



**System Analysis
and Information
Technology
SAIT 2018**

**May 21 – 24, 2018
Kyiv, Ukraine**



Institute for Applied System Analysis

National Academy of Sciences of Ukraine

Ministry of Education and Science of Ukraine

National Technical University of Ukraine "Igor Sikorsky Kyiv Polytechnic Institute"

Nataliya D. Pankratova (Ed.)

System Analysis and Information Technologies

20-th International Conference SAIT 2018

Kyiv, Ukraine, May 21 – 24, 2018

Proceedings



Institute for Applied System Analysis
at the Igor Sikorsky Kyiv Polytechnic Institute

UDC [519.7/.8:(004+007)](100)(06)
ББК 22.18я43+72я43
С40

Volume editor:

Nataliya D. Pankratova, Dr.Sc., Prof.

Editorial board:

Petro I. Bidyuk, Dr.Sc., Prof.

Nataliya D. Pankratova, Dr.Sc., Prof.

Anatoliy I. Petrenko, Dr.Sc., Prof.

Yuriy P. Zaichenko, Dr.Sc., Prof.

Illia O. Savchenko, Ph.D.

Revising:

Gennadii D. Kiselyov, Ph.D.

Illia O. Savchenko, Ph.D.

Oleksandr M. Terentiev, Ph.D.

Bohdan V. Chapaliuk

Design and typesetting:

Illia O. Savchenko

Mykhailo P. Makukha

System analysis and information technology: 20-th International conference
SAIT 2018, Kyiv, Ukraine, May 21 – 24, 2018. Proceedings. – ESC “IASA” NTUU “Igor
Sikorsky Kyiv Polytechnic Institute”, 2018. – 270 p.

С40 Системный анализ и информационные технологии: материалы 20-й
Международной научно-технической конференции SAIT 2018, Киев, 21 – 24 мая 2018 г. /
УНК “ИПСА” НТУУ “КПИ им. Игоря Сикорского”. – К.: УНК “ИПСА” НТУУ “КПИ”,
2018. – 270 с. – Текст: укр., рус., англ.

С40 Системний аналіз та інформаційні технології: матеріали 20-ї Міжнародної
науково-технічної конференції SAIT 2018, Київ, 21 – 24 травня 2018 р. / ННК “ІПСА”
НТУУ “КПІ ім. Ігоря Сікорського”. – К.: ННК “ІПСА” НТУУ “КПІ”, 2018. – 270 с. –
Текст: укр., рос., англ.

This book of abstracts includes issues connected with the research and development of complex systems of various nature in conditions of uncertainty and multifactor risks, Grid and high performance computing in science and education, intelligent systems for decision-making, progressive information technologies for needs of science, industry, economy, and environment. The problems of sustainable development and global threats estimation, forecast and foresight in tasks of planning and strategic decision making are investigated.

В сборнике рассматриваются вопросы, связанные с разработкой и исследованием сложных систем разной природы в условиях неопределенности и многофакторных рисков, Grid и систем высокопроизводительных вычислений в науке и образовании, интеллектуальных систем поддержки принятия решений, прогрессивных информационных технологий для потребностей науки, промышленности, экономики, окружающей среды. Исследуются вопросы устойчивого развития и оценивания глобальных угроз, прогноза и предвидения в задачах планирования и принятия стратегических решений на уровне регионов, больших городов, предприятий.

У збірнику розглядаються питання, що пов'язані з розробкою та дослідженням складних систем різної природи в умовах невизначеності та багатофакторних ризиків, нових інформаційних технологій, Grid і систем високопродуктивних обчислень в науці і освіті, інтелектуальних систем підтримки прийняття рішень, прогресивних інформаційних технологій для потреб науки, промисловості, економіки та навколишнього середовища. Досліджуються питання сталого розвитку та оцінювання глобальних загроз, прогнозу та передбачення в задачах планування та прийняття стратегічних рішень на рівні регіонів, великих міст, підприємств.

ISBN 978-617-7619-06-1



9

786177 619061

© Institute for Applied System Analysis at the Igor Sikorsky Kyiv
Polytechnic Institute, 2018

ISBN 978-617-7619-05-4 (print)

ISBN 978-617-7619-06-1 (ebook)

<http://sait.kpi.ua>

Table of contents · Содержание · Зміст

Plenary talks · Пленарные доклады · Пленарні доповіді	9
<i>Згуровський М.З., Перестюк М.М.</i> Моделирование линейного разлома цивилизаций в контексте их фундаментальных культурных відмінностей	10
<i>Погорілий С.Д., Кривий С.Л.</i> Методологія проектування застосувань в технології GPGPU	11
<i>Чижри́й А.А.</i> Методы прикладного нелинейного анализа в игровых задачах динамики	13
<i>Широков В.А.</i> Системный подход и эволюция мира	14
Section 1. System analysis of complex systems of various nature	
Секция 1. Системный анализ сложных систем разной природы	
Секція 1. Системний аналіз складних систем різної природи	16
<i>Aghaei Agh Ghamish Yaghoub, Donets G.A., Aghaei Agh Ghamish Ovi Nafas</i> Mathematical secures with different locks	18
<i>Baranovska L.V.</i> On one differential-difference game of approach with one evader and one pursuer	20
<i>Baranovska L.V., Baranovska G.G.</i> Quasilinear group pursuit differential games with delay	22
<i>Dychko A.O., Yermeev I.S.</i> Risk-oriented management of water ecosystems	23
<i>Kondratova L.P.</i> Application of Box-Jenkins models for forecasting the guaranteed functioning of complex technical systems in real-world conditions	24
<i>Malysheva M.O., Akhmedova D.N., Prymak I.K.</i> Method of modification of base graph-logical models	26
<i>Osaulevko V.M.</i> Model for sequence prediction based on dendritic spatiotemporal integration	28
<i>Shaptala R.V., Kyselova A.G., Kyselov G.D.</i> Evaluation of approaches to emotion classification	30
<i>Siryk S.V.</i> On the application of the mass lumping correction technique for convection-diffusion problems	31
<i>Soloviev V.N., Romanenko Y.V.</i> Quantum econophysics of bitcoin crises	33
<i>Ved O.V.</i> Basic Concepts of Catalytic Exhaust Gas Aftertreatment	35
<i>Бахрушин В.Є., Поджоваліхіна О.О., Логвіненко В.О.</i> Задача розподілу інвестицій в умовах статистичної невизначеності	36
<i>Беседовский М.О.</i> Модернизация маркетинговой стратегии Котлера на этапе рыночного спада в Украине	37
<i>Болдак А.О., Єфремов К.В., Пишинограєв І.О., Горюх О.С.</i> Використання результатів SWOT-аналізу в сценарному моделюванні	38
<i>Васильев В.И., Вишталъ Д.М., Любашенко Н.Д.</i> Алгоритм синтеза топологии сети по критерию доступности сетевой услуги	39
<i>Власов М.Д., Болдак А.О.</i> Оптимізація конфігурацій розподілених інформаційних систем	41
<i>Волкова В.Н., Логинова А.В., Черный Ю.Ю., Леонова А.Е.</i> Методика сравнительного анализа и выбора технологических инноваций четвертой промышленной революции	43
<i>Волошин О.Ф., Маляр М.М.</i> Нечітке математичне моделювання фінансових ризиків інноваційних проєктів	46

Годна А.В., Жданова О.Г., Сперкач М.О. Про один підхід складання розкладу виконання завдань паралельними пристроями з метою мінімізації сумарного відхилення моментів завершення від директивних строків	47
Горелова Г.В. О разработке инструментария когнитивного моделирования социально-экономических систем	49
Городецкий В.Г., Осадчук Н.П. Численно-аналитический метод решения обратной задачи для систем одного класса	51
Гришин И.Ю., Тимиргалева Р.Р. Метод автоматизации управления процессом обнаружения движущихся объектов радиолокационным комплексом	52
Денисенко А.И. Оптимизация геометрических параметров теплообменных элементов газовых котлов	53
Духновская К.К., Ковтун О.И. Динамическая мера TF-IDF	54
Євграфова К.Л., Бідюк П.І. Модифікація підходу до формування ціни у роздрібній торгівлі	56
Жданова О.Г., Маленко А.О., Сперкач М.О. Складання розкладу виконання завдань паралельними пропорційними пристроями з метою мінімізації сумарного відхилення від спільного директивного строку	58
Заводник В.В., Косташ Т.В. Оценивание состояния демографических процессов на Украине	61
Зайченко А.Є., Савченко І.О. Система підтримки прийняття рішень на основі мережі морфологічних таблиць для аналізу лісової пожежі	63
Касумова Т.А. Роль и место космической информации в формировании общества и культуры информационного мира	65
Кірік О.Є. Системний підхід до формування управлінських рішень для великих ієрархічно структурованих систем	67
Кісельова О.М., Притоманова О.М., Падалко В.Г. Про оптимізацію параметрів нейронечіткої моделі експортних відносин між Україною та Китаєм	68
Комлев О.О. «Земна поверхня»: структура і планетарна роль	69
Кравчук Є.С., Сергеев-Горчинський О.О. Методи оцінки семантичної орієнтації тексту	70
Крамов А.А. Экстракция структурированных данных с множества HTML-страниц	72
Кулик В.В., Кудін Г.І., Коробова М.В. Інтерактивне управління відтворювальними процесами в національній економіці на основі системи спрощених балансів	74
Куценко А.С., Коваленко С.В., Товажнянский В.И. Обращение динамических систем для некоторых классов сигналов	75
Лабуткина Т.В. Комплекс математических моделей и методов для прогноза и анализа поликонфликтных сближений орбитальных объектов	76
Мілявський Ю.Л. Порівняння методів класичного та нечіткого керування в імпульсних процесах когнітивних карт	78
Назарага І.М. Матрична множинна регресія та класичні методи біометрії для прогнозування біологічних показників: приклади	79
Нестеренко В.И. Разработка стратегии развития группы демографических реестров в системе государственных реестров Украины	80
Панкратов В.А. Оценивание промышленного сектора Украины на основе методологии когнитивного моделирования	81
Панкратова Н.Д., Маняк Ю.В. Нечітко-множинний підхід в задачах когнітивного моделювання складних систем	83
Панкратова Н.Д., Сльота М.Р. Гарантоване функціонування складних технічних систем з урахуванням ресурсу допустимого ризику	85
Романкевич О.М., Манілевич Д.Ф. Генератор двійкових векторів постійної ваги для виконання статистичних експериментів	86
Слюсар А.В., Данилов В.Я. Система прийняття рішень на основі вейвлетної ідентифікації хвиль Елліотта	87

Тимофієва Н.К. Задачі комбінаторної оптимізації, цільова функція в яких уведена на нескінченній комбінаторній множині	89
Фатенко В.В. Узагальнення задачі Реньї про випадкове паркування	91
Цегелик Г.Г., Краснюк Р.П. Моделювання оптимального розміщення баз даних інформаційних систем за наявності серверів проміжного зберігання даних	93
Циганок В.В., Ройк П.Д. Визначення узгодженості оцінок експертів при підтримці прийняття групових рішень	96
Чернуха О.Ю., Білуцак Ю.І., Чучвара А.Є. Програмний комплекс для моделювання дифузії у тілі з пастками за каскадного розпаду мігруючих частинок	98
Шибирин А.Р., Савченко І.О. Система підтримки прийняття рішень на основі двохетапного модифікованого методу морфологічного аналізу	100
Шудренко Є.О., Болдак А.О. Практична реалізація методу Делфі	101

