта. Данный вектор будем называть управлением.

Как очевидно, состояние ξ_j рассматриваемой системы на j -м этапе управления характеризуется значениями набора характеристик $t_{пр}, p, p_q, \Lambda$, где $t_{пр}$ – момент принятия решения, p и p_q – соответственно участок зоны Ξ и номер ветви расположения освободившегося процессора P , Λ – множество обслуженных на момент времени t объектов; состояние системы на начальном этапе обслуживания определяется набором $\xi_0 = (0, z, u, \emptyset)$. Множество $\Phi(\xi_j)$ допустимых управлений $\{\nu_j, \omega_j, \eta_j\}$ на j -м этапе может быть получено из кинематических соображений.

Пусть x – вектор, Y – множество векторов той же размерности. Через $x \oplus Y$ обозначим совокупность всех векторов v , представимых в виде $v = x + y$ ($y \in Y$).

Обозначим через M множество векторов-оценок, а через $eff(M)$ – максимальное по включению подмножество недоминируемых в M векторов. Под действием управлений $\{\nu_j, \omega_j, \eta_j\}$ система переходит из состояния ξ_j в состояние ξ_{j+1} . Выбранные управления позволяют обеспечить любую оценку из совокупности $[(w^+(\nu_j), -w^-(\nu_j)) \oplus B(\xi_{j+1})]$.

Тогда

$$B(\xi_j) = eff \left\{ \bigcup_{\{\nu_j, \omega_j, \eta_j\} \in \Phi(\xi_j)} [(w^+(\nu_j), -w^-(\nu_j)) \oplus B(\xi_{j+1})] \right\}. \quad (2)$$

На финальном этапе обслуживания имеет место соотношение

$$B(\xi) = \{0, \sum_{j=1}^n w^-(v_j)\}. \quad (3)$$

Выражения (2) – (3) образуют рекуррентные соотношения динамического программирования, позволяющие реализовать синтез полной совокупности эффективных оценок и соответствующих им оптимально-компромиссных стратегий обслуживания в задаче (1).

Результаты вычислительных экспериментов и выводы.

Приведем в табл. 1 результаты численного эксперимента для случая обслуживания объектов потока $O(11)$ с характеристиками $s_1 = 10, s_2 = 12, s_3 = 9, u = 1, z = 5, T_{-1} = 1, T_{-2} = 2, T_{-3} = 2, T_{+1} = 1, T_{+2} = 1, T_{+3} = 1$; значения параметров объектов потока приведены в табл. 1.

Таблица 1. Результаты численного эксперимента

i	$t(i)$	$\tau(i)$	$d^-(i)$	$d^+(i)$	$\chi^-(i)$	$\chi^+(i)$	$w^+(i)$	$w^-(i)$
1	1	6	1	2	1	4	30	10
2	4	8	1	3	2	3	45	20
3	8	5	2	1	3	4	20	12
4	10	6	3	2	2	3	21	5
5	15	4	1	2	2	4	12	9
6	18	6	2	3	3	4	28	10
7	20	9	3	2	2	3	14	11
8	22	8	1	2	1	3	17	15
9	24	6	2	1	2	3	39	14
10	28	5	3	2	2	4	30	12
11	30	4	1	3	1	2	25	15

Полная совокупность эффективных оценок и соответствующие им оптимально-компромиссные стратегии обслуживания приведены в табл. 2. Продолжительность решения задачи синтеза на компьютере с процессором AMD Turion 1,8 ГГц составила 63 с.

Таблица 2. Эффективные оценки и оптимально-компромиссные стратегии

$(W^+(\rho), W^-(\rho))$	ρ
(217, 40)	(1, 3, 1), (2, 3, 1), (10, 3, 3), (11, 2, 1), (6, 1, 3), (3, 8, 1), (9, 10, 1)
(204, 37)	(1, 8, 1), (2, 6, 1), (8, 6, 1), (11, 2, 1), (6, 1, 3), (3, 8, 1), (9, 10, 1)
(188, 36)	(2, 11, 1), (5, 10, 1), (8, 9, 1), (10, 1, 3), (11, 3, 3), (3, 8, 1), (9, 10, 1)

В работе построена математическая модель обслуживания конечного детерминированного потока объектов в трехкомпонентной узловой рабочей зоне mobile-процессора при наличии двух критериев оценки эффективности управления обслуживанием. Разработан алгоритм

синтеза полной совокупности эффективных оценок и соответствующих им оптимально-компромиссных стратегий обслуживания. Приведены результаты вычислительного эксперимента, демонстрирующие возможность штатной реализации алгоритма в системах поддержки принятия решений при диспетчерском управлении транспортно-технологическими процессами рассматриваемого типа.

Список литературы: 1. Коган Д.И. Проблема синтеза оптимального расписания обслуживания бинарного потока объектов mobile-процессором / Д.И. Коган, Ю.С. Федосенко, А.В. Шеянов / Труды III Международной конференции "Дискретные модели в теории управляющих систем", Москва, 1998. – М.: Изд-во МГУ им. М.В. Ломоносова, 1998. – С. 43-46. 2. Коган Д.И. Задача синтеза оптимального расписания обслуживания бинарного потока объектов в рабочей зоне mobile-процессора / Д.И. Коган, Ю.С. Федосенко. – Вестник Нижегородского университета. Математическое моделирование и оптимальное управление. – 1999. – Вып. 1 (20). – С. 179-187. 3. Беллман Р. Прикладные задачи динамического программирования / Р. Беллман, С. Дрейфус. – М.: Наука, 1965. – 457 с. 4. Коган Д.И. Задачи синтеза оптимальных стратегий обслуживания стационарных объектов в одномерной рабочей зоне процессора / Д.И. Коган, Ю.С. Федосенко // Автоматика и телемеханика. – 2010. – № 10. – С. 50-62. 5. Подиновский В.В. Парето-оптимальные решения многокритериальных задач / В.В. Подиновский, В.Д. Нозин. – М.: Физматлит, 2007. – 256 с. 6. Лотов В.А. Многокритериальные задачи принятия решений / В.А. Лотов, И.И. Поспелова. – М: МАКС Пресс, 2008. – 197 с. 7. Коган Д.И. Динамическое программирование и дискретная многокритериальная оптимизация / Д.И. Коган. – Н.Новгород: Изд-во ННГУ, 2005. – 260 с.