Section 2. Intelligent systems for decision-making

Секция 2. Интеллектуальные системы принятия решений

Секція 2. Інтелектуальні системи прийняття рішень 104

Chertov O., Rudnyk T. Search of phony accounts on Facebook and Twitter	106
Chertov O.R., Shanin V.O., Bobyr A.O. Patterns of glycemia behavior for prediction of nocturnal hypoglycemia	107
Honcharevskyi D., Chertov O. Violating group anonymity using machine learning techniques	108
Lazuka V.T., Rodionov A.M. Combining Several Techniques of Adversarial Example Generation	109
Naderan M., Zaychenko Y.P. Diagnosing Lung Cancer Based on Deep Learning Algorithms: Review	111
Nikolaiev S.S., Tymoshenko Y.O. RBF neural network probabilistic skin color detector for real-time video applications	113
Onanko A.M. Prediction of Russian content on Ukrainian television	115
Tkachenko D.A. Modeling spatial information with convolutional and recurrent neural networks for video classification	116
Tkachenko D.A., Kharchenko K.V. Deep learning approach to audio analysis	118
Zaychenko Y., Galib H. Inductive Modeling Method GMDH in the Problems of Big Data Analysis	120
Акимов В.С. Алгоритми багаторівневого навчання для класифікації захворювань шкіри	121
Бакурова А.В., Терещенко Е.В., Нагулов С.В. Логічна класифікація панельних даних	123
Буценко Ю.П., Лабжинський В.А. Нейромережеві алгоритми розпізнавання та класифікації надзвичайних ситуацій на об'єктах критичної інфраструктури	124
Галушко М.О. Класифікація зображень за допомогою нейронних мереж на GPU	125
Гусак О.В., Семендяк Є.С. Адаптивний пошук гіперпараметрів за допомогою підходів машинного навчання	126
Дмитренко О.О., Ланде Д.В. Алгоритм пошуку оптимального сценарію впливу вершин у когнітивних картах заснований на принципі Парето	127
Доманецька І.М., Хроленко Я.О. Нейромережеві технології у задачах автоматизації контролю знань для on-line сервісів відкритої освіти	129
Ерохин Г.В. Прогнозирование курса криптовалют на основе анализа тональности новостей и записей в социальных сетях	130
Жданова О.Г., Попенко В.Д. Економічна модель прийняття рішень щодо народження дітей у сім'ї.	131
Забелин С.И. Интеллектуальный анализ вулканических газов с помощью интеллектуальных методов анализа данных	134
Зайченко Е.Ю., Ови Нафас Агаи аг Гамиш Анализ и оптимизация компьютерных сетей нового поколения в условиях неопределенности	135
Караюз І.В., Бідюк П.І. Інтегрована модель для прогнозування макроекономічних процесів	136

<i>Клімова О.В.</i> Методи прогнозування рекламної активності на телебаченні України . . .	138
<i>Кондратенко Н.Р., Снігур О.О.</i> Використання інтервальних нечітких множин типу 2 в задачах ідентифікації складних об'єктів	139
<i>Кузнецова Н.В., Фомін О.В.</i> Прогнозування ризику втрати користувачів онлайн-платформи	141
<i>Луданов Д.К., Зайченко Ю.П.</i> Нейронні мережі та їх застосування	143
<i>Павлов Д.Г., Чертов О.Р.</i> Використання нечіткої логіки для визначення мережевого шахрайства на прикладі шаблону інформаційної атаки	144
<i>Поворознюк А.И., Филатова А.Е., Шехна Х.</i> Согласованная морфологическая фильтрация полутоновых изображений на основе фрактального анализа	146
<i>Пономаренко Р.М.</i> Нечеткий логический вывод на основе многоуровневого параллелизма	148
<i>Путренко В.В., Пащинська Н.М., Назаренко С.Ю.</i> Використання 3D Space-Time Cube для інтелектуального аналізу даних	150
<i>Роговий А.В., Бідюк П.І.</i> Побудова скорингових моделей для оцінки платоспроможності клієнтів телекомунікаційного оператора	152
<i>Савченко Д.А.</i> СППР для прогнозирования уровня продаж торговой фирмы	154
<i>Тимошук О.Л., Дорундяк К.М.</i> Порівняльний аналіз методів прогнозування банкрутства підприємств України	155
<i>Фіногенов О.Д., Грачова О.А.</i> Аналіз даних сайтів погоди на основі методу аналізу ієрархій	156
<i>Чапалюк Б.В., Зайченко Ю.П.</i> Автоматична медична діагностика на базі зображень комп'ютерної томографії	157
<i>Чечельницкая А.Б.</i> Прогнозирование цен на авиабилеты бюджетной авиакомпании Ryanair с помощью нейронных сетей и авторегрессии со скользящим средним . . .	158
<i>Чечула М.В., Погорілий С.Д.</i> Розробка методів роботи з BigData на мобільних пристроях з використанням технології GPGPU	159
<i>Шулькевич Т.В., Селін Ю.М.</i> Математичний апарат для інтелектуального аналізу і прогнозування нелінійних нестационарних процесів. Досвід застосування	160

Section 3. High Performance Computing and Microsystems Engineering

Секция 3. Высокопроизводительные вычисления и разработка микросистем

Секція 3. Високопродуктивні обчислення та розробка мікросистем 162

<i>Pechurin N.K., Kondratova L.P., Pechurin S.N.</i> Leontiev's model for Internet of things information loading estimation	164
<i>Sergiyenko A.M., Hasan M.J., Sergiyenko P.A.</i> Square root calculations in FPGA	165
<i>Vinogradov Ju.N., Sergiyenko A.M., Quadir S.H.</i> Minimized FIR Filter Design Implemented in FPGA	167
<i>Галета П.О.</i> Застосування технології Blockchain у Mesh мережах на прикладі розподіленого DNS-сервісу	169
<i>Івченко Д.А.</i> Використання Neo4j і GraphQL на базі GraphDB для створення мікросервісу	171
<i>Кириуша Б.А., Колінко А.М.</i> Взаємодія мультиагентної системи IoT девайсів з виділеним сервером	175
<i>Круш І.В., Михалько В.Г.</i> Адаптивний точковий пошук з використанням методів машинного навчання	177
<i>Лозінський А.П.</i> Багаторівнева архітектура хмарної платформи	179
<i>Назарова І.А.</i> Масштабованість паралельних розподілених обчислень на основі ізоєфективного аналізу	181
<i>Орехов О.А., Орехова Н.А.</i> Порівняльний аналіз Blockchain-архітектури на основі дерев Меркле та протоколу PHANTOM на основі архітектури BlockDAG	183
<i>Останчук Я.М.</i> Стратегія розвитку CRM систем як сервісу у хмарі Microsoft Azure . .	184
<i>Сапсай Т.Г., Довганюк А.О.</i> О последовательном диагностировании многопроцессорных систем	186
<i>Слинько М.С.</i> Формалізоване проектування застосувань для графічних відеоадаптерів	187

Section 4. Progressive information technologies**Секция 4. Прогрессивные информационные технологии****Секція 4. Прогресивні інформаційні технології****190**

<i>Basyuk T.M.</i> Popularization of Internet resources by using “featured snippets”	192
<i>Kalinovsky Ya.A., Boyarinova Yu.E., Khitsko Ya.V., Sukalo A.S.</i> Principles of constructing algorithms for processing digital signals using hypercomplex number systems	194
<i>Kasyanchuk I.V.</i> Education System to Improve Remembering of Information	196
<i>Kopp A.M., Orlovskyi D.L.</i> An approach to measure similarity of business process models	198
<i>Kurylenko O.M.</i> Credit scoring models development	200
<i>Moroz A.M., Globa L.S.</i> Data Mining and its applications	202
<i>Olefir O.S., Chernenko V.M.</i> System of logistics of material resources based on social networks	204
<i>Sergiyenko A.M., Qasim M.R.</i> Modeling of the wave propagation in the solid bar	205
<i>Tarasenko-Klyatchenko O.V., Tuparieva V.A.</i> Algorithm of authentication of objects of the network on the basis of analysis of signal parameters on the physical level	207
<i>Terentiev O.M., Makohon R.O.</i> Analysis of causality Bayesian network usage for data-mining purposes	208
<i>Арчвадзе Н.Н., Пховелишвили М.Г.</i> Применение параллельных данных для прогнозирования сложных процессов	210
<i>Бакун С.А., Терентьев О.М.</i> Методи перетворення категоріальних змінних в числові	211
<i>Васильев В.І., Вішталъ Д.М., Любашенко Н.Д.</i> RDF-модель порталу підтримки розміщення і використання даних у Всесвітній мережі	213
<i>Видолоб А.В.</i> Забезпечення інформаційної безпеки системи інтернету речей	215
<i>Виклюк Я.І., Гусак О.М.</i> Шляхи підвищення ефективності протипожежного моніторингу лісу	216
<i>Гіоргізова-Гай В.Ш., Шеренковський А.О.</i> Шлюзи в системах IoT	217
<i>Гнатенко В.Ю., Ситников В.С., Ступень П.В.</i> Электрическая модель с идеальными элементами для поиска максимального потока по транспортной сети	219
<i>Городько Н.А., Бояринова Ю.Е.</i> Адаптивные системы организационного обеспечения критических инфраструктур	221
<i>Грушенко К.Ю., Петрашенко А.В.</i> Модифікований алгоритм виявлення руху з камер відеоспостереження	222
<i>Зам'ятін Д.С., Книш П.К.</i> Визначення місцеположення джерела акустичного сигналу	223
<i>Зам'ятін Д.С., Перевертайло А.І.</i> Розподіленні обчислення координат місцезнаходження об'єктів за акустичним сигналом у мережі сенсорів	224
<i>Капшук О.О.</i> Використання біометричних технологій при дистанційному навчанні в мережевих інформаційно-освітніх середовищах	225
<i>Киричек Г.Г., Курай В.І.</i> Система розпізнавання об'єктів	226
<i>Кінда В.В.</i> Модель оцінки страхового випадку із використанням розподілу Твіді	228
<i>Кісілючук Б.Я., Тесленко О.К.</i> Ключі алгоритму шифрування на базі підстановок довільної розрядності	230
<i>Краштан Т.Є.</i> Аналіз існуючих алгоритмів колоризації зображень	232
<i>Кривоніс А.В., Волонтир О.О.</i> Поєднання Machine Learning та VR	234
<i>Кузнецова Ю.А., Цешевський В.В.</i> Аналіз застосування метапрограмування для проекту Mikz Market	235
<i>Лимар Б.О., Олефір О.С.</i> Метод виявлення паркувальних місць в інтелектуальній системі паркування	237
<i>Лисецкий Ю.М.</i> Защита информации: системы резервного копирования	238
<i>Литвинов В.А., Майстренко С.Я., Хурицлава К.В., Костенко С.В.</i> Оценка контролируемых и корректирующих свойств референтного словаря в системе обнаружения и исправления типовых ошибок пользователя	240
<i>Литвинюк А.А., Левчук С.Б.</i> Система автоматичної оптимізації асортименту в ритейлі	241

<i>Логін В.В., Бідюк П.І.</i> Моделювання та прогнозування характеристик трафіку цифрової Інтернет-реклами	243
<i>Лютенко І.В., Курасов О.І., Соколов Д.В.</i> Використання методів послідовного ранжування альтернатив для оцінювання веб-сайтів	245
<i>Максим К.Є.</i> Дослідження вразливостей спекулятивного виконання інструкцій на процесорах ARM	246
<i>Матвиенко Ю.А., Зинченко А.А.</i> Разработка программного обеспечения для оптимизации методики алергодиагностики к стоматологическим анестетикам	247
<i>Огурцов М.І.</i> Розробка протоколу захищеного обміну даними для спеціальних мереж .	249
<i>Орлова М.М., Шитий Д.В.</i> Способи та засоби підтримки багатопотокої передачі в комп'ютерних мережах	251
<i>Пашаева М.М.</i> Классификация различных типов растительного покрова на основе космических снимков и технологии ГИС	252
<i>Петрішенко С.О.</i> Сучасна концепція Інтернету речей	256
<i>Прогонов Д.О.</i> Теоретико-інформаційні оцінки стійкості методів UNIWARD до стегоаналізу	258
<i>Продан А.О.</i> Оптимізація пошуку та індексації документів формату XML	259
<i>Романкевич В.О., Липка Т.Б.</i> Покращення якісних характеристик стеганосистем з використанням зображень в якості контейнера	260
<i>Романов В.В., Шахун О.В.</i> Інформаційні технології в освіті	262
<i>Северін С.І., Тесленко О.К.</i> Вибір алгоритму ущільнення для багаторозрядних підстановок	264
<i>Сопілков М.Р.</i> Прогнозування виникнення збройних конфліктів з використанням ймовірнісно-статистичних методів	266
Authors · Авторы · Автори	268

Plenary talks

System analysis of
complex systems of
various nature

Intelligent systems for
decision-making

High Performance
Computing and
Microsystems Engineering

Progressive information
technologies

Plenary talks



Згуровський М.З., Перестюк М.М.

КПІ ім. Ігоря Сікорського, Київ, Україна

Моделювання ліній розлому цивілізацій в контексті їх фундаментальних культурних відмінностей

Метою доповіді є продовження досліджень, проведених у 2008 році [1], з часовим інтервалом 10 років, які ґрунтуються на застосуванні кількісного і якісного аналізу відмінностей між світовими культурами на основі підходу С. Хантінгтона [2].

В процесі проведення зазначеного моделювання виконано наступні етапи [1]:

- 1) Уточнення цивілізаційного розподілу, запропонованого С. Хантінгтоном;
- 2) Розробка системи критеріїв для оцінки культурних відмінностей між цивілізаціями;
- 3) Розробка і проведення експертного оцінювання цивілізаційних відмінностей;
- 4) Обробка даних оцінювання і розрахунок значень ліній “розломів” між цивілізаціями.

В результаті уточнення цивілізаційного розподілу 192 країни світу (члени ООН) було згруповано в 12 цивілізацій: 1. Західна-північноамериканська; 2. Ісламська-арабська; 3. Слов'янська-східноправославна; 4. Західна-європейська; 5. Ісламська-тюркська; 6. Слов'янська-західнокатолицька; 7. Конфуціанська; 8. Ісламська-малайська; 9. Латиноамериканська; 10. Японська; 11. Індуїстська; 12. Африканська.

За підсумками роботи групи експертів було сформульовано 8 базових критеріїв, які найбільш повно характеризують культурні відмінності цивілізацій: **1. Цінність життя людини** (діапазон: “Життя людини нічого не варто” – “Життя людини – найвища цінність”); **2. Свобода особистості в суспільстві** (ступінь свободи переміщення, свободи вираження поглядів, свободи в особистому житті і т.п.); **3. Статус жінки в суспільстві** (діапазон: “Повне домінування чоловічої статі” – “Тендерний паритет” – “Повне домінування жінки”); **4. Ступінь проникнення релігії в життя** (діапазон: “Релігійні та церковні інститути взагалі не впливають на життя людей” – “Релігійні та церковні інститути визначальним чином впливають на життя людей”); **5. Етнічна однорідність** (ступінь толерантності міжетнічних відносин всередині цивілізації); **6. Відкритість і традиціоналізм в культурі і мисленні** (схильність до змін в традиціях і способах мислення); **7. Радикалізм в політичному житті** (стабільність політичного життя).

Для проведення експертного оцінювання культурних відмінностей між дванадцятьма цивілізаціями на основі використання вищенаведених восьми критеріїв була створена група з 14 експертів, які мали досвід міжнародної діяльності в групах країн з вищевказаними цивілізаціями. Експерти оцінювали культурні відмінності між цивілізаціями (розломи) для кожного з критеріїв, привласнюючи кількісні значення цим розломів, з використанням системи оцінок за шкалою Міллера. На основі визначення цих оцінок, за представленими 8 критеріями, були сформовані профілі відмінностей між цивілізаціями, і визначені “зближення” і “розломи” між ними.

Отримані результати свідчать про існуючі культурні відмінності між різними цивілізаціями. Потенційні конфлікти можуть відбуватися між цивілізаціями, перш за все по лініях розломів, кількісні значення яких є найбільшими. І навпаки: потенційні об'єднання цивілізацій можуть відбуватися по лініях розломів, кількісні значення яких є найменшими.

Використовуючи результати кластеризації відстаней, чисельні значення загальних розломів між цивілізаціями і сукупні відмінності цивілізацій, побудовано перелік можливих об'єднань і конфліктів між світовими культурами.

Показано, що з точки зору протистоянь, найбільшу схильність до них мають Західний і Ісламський блоки цивілізацій. При цьому, Західний і Слов'янський блоки мають великі розломи з Ісламською, Індуїстською і Африканською цивілізаціями. Схильність до протистояння між Африканською і Індуїстською цивілізаціями можна віднести на рахунок сильного, до сих пір, впливу таких явищ, як расизм і сегрегація.

Література. 1. Michael Zgurovsky, Alexis Pasichny. Modelling of the civilizations' break lines in context of their fundamental cultural differences // System Research and Information Technologies, No1, 2009. – 18 p. 2. Хантінгтон С. Зіткнення цивілізацій // М., Поліс. – 1994. – №1. – С. 33–48.

Погорілий С.Д., Кривий С.Л.

Київський національний університет ім. Тараса Шевченка, Київ, Україна

Методологія проектування застосувань в технології GPGPU

Проведено аналіз тенденцій розвитку архітектур новітніх відеоадаптерів від компанії NVIDIA. Відзначено складності проектування застосувань в технології GPGPU, на основі чого обґрунтовано необхідність застосування формалізованих методів. Запропоновано сукупність методів розв'язання цієї задачі.

Вступ. Сьогодення високопродуктивних обчислень лежить в галузі систем з масовим паралелізмом. Існує широкий клас таких систем, але в останнє десятиріччя все більше починають домінувати ті, що ґрунтуються на технології GPGPU (General Purpose Graphics Processing Unit). Зростаюча складність програмного і технічного забезпечення таких IT-інфраструктур потребує розробки методів їх проектування, обґрунтування та коректної реалізації.

Апаратна платформа графічного процесору (GPU) від Nvidia складається з набору поточкових мультипроцесорів SM (streaming multiprocessor) [1–3]. Блоки потоків розподіляються між SM у такий спосіб, що всі потоки в межах блоку виконуються на одному мультипроцесорі, і в той же час, одному SM може бути поставлено у відповідність більше, ніж 1 блок. В свою чергу, мультипроцесор розділяє блок на групи по 32 потоки. Ця операція виконується для кожного блоку окремо, а утворені групи називаються варпами (warp). Всі потоки в межах варпу в один момент часу виконують одну і ту саму інструкцію. Таким чином, фізично паралельно виконуються лише потоки в межах одного варпу. Будь-яке розгалуження в інструкціях (умовні оператори тощо) призводить до того, що кожен шлях виконання буде обчислюватися окремо і послідовно.

Важливою особливістю компанії Nvidia є та, що їй притаманна висока динаміка модернізації існуючих і створення нових архітектур відеоадаптерів. На цей день найбільш продуктивною є архітектура Volta [6] – це найпотужніша в світі архітектура GPU, покликана стати каталізатором нових хвиль досягнень в області штучного інтелекту і високопродуктивних обчислень. Перший процесор на базі Volta – це GPU для дата-центрів Tesla V100, який забезпечує надвисоку швидкість і масштабованість навчання глибоких нейронних мереж, а також прискорює високопродуктивні і графічні обчислення. В основі Volta знаходиться 21 млрд. транзисторів, що забезпечують продуктивність в завданнях глибокого навчання, еквівалентну 100 CPU. Пікова продуктивність Volta в 5 разів вище архітектури Pascal – поточної графічної архітектури NVIDIA, і в 15 разів вище Maxwell, представленої два роки тому (ці цифри вчетверо більше того, що передбачав закон Мура).

Складність IT-інфраструктур, побудованих на технології GPGPU (багаторівнева структура пам'яті, велика кількість ядер GPU тощо), унеможливорює інтуїтивне проектування і вимагає наявності формалізованого математичного апарату і методів представлення застосування. В роботі надано аналіз таких методів.

Розглянуто зручний математичний апарат систем алгоритмічних (мікропрограмних) алгебр (САА) Глушкова. Однак, хоч САА і дозволяють досягнути структуру застосування і зручно проводити еквівалентні трансформації схем алгоритмів застосування без прив'язки до апаратної архітектури, в них відсутня можливість поглибленого аналізу застосування, що пояснюється достатньо високим рівнем абстрагування схем алгоритмів, що проектуються.

Відзначено переваги застосування мереж Петрі (звичайних, темпоральних та кольорових). Важливою перевагою такого підходу є наявність програмної підтримки процесу проектування за допомогою спеціальних пакетів прикладних програм.

В роботі запропоновано підхід до розробки систем GPGPU на основі розмічених транзицій-

них систем (РТС) з використанням різного типу їх композицій на вибраному рівні абстракції [4]. Отримана в такий спосіб РТС аналізується, корегується, перевіряються її очікувані властивості. Перевагою підходу є можливість редукції РТС до мережі Петрі. Пропонована техніка апробована прикладами синтезу сервісів ІТ-інфраструктур [5] типу PaaS (Platform as a Service), IaaS (Infrastructure as a Service), нейроподібних мереж для задач класифікації, моделювання паралельних та розподілених систем тощо [5].

Література. 1. С.Д. Погорілий, Д.Ю. Вітель, О.А. Верещинський. Новітні архітектури відеоадаптерів. Технологія GPGPU. Частина 1. Реєстрація, зберігання і обробка даних, 2012, – Т.14, – №4. 2. С.Д. Погорілий, Д.Ю. Вітель, О.А. Верещинський. Новітні архітектури відеоадаптерів. Технологія GPGPU. Частина 2. Реєстрація, зберігання і обробка даних, 2013, – Т.15, – №1. 3. С.Д. Погорельий, Ю.В. Бойко, М.И. Трибрат, Д.Б. Грязнов. Анализ методов повышения производительности компьютеров с использованием графических процессоров и программно-аппаратной платформы CUDA. Математичні машини і системи, 2010, №1. 4. С.Л. Кривий, В.М. Опанасенко, С.Д. Погорілий, С.Ф. Теленик. Композиційний метод проектування ІТ-інфраструктур з використанням розмічених транзиційних систем. Сучасна інформатика: проблеми, досягнення та перспективи розвитку. Тези доповідей Міжнародної наукової конференції, присвяченої 60-річчю заснування Інституту кібернетики імені В.М. Глушкова НАН України. Україна, Київ, 13–15 грудня 2017 р. 5. S.L. Kryvyi, Y.V. Boyko, S.D. Pogorilyy, O.F. Boretskyi, M.M. Glybovets. Design of Grid Structures on the Basis of Transition Systems with the Substantiation of the Correctness of Their Operation. Cybernetics and Systems Analysis. January 2017, Volume 53, Issue 1, pp.105–114. Springer Science+Business Media New York 2017. 6. NVIDIA Volta и другие анонсы NVIDIA для ИИ – GTC 2017. [Електронний ресурс]. – Режим доступу: <http://www.nvidia.ru/object/nvidia-ai-revolution-platform-ru.html>.

Чикрий А.А.

Институт кибернетики им. В.М.Глушкова НАН Украины, Киев, Украина

Методы прикладного нелинейного анализа в игровых задачах динамики

Рассматриваются конфликтно управляемые процессы о сближении траектории с заданным терминальным множеством за конечное время. Исходным объектом для исследования является представление решения динамической системы, что позволяет в единой схеме рассмотреть процессы, описываемые обыкновенными дифференциальными, интегральными, интегро-дифференциальными, дифференциально-разностными уравнениями, уравнениями с дробными производными, импульсными системами. Особенность представления состоит в том, что блок начальных данных отделен от блока управления. Такая ситуация имеет место, в частности, в случае квазилинейных процессов.

Используя аппарат многозначных отображений, их селекторов и технику выпуклого анализа, получены достаточные условия завершения игры на основе метода разрешающих функций [1,2]. Эти скалярные функции естественным образом оценивают качество принятых решений, что полностью соответствует нашим представлениям об игре. Предложенный аналитический метод развит и в случае матричных разрешающих функций [3], введены верхние и нижние разрешающие функции различных типов [4], что существенно расширяет возможности метода. В рамках разработанной идеологии решены задачи группового и поочередного преследования, охвачен случай фазовых ограничений [2]. Для иллюстрации возможностей метода приведены многочисленные примеры игровых ситуаций.