Статья представлена д.ф.-м.н. проф. М.И. Фейгин

УДК 681.31+519.8

Бікритеріальна модель обслуговування потоку об'єктів в розгалуженій робочій зоні обслуговуючого процесора / Цвєтков А.І. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 175 – 180.

Розглядається модель обслуговування детермінованого потоку об'єктів, проходящих транзитом вузлову робочу зону обслуговуючого mobile-процесора. Ефективність стратегій обслуговування оцінюється за значенням двох критеріїв. Пропонується алгоритм синтезу оптимальних по Парето стратегій обслуговування. Наводиться результат чисельного експерименту. Іл.: 1. Бібліогр.: 7 назв.

Ключові слова: детермінований потік об'єктів, стратегія обслуговування, оптимальність по Парето.

UDC 681.31+519.8

The bicriteria service model for the flow of objects in the ramified work area of service processor / Tsvetkov A.I. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 175 – 180.

The service model for the deterministic flow of objects going through nodal work area of service mobile processor is described. Values of two criteria have been taken into account when assessing the quality of the control policies. The algorithm of the Pareto optimal service policies is offered. The numerical experiment results are given. Figs.: 1. Refs.: 7 titles.

Keywords: deterministic flow of objects, service policy, Pareto optimality.

Поступила в редакцию 15.02.2011

С.В. ЧОПОРОВ, зав. лаб. геоинформационных систем Запорожского национального университета, Запорожье,
С.И. ГОМЕНЮК, д.т.н., проф. Запорожского национального университета, Запорожье

ПРОБЛЕМНО-ОРИЕНТИРОВАННЫЙ ЯЗЫК ГЕОМЕТРИЧЕСКОГО МОДЕЛИРОВАНИЯ НА БАЗЕ ТЕОРИИ R-ФУНКЦИЙ

В работе рассмотрена проблема геометрического моделирования сложных объектов на базе теории R-функций. Предложен подход к формализации описания геометрической модели на основе проблемно-ориентированного языка. Ил.: 2. Библиогр.: 9 назв.

Ключевые слова: геометрическая модель, R-функция, проблемно-ориентированный язык.

Постановка проблемы. Развитие современного производства требует качественной эволюции системы технологического анализа. Высокая конкуренция на рынке машиностроения и инженерных технологий ведет украинских разработчиков к необходимости проектирования все более сложных и неординарных решений. Дороговизна ошибки при реализации таких решений на практике делает все более актуальным и значимым их компьютерное моделирование и внесение оптимизационных изменений по результатам моделирования.

В современной технике одним из наиболее трудоемких является этап анализа функциональных характеристик, механических свойств, прочности и долговечности проектируемых конструкций, сооружений, машин и механизмов. В современных САПР данный этап, как правило, связывают с решением двух групп проблем: 1) автоматизация подготовки геометрической модели объекта и 2) автоматизация вычислений физических характеристик исследуемого объекта. Исследование второй группы показывает, что большое количество возникающих на практике задач (в частности анализ напряженно-деформированного состояния) связано с необходимостью решения систем дифференциальных или интегральных уравнений. Для их решения современные САПР используют различные вычислительные методы, основанные на идеи перехода от непрерывной задачи к дискретной, например, метод конечных элементов [1 – 3]. При этом необходимо построение различных математических моделей, в составляющие части которых, как правило, входят сложные геометрические модели исследуемых объектов. Следовательно, актуальной является проблема автоматизации процесса построения геометрических моделей сложной формы, которая может

быть условно разделена на две составляющие: 1) формализация описания геометрической модели сложного объекта; 2) автоматизация построения на ее базе адекватной дискретной модели.

Анализ литературы. Анализ современных САПР в машиностроении показывает, что наиболее распространенными подходами к геометрическому моделированию сложных тел являются:

- инженерные чертежи;
- граничное представление;
- конструктивная блочная геометрия;
- функциональное представление.

Первый подход выглядит наиболее привлекательно в глазах инженера, позволяя ему строить привычные чертежи геометрических проекций, и получил свое развитие в таких системах как AutoCAD, КОМПАС и т.д. Однако, при проектировании нестандартных решений не всегда имеется возможность выделить геометрические проекции, использование которых будет в полной мере отображать все особенности конструкции.

Возможность представлять тело совокупностью ограничивающих его объем оболочек и логических операций над ними, универсальность используемых структур делают второй подход одним из наиболее применимых в компьютерной графике [4, 5]. Граничное представление лежит в основе ядра геометрического моделирования Parasolid и системы Romulus. Однако необходимо отметить, что получение системы оболочек, описывающих сложное тело, является весьма трудоемкой задачей.

Конструктивная блочная геометрия (Constructive Solid Geometry, CSG) – подход, позволяющий создать геометрическую модель сложного объекта с помощью булевых операций над элементами некоторого множества более простых объектов, формирующих библиотеку примитивов. Недостатком подхода является относительная сложность получения границы, адекватно отражающей моделируемый объект. Также ограниченность набора базовых примитивов делает ограниченной область определения такого подхода.

Функциональное представление – подход, который основан на идее моделирования геометрической структуры тела с помощью математических функций или соотношений. Одним из наиболее распространенных тут является использование неявных функций, которые (как правило) больше нуля внутри области, равны нулю на границе и меньше нуля вне области. Использование теории R-функций, разработанной В.Л. Рвачевым [6 – 8], позволяет получать сложные неявные функции, соответствующие сложным областям, конструктивно. Теоретически такой подход, сочетая в себе богатство аппарата

математических функций, является наиболее универсальным. Однако, на практике его применение затрудняется сложностью умозрительного восприятия и верификации получаемых в итоге формул.