Литература. 1. Chikrii A.A. On analytical method in dynamic pursuit games // Proceedings of the Steklov Institute of Mathematics. – 2010. – Vol. 271. – P. 69–85. 2. Chikrii A.A. Conflict controlled processes. Springer Science and Business Media. 2013, 424 p. 3. Chikrii A.A., Chikrii G.Ts. Matrix resolving functions in game problems of dynamics // Proceedings of the Steklov Institute of Mathematics. 2015. Vol. 291, suppl. 1. – P. 56–65. 4. Чикрий А.А. Верхняя и нижняя разрешающие функции в игровых задачах динамики // Труды ИММ УрО РАН. – 2017. – № 1. – С. 293–305.

Широков В.А.

Украинский языково-информационный фонд НАН Украины, Киев, Украина

Системный подход и эволюция мира

1. Большинство мыслящих людей убеждено, что мир, в котором мы живем, развивается, эволюционирует. Постоянно обновляющиеся данные науки не опровергают, а лишь подтверждают данное убеждение. Однако в каком направлении происходит эта эволюция, каким законам она подчиняется и вообще, существуют ли такие общие закономерности, действующие на всём, так сказать, диапазоне Мироздания, и к чему в конце концов это может привести? – с ответом на эти вопросы дело обстоит гораздо сложнее.

2. Автор исходит из внутреннего убеждения, что Мироздание не случайно, а эволюция является одним из основных свойств Мира, возможно – вообще основным. Мы полагаем, что весь объем Универсума – от его Начала (и даже перед) до самого Конца (и даже после) – управляется некими общими закономерностями, которые проявляются в формах, специфичных для каждого эволюционного этапа, каждой фазы эволюции.

3. Такой подход диктует необходимость рассматривать весь Мир в его целостности – “без исключений и изъятий” – Мир с большой буквы, что создает для автора массу проблем, как принципиального, так и технического характера.

4. В работе [1] были исследованы свойства эволюции Мира, исходя из довольно общих предположений. Данная работа призвана продемонстрировать, что в их основе лежат некоторые системные представления. Наше определение понятия системы, которому мы следуем в данной работе, можно изобразить в виде символического равенства:

$$C = C_1 + C_2 + C_3,$$

где “С” левой части обозначает понятие “система”, а правая часть демонстрирует наличие и взаимодействие основных образующих компонент этого понятия, а именно C_1 “структуру”, C_2 “субстанцию” и C_3 “субъект”.

5. Мы полагаем, что “эволюция” является фундаментальным, первичным понятием (сравнимо с понятием множества в математике), которому строгое определение дать невозможно. Но можно привести примеры, дать некоторые пояснения, исследовать её свойства и установить закономерности, чем мы и предполагаем сейчас заняться.

6. Основным понятием общей теории эволюции мы полагаем понятие “**механизма эволюции**”. Краткое содержание данного понятия раскрывается в следующих утверждениях:

- а) эволюция мира осуществляется в определенных *дискретных* формах. Эти дискретные формы мы и будем называть “*механизмами эволюции*”;
- б) каждый *механизм эволюции* характеризуется определенным, специфичным для него комплексом материально, “субстанциально” выраженных признаков, позволяющих его индивидуализировать и выделить среди других механизмов;
- в) однако все *механизмы* имеют некоторые общие свойства, что позволяет их типизировать и отнести к одному классу явлений, а именно – к классу “ЭВОЛЮЦИЯ”.

7. В данной работе мы выделяем и будем рассматривать (с разной степенью детализации) следующие эволюционные механизмы:

1. Физический.
2. Химический.
3. Генетический.
4. Нейронный.
5. Коммуникационный.
6. Производственный, состоящий из следующих:
 - 6.1. Производственно-потребительский.
 - 6.2. Производственно-обменный.
 - 6.3. Производственно-финансовый.
 - 6.4. Производственно-финансово-информационный.
 - 6.5. Производственно-финансово-информационно-сетевой.

} экономические

Отметим, что каждый из перечисленных механизмов обладает своей “тонкой структурой”. Здесь мы привели лишь отдельные этапы производственного механизма; в статье [1] отмечены 10 этапов физического механизма.

8. Содержание основного постулата общей теории эволюции сводится к следующему утверждению: причиной и двигателем эволюционных процессов является **сложность** эволюционирующей системы. Это означает, что увеличение **сложности** системы является необходимым условием её эволюции: если система эволюционирует, то ее **сложность** увеличивается.

9. В данной работе мы исходим из представления, что сложность является не просто оценочной характеристикой, а фундаментальным объективным свойством явлений и процессов. По нашему мнению, существует глубокая аналогия между определениями понятия сложности, с одной стороны, и энергии, также являющейся универсальной объективной характеристикой вещей, с другой. Оба отмеченных понятия – и сложность и энергия – являются исключительно многоаспектными и по-разному проявляются в различных ситуациях.

10. В работе приведены данные по динамике сложности для различных организмов и их сообществ, реализующейся посредством трех механизмов эволюции (генетического, нейронного и коммуникационного). Их мерами сложности выступают, соответственно, объем генома (для генетического), объем нейронной системы (для нейронного) и объем “лексической” системы (для коммуникационного механизма).

11. Все механизмы эволюции имеют общую черту: разворачивание любого из них приводит к усложнению эволюционирующей системы. При этом каждый эволюционный механизм – механизм эволюции через усложнение системы – имеет свой естественный предел. Он “запрограммирован” на достижение некоего, характерного именно для данного механизма максимума сложности. Если эволюционирующая система подходит к этому пределу (ее ресурсы усложнения приближаются к исчерпанию), она попадает в так называемую “ловушку сложности”, вступает в критическую фазу и пытается “найти” другие механизмы, способные обеспечить ее усложнение, и следовательно – продолжить процесс эволюции.

Темпоральная динамика всех механизмов эволюции имеет общий характер, схематически изображенный на рис. 1.

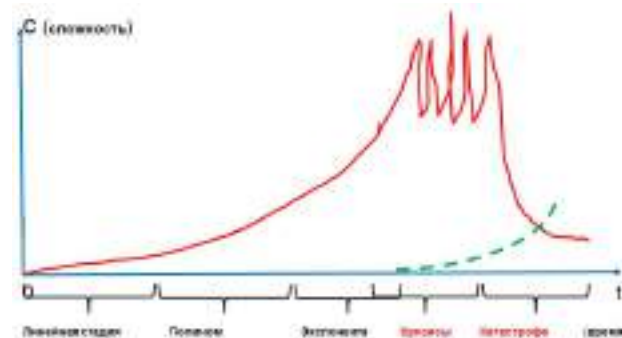


Рис. 1. Обобщенная темпоральная динамика механизма эволюции

12. Таким образом, процессы эволюции действительно демонстрируют системные черты в духе сформулированного нами понятия системы. В то же время наивно представлять процесс эволюции как постоянное, механическое увеличение соответствующей сложности. Параллельно действуют как эволюционные, так и антиэволюционные процессы. В процессе эволюции могут возникать тупиковые ветви. Сама сложность и ее динамика обладает определенной структурой. Более подробно эта феноменология исследована в работе [1] в которой исследованы также и процессы эволюции человеческих сообществ.

Литература. 1. В.А.Широков. Эволюция как универсальный естественный закон (Прологомены к будущей общей теории эволюции). Бионика интеллекта. ISSN 0555 2656. Часть I. №1(88), 2017. Часть II. №2(89), 2017. Часть III. №1(90), 2018.

Plenary talks

System analysis of
complex systems of
various nature

Intelligent systems for
decision-making

High Performance
Computing and
Microsystems Engineering

Progressive information
technologies

1

System analysis of complex systems of various nature



Section 1

System analysis of complex systems of various nature

1. System analysis methods for complex systems of various nature in conditions of uncertainty and risks.
2. Mathematical methods, models and technologies for complex systems' research.
3. Technology foresight system methodology in problems of planning and strategic decisions' making.
4. Theory and methods of optimal decision-making.
5. Problem oriented methods of complex systems' analysis and designing in conditions of uncertainty and risks.
6. Nonlinear problems of system analysis.
7. System methodology of sustainable development.
8. Text analytics and natural language processing.

Секция 1

Системный анализ сложных систем разной природы

1. Методы системного анализа сложных систем разной природы в условиях неопределенности и рисков.
2. Математические методы, модели и технологии исследования сложных систем.
3. Системная методология технологического предвидения в задачах планирования и принятия стратегических решений.
4. Теория и методы принятия оптимальных решений.
5. Проблемно-ориентированные методы анализа и проектирования сложных систем в условиях неопределенности и рисков.
6. Нелинейные задачи системного анализа.
7. Системная методология устойчивого развития.
8. Текстовая аналитика и обработка естественного языка.

Секція 1

Системний аналіз складних систем різної природи

1. Методи системного аналізу складних систем різної природи в умовах невизначеності та ризиків.
2. Математичні методи, моделі та технології дослідження складних систем.
3. Системна методологія технологічного передбачення в задачах планування та прийняття стратегічних рішень.
4. Теорія та методи прийняття оптимальних рішень.
5. Проблемно-орієнтовані методи аналізу та проектування складних систем за умов невизначеності та ризиків.
6. Нелінійні задачі системного аналізу.
7. Системна методологія сталого розвитку.
8. Текстова аналітика та обробка природної мови.

Aghaei Agh Ghamish Yaghoub^{1,2} Donets G.A.^{1,2} Aghaei Agh Ghamish Ovi Nafas^{1,2}
¹Igor Sikorsky Kyiv Polytechnic Institute, Kyiv, Ukraine; ²V.M. Glushkov Institute of Cybernetics of NAS of Ukraine, Kyiv, Ukraine

Mathematical secures with different locks

Mathematic safes with different types of locks are considered. It is shown that the solution is achieved in a system of linear equations under different modules. The appropriate examples are given.

A mathematical safe is a system $S < Z, b < Z >>$, consisting of many locks $Z = \{z_1, z_2, \dots, z_N\}$.

Vector safe states $b = (b_1, b_2, \dots, b_n)$, where $b_i \in \{0, 1, \dots, k_i - 1\}$, the state of the i -th lock and the set $< Z > = \{Z_1, Z_2, \dots, Z_N\}$, Where Z_i a lot of locks Adjacent to z_i As a result of one turn the key clockwise in the lock z_i , All adjacent locks go to state one more by the corresponding module. A safe is considered open if it is in a state $b = (0, 0, \dots, 0)$, It is necessary to find for each lock so many turns with a key to open the safe. Consider in the general case the safe defined above, in which all locks are arranged in the form of a rectangular matrix of the size $m \times n$.

Let $l = n(i - 1) + j$ $i = 1, 2, \dots, m; j = 1, 2, \dots, n$. denote Z_i – a lot of locks, combining locks i -th lines and j -th column, and let all locks have an arbitrary number of states and belong to the same type, that is $k_l = L$. The problem reduces to solving a system of linear equations

$$A\vec{x} + \vec{b} \equiv 0(\text{mod}K) \quad (1)$$

where the matrix A of size $mn \times mn$ consists of m^2 cells:

$$\begin{bmatrix} \mathfrak{Z}_n & E_n & E_{n..} & E_n \\ E_n & Im_n & E_{n..} & E_n \\ E_n & E_n & Im_{n..} & E_n \\ E_n & E_n & E_{n..} & Im_n \end{bmatrix}$$

There are safes in which the states of some locks do not coincide with the states of others; safes with unlike locks.

Consider a simple safe with a matrix of states of locks $b = \begin{bmatrix} 2 & 4 \\ 5 & 1 \end{bmatrix}$, in which the locks of the first row have 5 states, and the locks of the second row have 7 states. According to (1), we obtain a system of equations

$$\begin{aligned} x_1 + x_2 + x_3 &= -2(\text{mod}5), \\ x_2 + x_3 + x_4 &= -4(\text{mod}5), \\ x_1 + x_3 + x_4 &= -5(\text{mod}7), \\ x_2 + x_3 + x_4 &= -1(\text{mod}7), \end{aligned} \quad (2)$$

in which $X = (x_1, x_2, x_3, x_4)$, a $B = (2, 4, 5, 1)$ denote $S_1 = (k_1 + k_2)$, $S_2 = (k_3 + k_4)$.

Adding the first two lines and the last ones, we obtain an abbreviated system

$$\begin{aligned} 2S_1 + S_2 &= -6(\text{mod}5), \\ S_1 + 2S_2 &= -6(\text{mod}7). \end{aligned} \quad (3)$$

This system has a non-unique solution. The simplest of them $S_1 = -1, S_2 = 1$. From the system (3), substituting the values and successively obtain: $x_3 = -1, x_4 = 2, x_1 = 1, x_2 = -2$

There are also safes in which each lock is unique. Consider the following example for a safe with a matrix $b = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$. In this safe, the 1st lock has 2 states, the 2nd lock has 3 states, the 3rd lock

has 4 states and the 4th lock has 5. According to (1) we obtain the following system:

$$\begin{aligned}
 x_1 + x_2 + x_3 &= -1(mod2), \\
 x_1 + x_2 + x_4 &= 0(mod3), \\
 x_1 + x_3 + x_4 &= 0(mod4), \\
 x_2 + x_3 + x_4 &= -1(mod5).
 \end{aligned} \tag{4}$$

This system is equivalent to the following

$$\begin{aligned}
 x_1 + x_2 + x_3 &= 2\alpha - 1, \\
 x_1 + x_2 + x_4 &= 3\beta, \\
 x_1 + x_3 + x_4 &= 4\chi, \\
 x_2 + x_3 + x_4 &= 5\delta - 1,
 \end{aligned} \tag{5}$$

where $\alpha, \beta, \chi, \delta$ – arbitrary whole numbers.

For a greater number of different locks, it is necessary to apply specific methods for solving systems of equations.

Baranovska L. V.

Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

On one differential-difference game of approach with one evader and one pursuer

We consider a game of approach with controlled object whose dynamics is described by the linear differential-difference system with pure time-lag $\tau = \text{const} > 0$ in an Euclidean space $\mathbb{R}^n$

$$\dot{z}(t) = Bz(t - \tau) + \varphi(u, v), \quad z \in \mathbb{R}^n, \quad u \in U, \quad v \in V, \quad (1)$$

where B is a square constant matrix of order n ; U and V are nonempty compact sets; $\varphi: U \times V \rightarrow \mathbb{R}^n$, is jointly continuous in its variables; u and v are the control parameters of the pursuer and the evader chosen from the control domains U and V , respectively; z is the state vector, involving geometric coordinates, velocities, accelerations of the pursuer and the evader.

The initial condition

$$z(t) = z^0(t), \quad -\tau \leq t \leq 0, \quad (2)$$

is absolutely continuous on $[-\tau, 0]$.

Definition 1. (see [1–3]). For each $k = 1, 2, \dots$, the time-delay exponential is defined as follows

$$\exp_\tau\{B, t\} = \begin{cases} \Theta, & -\infty < t < -\tau; \\ I, & -\tau \leq t < 0; \\ I + B \frac{t}{1!} + B^2 \frac{(t-\tau)^2}{2!} + \dots + B^k \frac{(t-(k-1)\tau)^k}{k!}, & (k-1)\tau \leq t \leq k\tau \end{cases}$$

Lemma 1. (see [3]). Let $z(t)$ be a continuous solution to the system (1) with pure time-lag under the initial condition (2). Then,

$$\begin{aligned} z(t) = & \exp_\tau\{B, t\} z^0(-\tau) + \int_{-\tau}^0 \exp_\tau\{B, t - \tau - s\} \dot{z}^0(s) ds \\ & + \int_0^t \exp_\tau\{B, t - \tau - s\} \varphi(u(s), v(s)) ds. \end{aligned}$$

The terminal set has cylindrical form (see [4–7]), i.e. $M^* = M_0 + M$, where M_0 is a linear subspace in $\mathbb{R}^n$ and M is a compact set from the orthogonal complement L of M_0 in $\mathbb{R}^n$.

The goal of the pursuer (u) is in the shortest time to bring a trajectory of the process to a certain closed set M^* ; the goal of the evader (v) is to avoid a trajectory of the process from meeting with the terminal set M^* on a whole semi-infinite interval of time or if is impossible to maximally postpone the moment of meeting.

The game is evolving on the closed time interval $[0, T]$, T is a moment when a trajectory of the process brings to a terminal set M^* , $T > 0$, such that $z(T) \in M^*$ or $\pi z(T) \in M$, where π is the orthogonal project, $\pi: \mathbb{R}^n \rightarrow L$.

Consider the set-valued mapping

$$\bar{W}(t, v) = \pi \exp_\tau\{B, t\} \varphi(U, v), \quad \bar{W}(t) = \bigcap_{v \in V} \bar{W}(t, v).$$

Condition 1 (Pontryagin's condition). The mapping $\bar{W}(t) \neq \emptyset$ for all $t \geq 0$.

For fixed $g(\cdot) \in G$ we put

$$\begin{aligned} \xi(t, z^0(\cdot), g(\cdot)) = \\ = \pi \exp_\tau\{B, t\} z^0(-\tau) + \int_{-\tau}^0 \pi \exp_\tau\{B, t - \tau - s\} \dot{z}^0(s) ds + \int_0^t g(s) ds. \end{aligned}$$

Denote so-called resolving function (see [4], [5])

$$\alpha(t, s, z^0(\cdot), v, g(\cdot)) = \sup\{\alpha \geq 0 : [\bar{W}(t - \tau - s, v) - g(t - \tau - s)] \cap \alpha[M - \xi(t, z^0(\cdot), g(\cdot))] \neq \emptyset\}.$$

Consider the function

$$T = T(z^0(\cdot), g(\cdot)) = \inf \left\{ t \geq 0 : \int_0^t \inf_{v \in V} \alpha(t, s, z^0(\cdot), v, g(\cdot)) ds \geq 1 \right\}, \quad g(\cdot) \in G. \quad (3)$$

Theorem 1. Let the conflict controlled process (1) with the initial condition (2) satisfy Condition 1, and let the set M be convex, for the given initial state $z^0(\cdot)$ and some selection $g^0(\cdot) \in G$ $T = T(z^0(\cdot), g^0(\cdot)) < +\infty$.

Then a trajectory of the process (1) can be brought by the pursuer from $z^0(\cdot)$ to the terminal set M^* at the moment T (3).

Let the differential-difference game of approach

$$\dot{z}(t) = bz(t - \tau) + u(t) - v(t), \quad z \in \mathbb{R}^n, \quad 0 < b < 1, \quad \tau > 0,$$

be given, where the initial states $z(t) = z^0$, $-\tau \leq t \leq 0$, $\|u\| \leq 1$, $\|v\| \leq 1 - b$, $\|z^0\| = 1$. The game is considered complete if $x = y$.

The terminal set $M^* = M = M_0 = \{z \in \mathbb{R}^n : z = 0\}$, $L = \mathbb{R}^n$, π being the identity operator.

Consider the set-valued mapping

$$\bar{W}(t, v) = \pi \exp_{\tau}\{B, t\} \varphi(U, v),$$

and verify the Pontryagin's condition:

$$\bar{W}(t, v) = \exp_{\tau}\{bI, t\}(S - v), \quad \bar{W}(t) = \{0\},$$

where I is a unit matrix of order n ; S is the unit ball centered at zero in the space L .

The condition $\bar{W}(t) = \{0\}$ uniquely determines the selection $g(t) = 0$.

We say that in the game from the initial state $z^0 \notin M^*$ it is possible to avoid meeting the terminal set if there exists a measurable function $v(t) \in V$, $t \geq 0$, such that $z(t) \notin M^*$ for any $t \geq 0$.

It is shown in (see [7]) that if in this example we put in $v(t) = (b - 1)z^0$, $t \in [0, T]$, then $\|z(t)\| \geq 1$, that is $z(t) \notin M^*$ for any $t \geq 0$. Thus, in such a game of two persons, it is possible to avoid meeting with the terminal set with any control of the pursuer, in spite of the fact that the dynamic capabilities of the pursuer are greater. But if the pursuers are several then it is shown that the pursuit game can be completed.

Thus, a scheme of the method of resolving functions for a class of differential-difference games of approach with pure time-lag for the case of one evader and one pursuer has been developed.

References. **1.** Khusainov, D.Y., Benditkis D.D., Diblik, J.: Weak Delay in Systems with an Aftereffect. Functional Differential Equations. 9(3-4), 385–404 (2002). **2.** Khusainov, D.Ya., Diblik, J., Ruzhichkova, M.: Linear dynamical systems with aftereffect. Representation of decisions, stability, control, stabilization. GP Inform-Analyt. Agency, Kiev (2015). **3.** Diblók, J., Khusainov, D.Y.: Representation of solutions of linear discrete systems with constant coefficients and pure delay. Advances in Difference Equations, 1687–1847 (2006). doi: 10.1155/ADE/2006/80825. **4.** Chikrii, A.A.: Conflict-Controlled Processes. Springer Science & Business Media (2013). **5.** Chikrii, A.A.: Analytical method in dynamic pursuit games. In: Proceedings of the Steklov Institute of Mathematics. 271(1), 69–85 (2010). **6.** Baranovskaya, L.V.: A method of resolving functions for one class of pursuit problems. Eastern-European Journal of Enterprise Technologies. 74(4), 4–8 (2015). doi: 10.15587/1729-4061.2015.39355. **7.** Baranovskaya, L.V.: Method of resolving functions for the differential-difference pursuit game for different-inertia objects. Advance in Dynamical Systems and Control. 69, 159–176 (2016). doi: 10.1007/978-3-319-40673-2.

Baranovska L.V.¹, Baranovska G.G.²

¹Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine; ²Igor Sikorsky Kyiv Polytechnic Institute, Kyiv, Ukraine

Quasilinear group pursuit differential games with delay

Let a control system be described by the differential equations with delay in an Euclidean space $\mathbb{R}^{n_i}$

$$\begin{aligned} \dot{z}_i(t) &= A_i z_i(t) + B_i z_i(t - \tau_i) + \varphi_i(u_i, v), \\ z_i &\in \mathbb{R}^{n_i}, \quad u_i \in U_i, \quad v \in V, \quad i = 1, 2, \dots, m, \end{aligned} \quad (1)$$

where A_i and B_i are square constant matrices of order n_i ; U_i and V are nonempty compact sets; the function $\varphi_i: U_i \times V \rightarrow \mathbb{R}^n$, are jointly continuous in totality of variables; $\tau_i = \text{const} > 0$.

A piece of the trajectory

$$z^t(\cdot) = (z_1^t(\cdot), \dots, z_m^t(\cdot))$$

where

$$z_i^t = \{z_i(t + s), -\tau_i \leq s \leq 0\}, \quad i = 1, 2, \dots, m,$$

is the state of the system (1) at the t .

The terminal set M^* has consists of cylindrical sets M_i^* , $i = 1, 2, \dots, m$, having the form

$$M_i^* = M_i^0 + M_i, \quad (2)$$

where M_i^0 is a linear subspace in $\mathbb{R}^{n_i}$ and M_i is a compact subset from the orthogonal complement L_i of M_i^0 in $\mathbb{R}^{n_i}$.

Permissible controls $u_i(s)$, $v(s)$ are the Lebesgue measurable functions that take values from compacts U_i and V , respectively. Denote

$$\begin{aligned} \Omega_{U_i} &= \{u_i(s) : u_i(s) \in U_i, s \in [0, +\infty)\}, \quad i = 1, 2, \dots, m, \\ \Omega_V &= \{v(s) : v(s) \in V, s \in [0, +\infty)\}. \end{aligned}$$

Functions

$$u_i(t) = u_i(z_i^0(\cdot), t, v(t)), \quad i = 1, 2, \dots, m,$$

such that $v(\cdot) \in \Omega_V$ implies $u_i(\cdot) \in \Omega_{U_i}$ are called countercontrols of pursuers corresponding to initial states $z_i^0(\cdot)$.

Consider the problem of group pursuit in the class of countercontrols for the process (1), (2). This problem consists in steering at least one of the trajectories of system (1) to the corresponding set (2) for any counteractions of the evader.