Цель статьи – разработка проблемно-ориентированного языка геометрического моделирования, позволяющего упростить процесс построения и верификации функциональных геометрических моделей на базе теории R-функций.

Основы метода R-функций. Пусть Ω – сложное тело, геометрическую модель которого необходимо получить. Наиболее общим методом определения множества точек X , образующих объект Ω , является определение предиката A , который может быть вычислен для каждой точки p пространства [13]:

$$X = \{p | A(p) = \text{true}\} . \quad (1)$$

Таким образом, X определено неявно и состоит из всех точек, удовлетворяющих условию, определенному предикатом A . Простейшей формой предиката является ограничение на знак некоторой действительной функции. Например, если $f(x, y) = Ax + By + C$, тогда $f(x, y) = 0$, $f(x, y) \geq 0$ и $f(x, y) < 0$ определяют прямую, закрытую полуплоскость и открытую полуплоскость, соответственно.

Для определения более сложных областей могут быть использованы R-функции в качестве логических операций над более простыми функциями. Наиболее распространенной на практике системой R-функций является:

$$\neg x \equiv -x, \quad x \wedge y \equiv x + y - \sqrt{x^2 + y^2}, \quad x \vee y \equiv x + y + \sqrt{x^2 + y^2}. \quad (2)$$

Словарь проблемно-ориентированного языка. Естественным решением проблемы автоматизации геометрического моделирования на базе функционального подхода и теории R-функций является разработка проблемно-ориентированного языка, позволяющего описывать математические формулы. Такой подход сочетает в себе гибкость описания и компактность хранения модельных данных. Однако, в работе [5] отмечается, что порядка 60% механических деталей могут быть представлены с помощью системы конструктивной блочной геометрии, в основе которой только прямоугольные балки и цилиндрические примитивы. Следовательно, разумным является формирование библиотеки R-функций наиболее распространенных объектов с последующей интеграцией этой библиотеки в словарь проблемного

языка, что позволит моделировать ряд стандартных деталей в терминах конструктивной блочной геометрии, а для моделирования элементов с нестандартной формой использовать определяемые пользователем функции.

Геометрическое моделирование многих инженерных деталей и конструкций может быть сведено к последовательным логическим операциям объединения и пересечения полуплоскостей или полупространств.

Полуплоскость заданная упорядоченной парой точек $A_1(x_1, y_1)$ и $A_2(x_2, y_2)$, и расположенная по правую сторону при движении от первой ко второй точке, может быть представлена формулой

$$F_{A_1 A_2}(x, y, x_1, y_1, x_2, y_2) = (x - x_1)(y_2 - y_1) - (y - y_1)(x_2 - x_1). \quad (3)$$

Аналогично, полупространство, заданное с помощью точки $P(x_p, y_p, z_p)$, принадлежащей граничной плоскости, и внешней нормалью $n = (x_n, y_n, z_n)$ можно определить формулой

$$F_{P\bar{n}}(x, y, z, x_p, y_p, z_p, x_n, y_n, z_n) = -x_n(x - x_p) - y_n(y - y_p) - z_n(z - z_p). \quad (4)$$

Область, ограниченная эллипсом с центром в точке (x_0, y_0) , большая полуось которого равна a , малая полуось – b , может быть представлена формулой

$$Ellipse(x, y, x_0, y_0, a, b) = 1 - \frac{(x - x_0)^2}{a^2} - \frac{(y - y_0)^2}{b^2}, \quad (5)$$

частным случаем которой является круг радиуса r с центром в точке (x_0, y_0)

$$Disk(x, y, r, x_0, y_0) = r^2 - (x - x_0)^2 - (y - y_0)^2. \quad (6)$$

Аналогично, в трехмерном пространстве функция, описывающая область, ограниченную эллипсоидом, центр которого находится в точке (x_0, y_0, z_0) , а полуоси равны a , b и c , может быть представлена формулой

$$Ellipsoid(x, y, z, x_0, y_0, z_0, a, b, c) = 1 - \frac{(x - x_0)^2}{a^2} - \frac{(y - y_0)^2}{b^2} - \frac{(z - z_0)^2}{c^2}, \quad (7)$$

частным случаем которой является шар радиуса r с центром в точке (x_0, y_0, z_0) :

$$Ball(x, y, z, r, x_0, y_0, z_0) = r^2 - (x - x_0)^2 - (y - y_0)^2 - (z - z_0)^2. \quad (8)$$

Выпуклый многоугольник, заданный упорядоченной последовательностью из $n \geq 3$ вершин $V = \{(x_i, y_i)\}$, при условии обхода вершин по часовой стрелке, может быть представлен конъюнкцией n полуплоскостей, последовательно образуемых парами вершин:

$$\begin{aligned} & \text{ConvexPolygon}(x, y, V = \{(x_i, y_i)\}, n) = \\ & = F_{A_1 A_2}(x, y, x_1, y_1, x_2, y_2) \wedge F_{A_1 A_2}(x, y, x_2, y_2, x_3, y_3) \wedge \dots \wedge \\ & \wedge F_{A_1 A_2}(x, y, x_{n-1}, y_{n-1}, x_n, y_n) \wedge F_{A_1 A_2}(x, y, x_n, y_n, x_0, y_0). \end{aligned} \quad (9)$$

Частным случаем выпуклого многоугольника является правильный n -угольник, вписанный в окружность радиуса r с центром в точке $C(x_0, y_0)$, первая вершина которого расположена на пересечении оси ординат и окружности:

$$\begin{aligned} V = \left\{ (\xi_i; \eta_i) : \xi_i = x_0 + r \sin \alpha_i, \eta_i = y_0 + r \cos \alpha_i, \alpha_i = \frac{2\pi}{n}(i-1), i = \overline{1, n} \right\}, \quad (10) \\ \text{RegularPolygon}(x, y, r, x_0, y_0, n) = \text{ConvexPolygon}(x, y, V, n). \end{aligned}$$