An approach to the solution of such pursuit game based on the method of resolving functions (see [1–3]) is proposed. The integral presentation of game solution based on the time-delay exponential (see [4]) is proposed at first time. The guaranteed capture time is found, and the corresponding control law is constructed.

References. 1. Chikrii, A.A.: Conflict-Controlled Processes. Springer Science & Business Media (2013). 2. Baranovskaya, G.G., Baranovskaya, L.V.: Group Pursuit in Quasilinear Differential-Difference Games. Journal of Automation and Information Sciences. **29**(1), 55–62 (1997). doi: 10.1615/JAutomatInfScien.v29.i1.70. 3. Baranovska, L.V.: On quasilinear differential-difference games of approach. Problemy upravleniya i informatiki. **4**, 5–18 (2017). 4. Khusainov, D.Ya., Diblik, J., Ruzhichkova, M.: Linear dynamical systems with aftereffect. Representation of decisions, stability, control, stabilization. GP Inform-Analyt. Agency, Kiev (2015).

Dychko A.O.¹, Yermeev I.S.²

¹Igor Sikorsky Kyiv Polytechnic Institute, Institute of Energy Saving and Energy Management, Kyiv, Ukraine; ²Taurida National V.I. Vernadsky University, Kyiv, Ukraine

Risk-oriented management of water ecosystems

The functioning of the water ecosystem management system is characterized by the lack of a systematic approach to the analysis of the state of natural reservoirs, as well as the accurate and operative measurement of the qualitative and quantitative indicators of environmental safety of the hydrosphere. Ensuring the reliability of ecosystem assessment data is an urgent problem, because on the basis of these data, a decision should be made on forecasting changes in the stability of water hydrobiocenoses and the necessary measures for protection of surface water bodies and the environment as a whole in the presence of the risk of natural or man-made accidents.

The management of ecosystems, including water ones, takes place under conditions of limited (incomplete) and indistinct information that affects the effectiveness of self-cleaning and self-healing processes of natural landscapes. Indicators of environmental safety assessment are not determined online, as a rule, and in addition, certain averaged data is used. At the same time, regulatory influences aimed at achieving effective management may become inadequate to the current situation and there is a risk of losing control of ecosystem sustainability processes. Modelling the processes of pollution spread and identifying the dynamics of ecological state changes can help to find effective measures to prevent negative impact of pollution on the environment by creating artificial barriers to the migration of these contaminants, or organizing active measures to compensate or neutralize them.

Under conditions of uncertainty, this risk assessment should be simplified to cross the relevant probabilities of events and consequences of the implementation on the matrix, which indicates which zone it belongs to. As other risks also potentially exist and influence each other, increasing the overall risk assessment, it is necessary to determine a corrective coefficient which takes into account the non-formalizing (or exact) determinants of additional impacts. For that, elements of the theory of fuzzy sets are also used, and all the effects are presented in the same scale.

When the risk is identified and evaluated, the following approaches to risk management are used

- avoidance (exclusion);
- reduction (softening);
- restraint;
- transfer (insurance).

While compiling for each source of risk, it is necessary to specify areas in which the consequences of the implementation by definition are:

- permissible, that is, they may take place, practically without affecting the functioning, or because of their incredible implementation;
- unacceptable;
- relatively permissible, which are reduced to the minimum possible in concrete real conditions.

Risk Indicator (Risk Price) VaR indicates maximum losses that will not be exceeded for a given probability and confidence interval in a given time interval.

The optimal risk assessment for VaR_o will be a disjunctive mark that characterizes the maximum risk value, that is, one that meets the following condition:

$$VaR_o = \max(VaR_j),$$

where VaR_j in the classical form corresponds to the product of the probability of the implementation of the event associated with the risk, and the amount of losses in case of its implementation.

Corrective coefficient can be presented in the form

$$\mu = 1 + (r - 1)/2r,$$

and an integral risk assessment (if there is no data on the distribution law) – respectively as

$$VaR_{oi} = \mu VaR_o, \text{ or } VaR_{oi} = \mu(2,33V\sigma),$$

where V – maximum possible loss in the case of risk implementation R , σ – standard deviation, r – the number of identified risks, T – the term during which the risk is determined, in weeks, months, or years, 2.33 – the coefficient that corresponds to the confidence interval 99%.

Kondratova L.P.

Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

Application of Box-Jenkins models for forecasting the guaranteed functioning of complex technical systems in real-world conditions

The strategy of the guaranteed operation of a complex technical system (CES) in real-world operating conditions, implemented in the form of an information platform for technical diagnostics, involves the use of the principles of timeliness and reliability of assessing and forecasting risk situations, guaranteeing timely detection, recognition and assessment of the risk of abnormal regime for the forecast period complex system operation [1, 2].

It is proposed to use the Box-Jenkins models that have been widely used for predicting stationary and non-stationary processes: an autoregressive model, a moving average model, an autoregressive model with a moving average, an autoregressive model with an integrated moving average. The corresponding prediction model is chosen as the result of fulfilling hypothesis testing conditions about the constancy of the mathematical expectation and variance using the Student and Fisher distribution quantiles [3, 4].

Box-Jenkins models are implemented to estimate the predicted values of CES performance indicators of $y_i^p \mid i = \overline{1, m}$ or parameters of $x_j^p \mid j = \overline{1, n}$ with the subsequent restoration of functional dependencies $\tilde{y}_i = \phi_i(\{x_j^p \mid j = \overline{1, n}\}, F_i(\{\rho_{qk} \mid k = \overline{1, n_q}\}))$, $i = \overline{1, m}$ in some time interval of $D^+ = \{t_p \mid t_0^+ < t_p \leq t^+\}$ for forecasting taking into account the level of impact of $F_i(\{\rho_{qk} \mid k = \overline{1, n_q}\})$ to destabilizing risk factors.

This technique is based on using a discrete sample of values of real performance indicators and restored functional dependencies of $\tilde{y}_i = \phi_i(\{x_j \mid j = \overline{1, n}\}, F_i(\{\rho_{qk} \mid k = \overline{1, n_q}\}))$, $i = \overline{1, m}$ in the initial time interval of $D_0 = \{t \mid t_0^- \leq t_0 \leq t_0^+\}$.

There is introduced a procedure for detecting a possible failure of sensor (the step functions of the 1st and 2nd levels [1] and the Bollinger indicator [5]). The use both step functions and Bollinger lines (bands) that characterize the instrument of technical analysis and a technical indicator reflecting the current deviations of the observed value, are oriented on a decrease in the level of dependence on the error of the measured indicators.

Predicted values of $y_i^p \mid i = \overline{1, m}$ and $x_j^p \mid j = \overline{1, n}$ are estimated by the relationships describing the operation of CES using models of AR(r) autoregression, MA(q) moving average, ARMA(r,q) autoregression with moving average, ARIMA(r,d,q) autoregression with an integrated moving average, respectively in form [3]:

$$y_i^p[t_k] = \sum_{\tau=1}^r \mu_\tau \cdot y_i[t_k - \tau] + \epsilon_i[t_k], i = \overline{1, m}; x_j^p[t_k] = \sum_{\tau=1}^r \mu_\tau \cdot x_j[t_k - \tau] + \epsilon_j[t_k], j = \overline{1, n}; \quad (1)$$

$$y_i^p[t_k] = \epsilon_i[t_k] - \sum_{\tau=1}^q \theta_\tau \cdot \epsilon_i[t_k - \tau], i = \overline{1, m}; x_j^p[t_k] = \epsilon_j[t_k] - \sum_{\tau=1}^q \theta_\tau \cdot \epsilon_j[t_k - \tau], j = \overline{1, n}; \quad (2)$$

$$y_i^p[t_k] = \sum_{\tau=1}^r \mu_\tau \cdot y_i[t_k - \tau] + \epsilon_i[t_k] - \sum_{\tau=1}^q \theta_\tau \cdot \epsilon_i[t_k - \tau], i = \overline{1, m}; \quad (3)$$

$$x_j^p[t_k] = \sum_{\tau=1}^r \mu_\tau \cdot x_j[t_k - \tau] + \epsilon_j[t_k] - \sum_{\tau=1}^q \theta_\tau \cdot \epsilon_j[t_k - \tau], j = \overline{1, n};$$

$$y_i^p[t_k] = \sum_{\tau=1}^r \mu_\tau \cdot y_i[t_k - \tau] + \epsilon_i[t_k] - \sum_{\tau=1}^q \theta_\tau \cdot \epsilon_i[t_k - \tau] + \sum_{l=1}^{d-1} \Delta^l y_i[t_k - l] + y_i[t_k - 1], i = \overline{1, m}; \quad (4)$$

$$x_j^p[t_k] = \sum_{\tau=1}^r \mu_\tau \cdot x_j[t_k - \tau] + \epsilon_j[t_k] - \sum_{\tau=1}^q \theta_\tau \cdot \epsilon_j[t_k - \tau] + \sum_{l=1}^{d-1} \Delta^l x_j[t_k - l] + x_j[t_k - 1], j = \overline{1, n};$$

In (1) – (4) $\mu_1, \dots, \mu_r, \theta_1, \dots, \theta_q$ characterize the coefficients of autoregression and moving average, respectively, determined using, for example, the method of least squares; r, q, d characterize given orders of autoregression, moving average and differences; $\epsilon_i[t_k], \epsilon_j[t_k]$ are the random errors at time points $t_k \in D_0 \cup D^+$.

Predicted values of $x_j^p \mid j = \overline{1, n}$ are used to restore the functional dependencies of $\tilde{y}_i = \phi_i(\{x_j^p \mid j = \overline{1, n}\}, F_i(\{\rho_{qk} \mid k = \overline{1, n_q}\}))$, $i = \overline{1, m}$ in time interval of forecasting of $D^+ = \{t_p \mid t_0^+ < t_p \leq t^+\}$.

Study of CES operation process using the Box-Jenkins models was carried out on the example of a closed water supply system, whose state at each time moment of a given operating period is characterized by the indicator values of the water level y_1, y_2 in the inlet and outlet reservoirs, and the temperature y_3 and the water pressure y_4 at the inlet of technological object. The indicators of y_1, y_2, y_3, y_4 are the functions of the system's parameters that represent the total water consumption of x_1 , the number of included pumps of x_2 , the speed of the regulated pump of x_3 , the temperature of the cooling water of x_4 , the pressure at the outlet of the pumping unit of x_5 , the power of the heat losses of x_6 , the set temperature of the process unit of x_7 , the hydraulic valve resistance of x_8 .

Using the Fisher, Student, Cochran-Cox criteria, it was revealed the stationarity of the time series' behavior for some performance indicators and CES parameters.

The forecast is built in real time on the basis of AR(r), MA(q), ARMA(r,q), ARIMA(r,d,q) models, taking into account the restored in time interval of $D_0 = \{t \mid t_0^- \leq t_0 \leq t_0^+\}$ functional dependencies of $\tilde{y}_i = \phi_i(\{x_j \mid j = \overline{1, n}\})$, $i = \overline{1, m}$. The accuracy of prediction for the performance indicators of $y_i^p \mid i = \overline{1, 4}$ and the functional dependencies of $\tilde{y}_i = \phi_i(\{x_j^p \mid j = \overline{1, 8}\}, F_i(\{\rho_{qk} \mid k = \overline{1, 8}\}))$, $i = \overline{1, 4}$ composed 2.31%–15.6% compared to a discrete sample of real indicators without using the procedure for detecting a possible sensor failure. Using this procedure, the accuracy of the prediction composed 0.03%–11.3%). This result guarantees the detection of the reasons for the possible transition of the CES to an abnormal mode of operation and also their elimination with introducing the margin of permissible risk [6].

Thus, the use for predicting Box-Jenkins models guarantees the survivability of the operation of the CES in real operating conditions under the influence of destabilizing factors.

References. **1.** Pankratova N.D. System strategy for guaranteed safety of complex engineering systems // Cybernetics and System Analysis. – 2010. – V. 46, N 2. – P. 243–251. **2.** Pankratova N.D. System coordination of survivability and safety complex engineering objects operation // Computer Science Journal of Moldova, 2014, vol. N3.- P.14–27. **3.** Davydyuk Ye.S. Box-Jenkins Models Box-Jenkins models as a means of analyzing financial series // Perspective innovations in science, education, production and transport'2013, S'World, 17–26 December, 2013. – <http://www.sworld.com.ua/index.php/ru/conference/the-content-of-conferences/archives-of-individual-conferences/dec-2013>. **4.** Nosko V.P. Econometrics: An to Regression Analysis of Time Series: Tutorial. – M., 2002. – 254 p. (www.iet.ru). **5.** Colby R. The Encyclopedia of Technical Market Indicators. – M.: “Alpina Publisher”, 2011. – 840 p. **6.** Pankratova N.D., Kondratova L.P. System evaluation of engineering objects' operating taking into account the margin of permissible risk // Eastern-European Journal of Enterprise Technologies. – 2016. – T.3, №4(81). – P. 13–19. – DOI:10.15587/1729-4061.2016.71126.

Malysheva M.O., Akhmedova D.N., Prymak I.K.

Igor Sikorsky Kyiv Polytechnic Institute, Faculty of Applied Mathematics, Kyiv, Ukraine

Method of modification of base graph-logical models

Fault-tolerant multiprocessor systems capable of reconfiguration have become widespread with their main advantage of high reliability. Such systems continue to function during the failure of certain modules and loses its ability to operate when another set of modules fails, and these sets can intersect. Different mathematical models are used to solve the problem of reflecting the reaction of multiprocessor systems to the failures of their modules.

In some cases, it is relatively easy to describe the logic of fault-tolerant functioning of the system (for example, in [1], k-out-of-n systems and sequential k-out-of-n systems are considered) but it is difficult to do in the general case. Practical problems require non-base models to be used for correct reflection of multiprocessor system reaction to the modules' refusal. This paper concerns graph-logical models (GL-models) [2].

The GL-model is a non-oriented cyclic graph whose edges have certain Boolean functions assigned, whose arguments are indicator variables x_i ($i = 1, 2, \dots, n$), representing the operability of the corresponding component of the system ("0" represents a failure, "1" represents operability). The edge is considered removed from the graph when its corresponding function is equal to zero. The connectivity of the graph represents the operability of the system as a whole. Systems which consists of n modules and maintain operability when no more than m of its components fail are called base systems and are denoted as $K(m, n)$ [3].

Models for non-base systems can be constructed by transforming base models. One of the ways of doing so is the introduction of additional edges which block the loss of connectivity of the graph of the GL-model when the corresponding state vectors of the system appear. In this paper case of increasing the fault-tolerance is considered, defining $W = w_1, w_2, \dots, w_z$ as a set of system's state vectors that cause the system to maintain operability.

Consider the basic GL-model $K(2, 6)$ [2]:

$$f_1 = x_1 \vee x_2 x_3$$

$$f_2 = x_2 \vee x_3 x_4$$

$$f_3 = x_3 \vee x_4 x_5$$

$$f_4 = x_4 \vee x_5 x_6$$

$$f_5 = x_5 \vee x_6 x_1$$

$$f_6 = x_6 \vee x_1 x_2$$

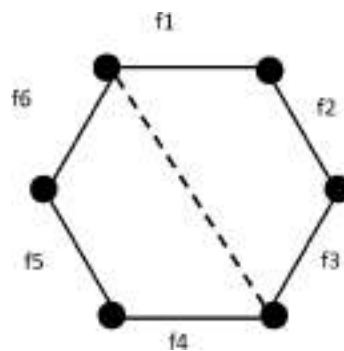


Figure 1. Graph of $K(2, 6)$ GL-model with additional edge

In order to determine the optimal placement of additional edges in [4], it is proposed to divide a graph into a subset of edges S – non-empty non-disjoint subsets of the base model edges, such that the graph connectivity is preserved in the loss of any pair of edges belonging to different subsets, and the graph is rendered disconnected by the loss of any pair of edges belonging to any of the subsets containing more than one edge.

Thus, it is possible to define pairs of basic edges that will be “protected” by an additional edge, that is, loss of these edges will not lead to loss of the graph of connectivity. In the example in figure 1, such pairs will be $f_1f_4, f_1f_5, f_1f_6, f_2f_4, f_2f_5, f_1f_6, f_3f_4, f_3f_5, f_3f_6$. Analyzing these pairs, one can choose the additional edges that are required to block the given set W . Since the basic $K(m, n)$ models do not determine the sequence of the notation of edge functions on the graph, one can, if necessary, change the order of their designations to obtain new combinations of “protected” edges. However, for certain sets of W the formation of so-called pairwise edge cycles can be observed.

The idea of pairwise edge cycles is as follows. Suppose that there are system’s state vectors w_1, w_2, w_3 in W . Suppose that when w_1 appears the edges a_i and a_j are removed from the graph of the GL-model, with the appearance of $w_1 - a_j$ and a_k , and with the appearance of $w_3 - i$ and k . The edges a_i, a_j and a_k form a cycle, preventing the blocking of such sets by a single additional edge.

Since the addition of the edges complicates the establishment of the connectivity of the graph, the problem of the preliminary determination of pairwise edge cycles and their correct processing is one of the major drawbacks of the method of converting GL-models by introducing internal edges. To overcome this defect the pairwise cycles that can appear should be analyzed, including those which share edges or vertices, and the subsequent optimization to the adding of the internal edges when the target set of vectors is blocked should be carried out.

References. 1. A.E. Gera (2004). Combined k-out-of-n:G, and Consecutive kc-out-of-n:G systems. IEEE Transactions on Reliability, R-53, 523–531. 2. Романкевич О.М. (1998). До питання побудови моделі поведінки багатомодульних систем // Наукові вісті НТУУ “КПІ” №1, 38–40. 3. Романкевич А.М., Карачун Л.Ф., Романкевич В.А. (2001). Графо-логические модели для анализа сложных отказоустойчивых вычислительных систем. // Электронное моделирование. Т. 23, №1, 102–111. 4. Романкевич А.М., Иванов В.В., Романкевич В.А. (2004). Анализ отказоустойчивых многомодульных систем со сложным распределением отказов на основе циклических GL-моделей. // Электронное моделирование. Т. 26, №5, 67–81.

and predict the next state of the network $x^j(t_s)$. Fig. 1 presents an image for a learning lateral connections. The dots show the numbers of active neurons for different moments of time.

The presented model differs from the traditional that compute activation based only on a previous state of the network since we track activation from more distant moments of time. As well, it is different from the standard recurrent neural network, since it uses binary neurons, nondifferentiable activation function and unsupervised learning procedure of creating clusters G_m^j . Furthermore, it is more biologically plausible. Similar ideas were presented earlier in [7]; however, we use different activation function and learning rule, that allowed us to gain much higher memory capacity.

Results. We made a theoretical analysis of memory capacity of the proposed model. We showed that for a sparsely active network, detection of a coincidence through time leads to the high capacity of sequence memory and each cell can store only a small amount of synapses per sequence. It is important because it determines the accuracy of the prediction. Mathematically, by tracking coincident activation of previously active neurons, the cell increases the dimension of the network that allows to better separate confusing patterns of activation.

We performed a simulation of the proposed model and confirmed theoretical conclusions. Furthermore, we investigated how the capacity of the model depends on parameters. The longer the history time the neuron can access the higher the capacity. Also, the more neurons from the one moment of time belong to a cluster the larger the capacity. Interestingly, for some combination of parameters, the neural network can recognize and predict hundreds of thousands of sequences.

Discussion and Conclusions. The idea that the single neuron can learn multiple spatiotemporal sequences on behavioral time scale still needs more experimental verification. From our theoretical and simulation analysis we can see that this yet hypothetical idea leads to a much greater computational power of a biological neuron and the network overall.

We captured only minimal biological details, namely multiple coincidence detections through time, but other significant elements are missing. The next natural extension of the model should be an inclusion of forgetting connections and adjusting for limited resources. Expectedly, it will select the most frequent transition in the input stream and reflect its probability distribution.

Overall, we showed that presented model, inspired by recent experimental findings from dendritic computation, provides a high capacity of sequence memory thus gives high accuracy for a sequence prediction.

References. 1. Branco, T. (2014). Computing Temporal Sequence with Dendrites, 11, 245–257. <http://doi.org/10.1007/978-1-4614-8094-5>. 2. Bhalla, U.S. (2017). Dendrites, Deep Learning, and Sequences in the Hippocampus. *Hippocampus*, 2014(6). <http://doi.org/10.1002/hipo.22806>. 3. Branco, T., Clark, B.A., Häusser, M. (2010). Dendritic discrimination of temporal input sequences in cortical neurons. *Science*, 329(5999), 1671–1675. <http://doi.org/10.1126/science.1189664>. 4. He, K., Huertas, M., Hong, S.Z., Tie, X.X., Hell, J.W., Shouval, H., Kirkwood, A. (2015). Distinct Eligibility Traces for LTP and LTD in Cortical Synapses. *Neuron*, 88(3), 528–538. <http://doi.org/10.1016/j.neuron.2015.09.037>. 5. Sun, R., Giles, C.L. (2001). Sequence learning: from recognition and prediction to sequential decision making. *IEEE Intelligent Systems*, 16(4), 67–70. <http://doi.org/10.1109/MIS.2001.1463065>. 6. Dasgupta, S., Stevens, C.F., Navlakha, S. (2017). A neural algorithm for a fundamental computing problem. *Science*, 358(6364), 793–796. <http://doi.org/10.1126/science.aam9868>. 7. Bose, J., Furber, S.B., Shapiro, J.L. (2005). An associative memory for the on-line recognition and prediction of temporal sequences. <http://doi.org/10.1109/IJCNN.2005.1556028>.

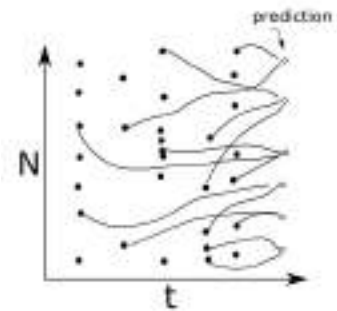


Figure 1. Illustration of an idea of connection through time

Shaptala R.V., Kyselova A.G., Kyselov G.D.

Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

Evaluation of approaches to emotion classification

Introduction. High-quality data, annotated with emotions, is difficult to obtain in large volumes, since emotion annotation is a rather subjective task and because of its complexity requires a lot of time and human resources. Moreover, the long-term goal of artificial intelligence includes making machines understand emotions [1]. That is why the creation of a highly accurate emotion classifier is vital for the advance of AI research.

Approaches to emotion classification. To construct the emotion classifier, several different sources of texts, annotated with emotions were used, including Hashtag Emotion Corpus [2] and Affective Text: Data Annotated for Emotions and Polarity [3]. To make the task better defined, the data sets were manually cleaned and several classes were merged, so the final target emotions are: sadness, anger, neutral, happiness and love. As decision trees, random forests and logistic regressions usually perform well on text classification tasks, they are investigated in the experiment. A recently featured FastText [4] library provides a modern approach of mixing hierarchical neural networks with bag of words to achieve state-of-the-art performances on other text-related tasks. That is why we include it in the emotion classification evaluation as well.

Evaluation results. Here are the details of trained models:

- Logistic regression got TF-IDF (term frequency-inverse document frequency) attributes extracted from texts as input; to accommodate many classes, the logistic regression was trained using OVR(one-vs-rest) method; to combat overfitting, L2 regularization was used.
- The decision tree had the following hyperparameters: maximum depth – 10, criterion – entropy, minimum number of samples required to split an internal node – 3.
- Randomized forests had the following hyperparameters: number of trees – 10, minimum number of samples required to split an internal node – 2, criterion – Gini.
- FastText trained with the following hyperparameters: vector size – 500, learning rate – 0.01, number of epochs – 50, number of negative samples – 100.

All hyperparameters were selected using grid search with 10-fold cross-validation. The final results of the evaluation on the test set are shown in Table 1.