R -функция, которая представляет параллелограмм с координатами нижнего левого угла $A(x_0, y_0)$, длиной основания a , высотой h и углом $0 < \alpha < \pi$ в точке A , может быть представлена формулой

$$\begin{aligned} V = \{A(x_0, y_0), B(x_0 + \Delta_x, y_0 + h), C(x_0 + a + \Delta_x, y_0 + h), D(x_0 + a, y_0)\}, \quad (11) \\ \text{Parallelogram}(x, y, x_0, y_0, a, h, \alpha) = \text{ConvexPolygon}(x, y, V, 4), \end{aligned}$$

где $\Delta_x = h \operatorname{ctg}(\alpha)$,

частым случаем которой при $\alpha = \pi / 2$ будет прямоугольник. Однако, с вычислительной точки зрения прямоугольник рационально представить в виде логического пересечения двух полос:

$$\text{Rectangle}(x, y, x_0, y_0, w, h) = \left[\left(\frac{w}{2} \right)^2 - \left(x - x_0 - \frac{w}{2} \right)^2 \right] \wedge \left[\left(\frac{h}{2} \right)^2 - \left(y - y_0 - \frac{h}{2} \right)^2 \right], \quad (12)$$

где (x_0, y_0) – координаты нижнего левого угла, w – ширина, h – высота.

Параллелепипед, определенный вершиной $A(x_0, y_0, z_0)$, шириной a , высотой h , глубиной d , углами α и β , можно рассматривать как пересечение двух полос пространства с сечением в форме параллелограмма: в плоскости xOy и плоскости zOy , может быть представлен формулой

$$\begin{aligned} & \text{Parallelepiped}(x, y, z, x_0, y_0, z_0, a, h, d, \alpha, \beta) = \\ & = \text{Parallelogram}(x, y, x_0, y_0, a, h, \alpha) \wedge \text{Parallelogram}(z, y, z_0, y_0, d, h, \beta). \end{aligned} \quad (13)$$

Используя рассуждения аналогичные (12), прямоугольный параллелепипед можно представить формулой:

$$\begin{aligned} \text{Box}(x, y, z, x_0, y_0, z_0, w, h, d) = & \left[\left(\frac{w}{2} \right)^2 - \left(x - x_0 - \frac{w}{2} \right)^2 \right] \wedge \\ & \wedge \left[\left(\frac{h}{2} \right)^2 - \left(y - y_0 - \frac{h}{2} \right)^2 \right] \wedge \left[\left(\frac{d}{2} \right)^2 - \left(z - z_0 - \frac{d}{2} \right)^2 \right]. \end{aligned} \quad (14)$$

Таким образом, взяв за основу язык на базе ECMAScript [9], можно определить в нем библиотеку глобальных объектов и функций для описания геометрической модели. В частности, в процессе разработки системы геометрического моделирования qMesher был определен глобальный объект *Geom*, поля которого хранят метаинформацию о свойствах геометрической модели (область определения, плотность начальной сетки, имя результирующей функции, степень оптимизации), и глобальные функции инкапсулирующие R-функции (3) – (14) и другие популярные на практике геометрические формы (например, сектор круга, прямоугольник со скругленными углами и прочие).

Например, на рис. 1 изображены описание на проблемно-ориентированном языке и визуализация геометрической модели, в форме объединения семи шаров.

```

1 Geom.dimension = 3;
2 Geom.x_min = -1.0;
3 Geom.x_max = 1.0;
4 Geom.y_min = -1.0;
5 Geom.y_max = 1.0;
6 Geom.z_min = -1.0;
7 Geom.z_max = 1.0;
8 Geom.nx = 20;
9 Geom.ny = 20;
10 Geom.nz = 20;
11 Geom.name = "balls";
12 Geom.optiterations = 2;
13
14 function balls(x, y, z)
15 {
16   var r1 = 0.7; // радиус центрального шара
17   var r2 = 0.3; // радиус окружающих шаров
18   var ball1 = Ball(x, y, z, r1, 0, 0, 0); // центральный шар
19   var ball2 = Ball(x, y, z, r2, 1, 0, 0);
20   var ball3 = Ball(x, y, z, r2, -1, 0, 0);
21   var ball4 = Ball(x, y, z, r2, 0, 1, 0);
22   var ball5 = Ball(x, y, z, r2, 0, -1, 0);
23   var ball6 = Ball(x, y, z, r2, 0, 0, 1);
24   var ball7 = Ball(x, y, z, r2, 0, 0, -1);
25   return Union(ball1, ball2, ball3, ball4, ball5, ball6, ball7);
26 }

```

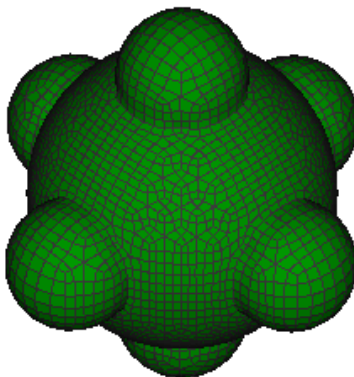


Рис. 1. Определение объекта на проблемно-ориентированном языке

Для построения объектов с элементами нестандартной формы могут быть использованы определенные пользователем R-функции (пример на рис. 2).

```

1 Geom.dimension = 3; // размерность пространства
2 Geom.x_min = -3.0; // минимум по оси Ox
3 Geom.x_max = 3.0; // максимум по оси Ox
4 Geom.y_min = -3.0; // минимум по оси Oy
5 Geom.y_max = 3.0; // максимум по оси Oy
6 Geom.z_min = -3.0; // минимум по оси Oz
7 Geom.z_max = 3.0; // максимум по оси Oz
8 Geom.nx = 40; // густота сетки вдоль оси Ox
9 Geom.ny = 40; // густота сетки вдоль оси Oy
10 Geom.nz = 40; // густота сетки вдоль оси Oz
11 Geom.optiterations = 5; // количество итераций оптимизации
12 function ribbon (x, y, z)
13 {
14   return Con ( Math.sin(x + y) - z + 1 , z - Math.sin(x + y) );
15 }
16 function r_main(x, y, z)
17 {
18   var w = 0 - x * x;
19   var h = 9 - y * y;
20   var r = ribbon (x, y, z);
21   return Intersection (w, h, r);
22 }

```

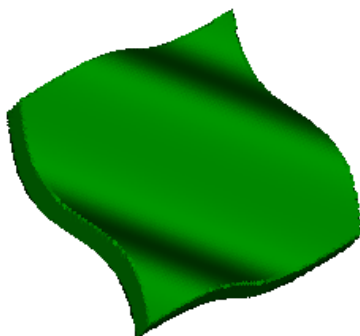


Рис. 2. Определение объекта с помощью функций пользователя

Выводы. В результате проделанной работы предложен подход к формализации описания геометрических моделей, соединяющих общность и универсальность функционального представления на базе теории R-функций с простотой и наглядностью методов конструктивной блочной геометрии.

Список литературы: 1. Городецкий А.С. Информационные технологии расчета и проектирования строительных конструкций. Учебное пособие / А.С. Городецкий, В.С. Шмуклер, А.В. Бондарев. – Харьков: НТУ "ХПИ", 2003. – 889 с. 2. Толоч В.А. Метод конечных элементов: теория, алгоритмы, реализация / В.А. Толоч, В.В. Киричевский, С.И. Гоменюк, С.Н. Гребенюк, Д.П. Бувайло. – К.: Наукова думка, 2003. – 316 с. 3. Smith I.M. Programming the finite element method / I.M. Smith, D.V. Griffiths. – England, Chichester: Wiley, 2004. – 646 p. 4. Голованов Н.Н. Геометрическое моделирование / Н.Н. Голованов. – М.: Изд-во физ.-мат. лит., 2002. – 472 с. 5. Agoston M.K. Computer graphics and geometric modeling: implementation and algorithms / M.K. Agoston. – London: Springer-Verlag, 2005. – 959 p. 6. Рвачев В.Л. Введение в теорию R-функций / В.Л. Рвачев, Т.И. Шейко // Проблемы машиностроения. – 2001. – Т. 4. – № 1–2. – С. 46–58. 7. Рвачев В.Л. Проблемно-ориентированные языки и системы для инженерных расчетов / В.Л. Рвачев, А.Н. Шевченко. – К.: Техніка, 1988. – 198 с. 8. Рвачев В.Л. Теория R-функций и некоторые ее приложения / В.Л. Рвачев. – К.: Наукова думка, 1982. – 552 с. 9. Standard ECMA-262. ECMAScript Language Specification [Электронный ресурс] – 2009. – 252 p. Режим доступа: <http://www.ecma-international.org/publications/files/ECMA-ST/ECMA-262.pdf>

УДК 681.5.001.63: 519.711

Проблемно-орієнтована мова геометричного моделювання на базі теорії R-функцій / Чопоров С.В., Гоменюк С.І. // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 181 – 188.

В роботі розглянута проблема геометричного моделювання складних об'єктів на базі теорії R-функцій. Запропоновано підхід до формалізації описання геометричної моделі на основі проблемно-орієнтованої мови. Іл.: 2. Бібліогр.: 9 назв.

Ключові слова: геометрическая модель, R-функція, проблемно-орієнтована мова.

UDC 681.5.001.63: 519.711

A domain-specific language for geometrical modeling on the basis of R-functions
/ **Choporov S.V., Gomenyuk S.I.** // Herald of the National Technical University "KhPI". Subject
issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 181 –
188.

The problem of geometrical modeling of complex objects on the basis of R-functions is described in the article. Authors propose a domain-specific language approach to formalization of description of geometrical models. Figs.: 2. Refs.: 9 titles.

Keywords: geometrical model, R-function, domain-specific language.

В.А. ЯЦКО, к.т.н., доц. Новосибирского государственного технического университета, Новосибирск

АНАЛИЗ СЕБЕСТОИМОСТИ С ИСПОЛЬЗОВАНИЕМ НЕЧЕТКИХ МОДЕЛЕЙ

Рассматривается новый подход к калькулированию себестоимости продукции с учетом множественного выбора баз распределения постоянных расходов. В результате применения данного подхода формируется множество нечетких переменных, исследование которых позволяет провести анализ производственной программы предприятия с целью ее оптимизации. Ил.: 1. Библиогр.: 11 назв.

Ключевые слова: калькулирование себестоимости, постоянные расходы, нечеткая переменная.

Постановка проблемы. Одной из проблем, связанных с анализом прибыльности предприятий, является неоднозначность и недостоверность получаемых оценок себестоимости как для отдельных единиц продукции, так и при оценке издержек для отдельных видов деятельности. Большинство известных методов калькулирования себестоимости продукции можно отнести к так называемым методам Absorption Costing, предполагающим полное распределение всех производственных затрат. К сожалению, оценки себестоимости, получаемые с использованием подобных методов, существенно зависят от выбора базы распределения постоянных расходов, что в некоторых случаях приводит к неверным выводам о доходности или убыточности производства отдельных видов продукции, что, в свою очередь, ведет к принятию ошибочных решений при разработке производственной программы предприятия.

Анализ литературы. Для включения постоянных издержек производства в состав себестоимости изделий применяют различные методы пропорционального распределения таких издержек [1, 2]. Наиболее часто используется метод разнесения постоянных издержек по единой ставке, когда для разнесения этих издержек выбирается какая-то единая для всего предприятия величина. Обычно в качестве базы распределения используется заработная плата производственных рабочих, значительно реже – машиночасы, стоимость основных производственных материалов, прямые затраты (заработная плата производственных рабочих + стоимость основных материалов), объем произведенной продукции, оптовая цена продукции и т.п. К достоинствам данного метода распределения постоянных издержек относится простота учета, что обуславливает его популярность. Однако, себестоимость

продукции, рассчитанная с использованием данного подхода, весьма значительно зависит от выбора базы распределения. Известны примеры, когда в процессе анализе себестоимости и рентабельности отдельных видов продукции при замене базы распределения постоянных расходов получался парадоксальный результат – прибыльные изделия оказывались убыточными и наоборот.

В [3, 4] приводятся примеры, когда при увеличении продаж "высокоприбыльных" продуктов общая рентабельность продаж снижается. Одной из попыток решения указанной проблемы является ABC-метод калькулирования, когда вместо единой базы распределения постоянных расходов вводится несколько "носителей затрат" (cost driver), а для каждого вида постоянных расходов можно указать свой носитель затрат. Однако, и при использовании ABC-метода остается проблема выбора наиболее подходящих баз распределения постоянных расходов. Неудачный выбор "носителя затрат" может существенно повлиять на значение себестоимости продукции.

Цель статьи – разработка нового подхода к калькулированию и анализу себестоимости продукции за счет введения нечетких моделей, позволяющих описать и проанализировать допустимое множество оценок себестоимости.

Обобщенная полная себестоимость. В работе [5] вместо калькулирования единственного варианта себестоимости продукции предлагается рассчитывать множество значений себестоимости для различных вариантов распределения накладных расходов. На основе сформированного множества оценок себестоимости определяется некоторая нечеткая переменная, представляющая возможный диапазон значений себестоимости.

Введем в рассмотрение обобщенную полную себестоимость единицы продукции i -го вида S_i

$$S_i = V_i + \sum_{j=1}^m \sum_{k=1}^{K_j} \alpha_{jk} \cdot D_{ijk}, \quad (1)$$

где V_i – прямые затраты, приходящиеся на единицу продукции i -го вида; m – число различных видов накладных затрат; K_j – число возможных вариантов разнесения j -го вида накладных затрат; α_{jk} – весовой

коэффициент для k -й базы распределения, $\alpha_{jk} \geq 0$, $\sum_{k=1}^{K_j} \alpha_{jk} = 1$; D_{ijk} –

величина накладных затрат j -го вида, приходящихся на единицу продукции i -го вида при использовании k -го носителя затрат.

Себестоимость продукции вида (1) представляет собой средневзвешенную величину всех возможных себестоимостей продукции i -го вида, которые могли бы быть рассчитаны при использовании различных баз распределения накладных расходов. Очевидно, что даже в случае, если бы для каждого вида накладных расходов использовалась бы только одна "своя" база распределения, то число возможных себестоимостей продукции i -го вида было бы равно

$$K^* = K_1 \cdot K_2 \cdot \dots \cdot K_m = \prod_{j=1}^m K_j.$$

На первый взгляд, такой подход мало пригоден для практического использования вследствие того, что возможно бесконечное множество различных значений себестоимости вида (1) при задании различных значений весовых коэффициентов α_{jk} . Однако, применение данного подхода позволяет проанализировать некоторые аспекты формирования себестоимости продукции. Рассмотрим более подробно данный подход применительно к калькулированию себестоимости одного вида продукции, поэтому в дальнейших выкладках будем упускать индекс i .

Без потери общности будем полагать, что $D_{j1} \leq D_{j2} \leq \dots \leq D_{jK_j}$ для $j = 1, \dots, m$. Очевидно, что $S \in \left[V + \sum_{j=1}^m D_{j1}, V + \sum_{j=1}^m D_{jK_j} \right]$. Тогда система ограничений

$$\begin{cases} S = V + \sum_{j=1}^m \sum_{k=1}^{K_j} \alpha_{jk} \cdot D_{jk}, \\ \sum_{k=1}^{K_j} \alpha_{jk} = 1, \alpha_{jk} \geq 0, \end{cases} \quad (2)$$

определяет симплекс (выпуклый многогранник) в пространстве размерностью $\sum_{j=1}^m K_j$, где в качестве переменных выступают весовые

коэффициенты α_{jk} , а себестоимость S выступает в качестве параметра. Объем симплекса является функцией от величины параметра S . На рис. приведен примерный график этой функции при $m = 2$ и $K_1 = 2, K_2 = 3$, что соответствует тому, что имеется два вида накладных расходов и для первого вида возможно два варианта распределения, для второго – три

варианта. Аналитическое выражение функции для вычисления объема симплекса может быть использовано в качестве функции принадлежности неотрицательной нечеткой переменной x с областью определения на интервале $\left[V + \sum_{j=1}^m D_{j1}, V + \sum_{j=1}^m D_{jK_j} \right]$.

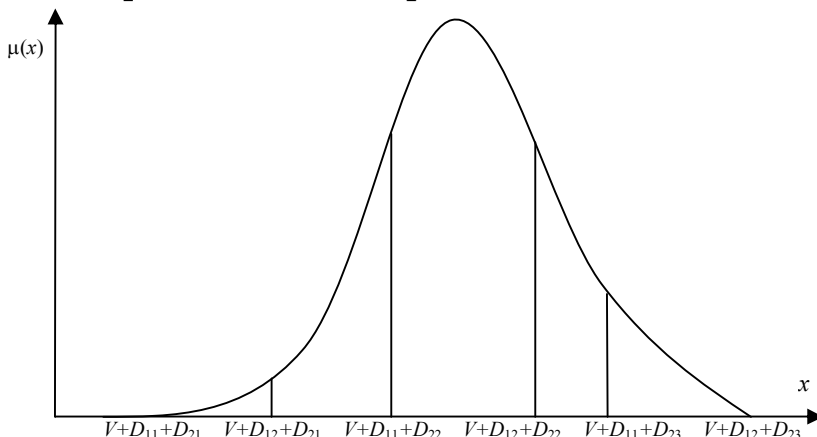


Рис. Примерный график функции принадлежности $\mu(x)$

Было получено следующее аналитическое выражение для функции принадлежности $\mu(x)$

$$\mu(x) = p \sum_{j_1=1}^{K_1} \sum_{j_2=1}^{K_2} \dots \sum_{j_m=1}^{K_m} \frac{1(x - V - \sum_{l=1}^m D_{lj_l}) \cdot (x - V - \sum_{l=1}^m D_{lj_l})^{p-1}}{\prod_{l=1}^m q_{lj_l}}, \quad (3)$$

где $1(x)$ – единичная функция Хевисайда, $1(x) = \begin{cases} 1, & x \geq 0, \\ 0, & x < 0, \end{cases}$

$q_{lj} = \prod_{\substack{i=1 \\ i \neq j}}^{K_j} (D_{li} - D_{lj})$; $p = \sum_{j=1}^m (K_j - 1) = \sum_{j=1}^m K_j - m$. В данном случае

приведено выражение для ненормированной функции принадлежности в традиционном понимании теории нечетких множеств (т.е. не выполняется

условие $\max_x \mu(x) = 1$), что в некоторой степени упрощает анализ

нечеткой переменной.

Представление себестоимости в виде нечеткой переменной позволяет повысить информативность анализа себестоимости продукции за счет интеграции в данном представлении множества возможных оценок себестоимости.

В последние годы резко возрос интерес к использованию теории нечетких множеств для анализа различных экономических явлений. Использование инструментов нечеткого анализа позволяет каким-либо образом "измерить" неопределенность экономического явления, а следовательно, обеспечить принятие более сбалансированных решений с учетом этой неопределенности.

К сожалению, "измерение" неопределенности экономических явлений практически всегда сводится к получению некоторой экспертной оценки области определения и функции принадлежности нечеткой переменной, что заведомо субъективно. В отличие от применявшихся ранее подходов [6 – 10], в данной работе область определения и функция принадлежности формируются аналитически на основе данных бухгалтерского управленческого учета, что обеспечивает большую степень объективности.

Для интегральной оценки безубыточности отдельных ассортиментных позиций был использован показатель "коэффициент надежности безубыточности" вида

$$M(x) = \int_{V + \sum_{j=1}^m D_{j1}}^x \mu(x) dx = \sum_{j_1=1}^{K_1} \sum_{j_2=1}^{K_2} \dots \sum_{j_m=1}^{K_m} \frac{1(x - V - \sum_{l=1}^m D_{lj_l}) \cdot (x - V - \sum_{l=1}^m D_{lj_l})^p}{\prod_{l=1}^m q_{lj_l}},$$

где аргумент x соответствует отпускной цене предприятия. Данный коэффициент может принимать значения в диапазоне от 0 до 1. Чем больше значение коэффициента для конкретного вида продукции, тем выше уверенность, что данная продукция является безубыточной.

Выводы. Рассмотренный в работе подход позволяет повысить информативность анализа себестоимости продукции за счет того, что множество возможных оценок себестоимости описывается посредством нечетких переменных. Использование формализованных процедур теории нечетких множеств повышает обоснованность принимаемых решений по оптимизации производственной программы предприятия.

Список литературы: 1. Шанк Дж.К. Стратегическое управление затратами / Дж.К. Шанк, В. Говиндараджан. – СПб.: ЗАО "Бизнес Микро", 1999. – 288 с. 2. Вахрушина М.А. Бухгалтерский управленческий учет / М.А. Вахрушина. – М.: Омега-Л, 2007. – 570 с. 3. Попова Л.В. Формирование учетно-аналитической системы затрат на промышленных предприятиях / Л.В. Попова, В.А. Константинов, И.А. Маслова, М.М. Коростелкин. – М.: Дело и сервис, 2007. – 224 с. 4. Аткинсон Э.А. Управленческий учет / Э.А. Аткинсон, Р.Д. Банкер, Р.С. Каплан, М.С. Янг. – М.: Издательский дом "Вильямс", 2007. – 880 с. 5. Кофман А. Введение в теорию нечетких множеств в управлении предприятиями / А. Кофман, Х. Хил Алуха. – Минск: Высшая школа, 1992. – 223 с. 6. Хил Лафуенте А.М. Финансовый анализ в условиях неопределенности / А.М. Хил Лафуенте. – Минск: Тэхналогія, 1998. – 150 с. 7. Nachtmann H. Fuzzy Activity Based Costing: A Methodology for Handling Uncertainty in Activity Based Costing Systems / H. Nachtmann, K.L. Needy // The Engineering Economist, 2001. – V. 46. – № 4. – P. 245-273. 8. Nachtmann H. Methods for Handling Uncertainty in Activity Based Costing Systems / H. Nachtmann, K.L. Needy // The Engineering Economist, 2003. – V. 48. – № 3. – P. 259-282. 9. Пегат А. Нечеткое моделирование и управление / А. Пегат. – М.: БИНОМ. Лаборатория знаний, 2009. – 798 с. 10. Недосекин А.О. Методологические основы моделирования финансовой деятельности с использованием нечетко-множественных описаний: автореф. дис. на соискание степени д-ра экон. наук: спец. 08.00.13 "Математические и инструментальные методы экономики" / А.О. Недосекин. – СПб., 2003. – 37 с. 11. Яцко В.А. Калькулирование себестоимости продукции с использованием аппарата теории нечетких множеств / В.А. Яцко // Проблемы современной экономики. – 2009. – № 4 (32). – С. 187-192.

Статья представлена профессором кафедры автоматика Новосибирского государственного технического университета, д.т.н. А.А. Воеводой.

УДК 338.512

Аналіз собівартості з використанням нечітких моделей / Яцко В.А // Вісник НТУ "ХПІ". Тематичний випуск: Інформатика і моделювання. – Харків: НТУ "ХПІ". – 2011. – № 17. – С. 189 – 194.

Розглядається новий підхід до калькулювання собівартості продукції з урахуванням множинного вибору баз розподілу постійних витрат. У результаті застосування даного підходу формується множина нечітких змінних, дослідження яких дозволяє провести аналіз виробничої програми підприємства з метою її оптимізації. Лл.: 1. Бібліогр.:11.

Ключові слова: калькулювання собівартості, постійні витрати, нечітка змінна.

UDC 338.512

The cost accounting with use fuzzy models / Yatsko V. A. // Herald of the National Technical University "KhPI". Subject issue: Information Science and Modelling. – Kharkov: NTU "KhPI". – 2011. – №. 17. – P. 189 – 194.

The new approach for cost accounting is considered taking into account the multiple-choice of base to allocate fixed expenses. As a result of this approach is formed by a set of fuzzy variables, the study of which allows to carry out the analysis of the production program of the enterprise in order to optimize it. Figs.: 1. Refs.: 11 titles.

Keywords: costing, fixed expenses, fuzzy variable.

Поступила в редакцію 15.02.2011

Содержание

Ахметшин А.М., Степаненко А.А. Интерференционный метод анализа радиологических изображений в пространстве признаков преобразования Радона	4
Великая Я.Г. Проблема эквивалентности в структурированных моделях вычислений	10
Гофман Є.О., Олійник А.О., Субботін С.О. Агентный метод синтеза деревьев решений	16
Дмитриенко В.Д., Заковоротный А.Ю., Носков В.И. Линеаризация математической модели дизель-поезда с тяговым асинхронным приводом	26
Жияков Е.Г., Болдышев А.В., Курлов А.В., Фирсова А.А., Эсауленко А.В. Об избирательном преобразовании частотных компонент речевых сигналов в задаче сжатия	37
Жияков Е.Г., Прохоренко Е.И., Болдышев А.В., Фирсова А.А., Фатова М.В. Сегментация речевых сигналов на основе анализа особенностей распределения долей энергии по частотным интервалам	44
Иванов В.Г., Ломоносов Ю.В., Любарский М.Г., Кошева Н.А., Гвозденко М.В., Мазниченко Н.И. Сжатие данных в системе Хаара с учетом объема первичной дискретизации	51
Иванов Д.Е. Применение алгоритмов симуляции отжига в задачах идентификации цифровых схем	60
Куимова А.С., Федосенко Ю.С. Синтез ограниченных по структуре оптимально-компромиссных стратегий обслуживания потока объектов	70
Курсін А.І. Метод нормалізації словоформ зі словником початкових форм слів без граматичної інформації	76
Любченко В.В., Шинкарьок О.С. Метод будування навчальної траєкторії в умовах мобільного навчання	81
Лысенко А.Н., Романов А.Ю. Ресурсоэффективный роутер для многопроцессорной сети на чипе	86
Марончук И.Е., Марончук И.И., Петраш А.Н. NANO-S: пакет для расчета энергетических спектров электронов в наногетероструктурах с квантовыми точками	93

Мишин Д.В., Монахова М.М. Об оптимизации администрирования корпоративных сетей передачи данных в условиях ограниченных административных ресурсов АСУП	101
Мороз Б.І., Коноваленко С.М. Деякі аспекти розвитку системи аналізу ризиків порушення митного законодавства	109
Никонов О.Я., Мнушка О.В., Савченко В.Н. Оценка точности вычислений специальных функций при разработке компьютерных программ математического моделирования	115
Плюта Н.В., Гоменюк С.І. Модель координаційної взаємодії в складній ієрархічно впорядкованій системі	122
Полякова М.Ю., Судаков Б.Н. Разработка подхода к созданию алгоритма синтаксического анализа естественно-языкового текста информационно-поисковых систем	128
Ризун Н.О. Концепция построения экспертной системы поддержки принятия решений по управлению учебным процессом в вузе	135
Самигулина Г.А., Чебейко С.В. Разработка технологии иммунносетевого моделирования для компьютерного молекулярного дизайна лекарственных препаратов	142
Субботин С.А. Методы формирования выборок для построения диагностических моделей по прецедентам	149
Судаков Б.М., Духглавов Д.Е., Володіна І.М. Метод структуризації предметної галузі в системах підтримки прийняття рішень	157
Телішевська А.В., Поворознюк А.І. Формалізація вхідної інформації для діагностики неврологічних захворювань	162
Филатова А.Е. Нелинейная фильтрация биомедицинских сигналов с локально сосредоточенными признаками в задаче структурной идентификации	168
Цветков А.И. Бикритериальная модель обслуживания потока объектов в разветвленной рабочей зоне обслуживающего процессора	175
Чопоров С.В., Гоменюк С.И. Проблемно-ориентированный язык геометрического моделирования на базе теории R-функций	181
Яцко В.А. Анализ себестоимости с использованием нечетких моделей	189

НАУКОВЕ ВИДАННЯ

ВІСНИК НАЦІОНАЛЬНОГО ТЕХНІЧНОГО УНІВЕРСИТЕТУ "ХПІ"

*Збірник наукових праць
Тематичний випуск
Інформатика і моделювання
Випуск 17*

Науковий редактор д.т.н. Дмитрієнко В.Д.
Технічний редактор к.т.н. Леонов С.Ю.
Відповідальний за випуск к.т.н. Обухова І.Б.

Дизайн та оформлення к.т.н. Гладких Т.В.

Обл.вид. № 79–11

Підп. до друку 27.05.2011 р. Формат 60х84 1/16. Папір Сору Рарег.
Друк-ризוגрафія. Гарнітура Таймс. Умов. друк. арк. 9,8.
Облік. вид. арк. 10,0. Наклад 300 прим.
Ціна договірна

НТУ "ХПІ", 61002, Харків, вул. Фрунзе, 21

Видавничий центр НТУ "ХПІ"
Свідоцтво ДК № 116 від 10.07.2000 р.

Типографія "Современная печать"
61024 Україна, г. Харьков, ул. Культуры, 22