Table 1. F1-micro score of classification models

Model	F1-micro
Decision Tree	0.3755
Random Forest	0.4063
Logistic Regression	0.5028
FastText	0.5118

Conclusion. FastText proved to be the most accurate model for emotion classification. It can be successfully used for emotion annotation in dialogue systems or for sentiment analysis in various social applications.

References. **1.** Bartneck, C., Lyons, M.J., & Saerbeck, M. (2017). The relationship between emotion models and artificial intelligence. arXiv preprint arXiv:1706.09554. **2.** Mohammad, S.M., & Kiritchenko, S. (2015). Using hashtags to capture fine emotion categories from tweets. Computational Intelligence, 31(2), 301–326. **3.** Strapparava, C., & Mihalcea, R. (2007, June). Semeval-2007 task 14: Affective text. In Proceedings of the 4th International Workshop on Semantic Evaluations (pp. 70–74). Association for Computational Linguistics. **4.** Joulin, A., Grave, E., Bojanowski, P., & Mikolov, T. (2016). Bag of tricks for efficient text classification. arXiv preprint arXiv:1607.01759.

Siryk S. V.

Igor Sikorsky Kyiv Polytechnic Institute, Faculty of Applied Mathematics, Kyiv, Ukraine

On the application of the mass lumping correction technique for convection-diffusion problems

We provide a careful Fourier analysis of the Guermond-Pasquetti mass lumping correction technique [1] applied to pure transport and convection-diffusion problems. In particular, it is found that increasing the number of corrections reduces the accuracy for problems with diffusion; however all the corrected schemes turn out to be more accurate than the consistent Galerkin formulation in this case. For the pure transport problems the situation is the opposite. We also investigate the difference between two numerical solutions – the consistent solution and the corrected ones, and show that increasing the number of corrections makes solutions of the corrected schemes closer to the consistent solution.

It is known (see [1–3]) that the application of mass lumping in Galerkin finite element methods may lead to dispersion and dissipative errors, and, accordingly, to significant errors in the numerical solution. The paper of Guermond J.-L. and Pasquetti R. [1] is devoted to the investigation and overcoming of this essential drawback of mass lumping.

The approach of [1] is based on the use of matrix series to approximate the matrix M^{-1} (inverse to the mass matrix of the Galerkin finite-element method): for this, M is represented as $M = \bar{M}(I - A)$ (where I is the identity matrix, $A \equiv \bar{M}^{-1}(\bar{M} - M)$, $\bar{M}$ is a diagonal matrix whose diagonal elements are equal to the sums of the elements of the corresponding rows of M), whence $M^{-1} = (I + A + A^2 + \dots)\bar{M}^{-1}$ (the Neumann series). In [1], authors note that even one correction (i.e., using $(I + A)\bar{M}^{-1}$ instead of M^{-1}) can significantly improve the accuracy of the numerical solution.

Consider the one-dimensional nonstationary convection-diffusion (CD) equation [3–5]

$$\frac{\partial u}{\partial t} + \lambda \frac{\partial u}{\partial x} - \kappa \frac{\partial^2 u}{\partial x^2} = 0, \quad (1)$$

where $\lambda = \lambda(t)$ and $\kappa = \kappa(t)$ are given, $\kappa \geq 0$, $\kappa + |\lambda| > 0$, and $u = u(t, x)$ is an unknown solution.

For the numerical solution we use the classical Galerkin finite-element method with linear elements, i.e., piecewise-linear continuous Lagrange trial functions (hat functions), on a uniform grid with step h [1, 3–5].

The generic form of the consistent (i.e., without mass lumping) Galerkin finite-element formulation leads to the system $M\dot{\vec{a}} = \vec{F}(t, \vec{a})$ (in its generic form). The last equation can be rewritten as $\dot{\vec{a}} = M^{-1}\vec{F}(t, \vec{a})$.

Definition 1. The system $\dot{\vec{a}} = (I + A + A^2 + \dots + A^n)\bar{M}^{-1}\vec{F}(t, \vec{a})$ is called the n -th corrected scheme.

Statement 1. The typical k -th equation (corresponding to the typical mesh node x_k) of the n -th corrected scheme has the form

$$\frac{da_k}{dt} + \lambda a_k^{(1)} - \kappa a_k^{(2)} - \frac{\lambda h^2}{6} a_k^{(3)} + \frac{\kappa h^2}{6} a_k^{(4)} + \dots + (-1)^n \frac{\lambda h^{2n}}{6^n} a_k^{(2n+1)} - (-1)^n \frac{\kappa h^{2n}}{6^n} a_k^{(2n+2)} = 0,$$

where $a_k^{(i)}$ is the central divided finite-difference of the i -th order at the node k .

We carry out a careful Fourier analysis of the schemes corrected by the Guermond-Pasquetti mass lumping correction technique (i.e., studying the behavior of flat monochromatic harmonics in the form $a_k(t) = e^{\omega t} e^{ikh p}$, where $i = \sqrt{-1}$, $p \in \mathbb{R}$ is the (spatial) wave number, $\omega \in \mathbb{C}$ is the quantity that determines the dissipative and dispersion characteristics of systems).

Let us denote the number ω for the n -th corrected scheme (see Statement 1) by ω_n , and the

number ω for the consistent Galerkin FEM by ω_G . It is easy to show that

$$\omega_G = \frac{-\frac{4\kappa}{h^2} \sin^2 \frac{ph}{2} - i\frac{\lambda}{h} \sin(ph)}{1 - \frac{2}{3} \sin^2 \frac{ph}{2}}.$$

In particular, we established the following statements.

Statement 2. The real and imaginary parts of ω_n can be expressed as follows:

$$\begin{aligned} \operatorname{Re}(\omega_n) &= -\frac{4\kappa}{h^2} \sin^2 \left(\frac{ph}{2} \right) - \frac{8\kappa}{3h^2} \sin^4 \left(\frac{ph}{2} \right) - \frac{16\kappa}{9h^2} \sin^6 \left(\frac{ph}{2} \right) - \dots - \frac{2^{n+2}\kappa}{3^n h^2} \sin^{2n+2} \left(\frac{ph}{2} \right), \\ \operatorname{Im}(\omega_n) &= \left(-\frac{\lambda}{h} - \frac{2\lambda}{3h} \sin^2 \left(\frac{ph}{2} \right) - \frac{4\lambda}{9h} \sin^4 \left(\frac{ph}{2} \right) - \dots - \frac{2^n \lambda}{3^n h} \sin^{2n} \left(\frac{ph}{2} \right) \right) \sin(ph). \end{aligned}$$

Substituting the ansatz $u(t, x) = e^{\bar{\omega}t} e^{ixp}$ into the equation (1) we obtain $\bar{\omega} = -\kappa p^2 - i\lambda p$.

Note that

$$\lim_{h \rightarrow 0} \omega_G = \lim_{h \rightarrow 0} \omega_n = \bar{\omega},$$

i.e., the solutions (found in the class of harmonics) of all considered numerical problems tend to the solution of the (1) if h tends to zero.

Statement 3. Let $\kappa > 0$, $n \geq 1$, $\lambda \in \mathbb{R}$. Then the inequality $|\omega_n - \bar{\omega}| < |\omega_{n+1} - \bar{\omega}|$ is true for all h sufficiently small.

Statement 4. Let $\kappa = 0$, $n \geq 1$, $\lambda \neq 0$. Then the inequality $|\omega_{n+1} - \bar{\omega}| < |\omega_n - \bar{\omega}|$ is true for all h sufficiently small.

Statement 5. Let $\kappa > 0$, $n \geq 1$, $\lambda \in \mathbb{R}$. Then the inequality $|\omega_n - \bar{\omega}| < |\omega_G - \bar{\omega}|$ is true for all h sufficiently small.

Statement 6. Let $\kappa = 0$, $n \geq 1$, $\lambda \neq 0$. Then the inequality $|\omega_G - \bar{\omega}| < |\omega_n - \bar{\omega}|$ is true for all h sufficiently small.

Statement 7. Let $n \geq 1$, and parameters λ and κ are arbitrary. Then the inequality $|\omega_G - \omega_{n+1}| < |\omega_G - \omega_n|$ is true for all h sufficiently small.

Thus, it is found that increasing the number of corrections reduces the accuracy for problems with diffusion (Statement 3); however all the corrected schemes turn out to be more accurate than the consistent Galerkin formulation in this case (Statement 5). For the pure transport problems the situation is the opposite (Statements 4 and 6, respectively). We also show that increasing the number of corrections makes solutions of the corrected schemes closer to the consistent solution (Statement 7).

We carried out numerical experiments for confirming the corresponding theoretical results obtained.

Following to [3], it is easy to go from semi-discrete formulations to families of difference schemes with weights and show their mean-square stability in the grid norm L_2 .

References. 1. Guermond J.-L., Pasquetti R. A correction technique for the dispersive effects of mass lumping for transport problems // *Comput. Methods Appl. Mech. Engrg.* – 2013. – Vol. 253. – P. 186–198. 2. Wendland E., Schulz H.E. Numerical experiments on mass lumping for the advection-diffusion equation // *Revista Minerva* – 2005. – Vol. 2, № 2. – P. 227–233. 3. Siryk S.V. Analysis of Lumped Approximations in the Finite-Element Method for Convection-Diffusion Problems // *Cybernetics and Systems Analysis* – 2013. – Vol. 49 (5). – P. 774–784. 4. Salnikov N.N., Siryk S.V. Construction of weight function of the Petrov-Galerkin method for convection-diffusion-reaction equations in the three-dimensional case // *Cybernetics and Systems Analysis* – 2014. – Vol. 50 (5). – P. 805–814. 5. Sirik S.V. Estimation of the accuracy of finite-element Petrov-Galerkin method in integrating the one-dimensional stationary convection-diffusion-reaction equation // *Ukrainian Mathematical Journal* – 2015. – Vol. 67 (7). – P. 1062–1090.

Soloviev V.N.¹, Romanenko Y.V.²

¹*Faculty of Physics and Mathematics, State Pedagogical University, Kryvii Rih, Ukraine;* ²*Igor Sikorsky Kyiv Polytechnic Institute, Faculty of Aerospace Systems, Kyiv, Ukraine*

Quantum econophysics of bitcoin crises

The attempts to create an adequate model of socio-economic critical events, which, as it has been historically proven, are almost permanent, were, are and will always be made. Actually, it is a supertask, impossible to solve. However, the potentially useful solutions, local in time or other socio-economic logistic coordinates, are possible. In fact, they have to be the object of interest for a real and effective economic science.

Econophysics is a young interdisciplinary scientific field, which developed and acquired its name at the end of the last century [1]. Quantum econophysics, a direction distinguished by the use of mathematical apparatus of quantum mechanics as well as its fundamental conceptual ideas and relativistic aspects, developed within its boundaries just a couple of years later, in the first decade of the 21-st century.

According to classical physics, immediate values of physical quantities, which describe the system status, not only exist, but can also be exactly measured. Although non-relativistic quantum mechanics doesn't reject the existence of immediate values of classic physical quantities, it postulates that not all of them can be measured simultaneously (Heisenberg uncertainty ratio). Relativistic quantum mechanics denies the existence of immediate values for all kinds of physical quantities, and, therefore, the notion of system status seizes to be algebristic.

In paper [2] we have suggested a new paradigm of complex systems modelling based on the ideas of quantum as well as relativistic mechanics. It has been revealed that the use of quantum-mechanical analogies (such as the uncertainty principle, notion of the operator, and quantum measurement interpretation) can be applied to describing socio-economic processes. Methodological and philosophical analysis of fundamental physical notions and constants, such as time, space and spatial coordinates, mass, Planck's constant, light velocity from the point of view of modern theoretical physics provides an opportunity to search of adequate and useful analogues in socio-economic phenomena and processes.

Suppose there is a set of M time series, each of N samples, that correspond to the single distance T , with an equal minimal time step Δt_{min} :

$$X_i(t_n), t_n = \Delta t_{min} \cdot n; n = 0, 1, 2, \dots, N - 1; i = 1, 2, \dots, M. \quad (1)$$

To bring all series to the unified and non-dimensional representation, accurate to the additive constant, we normalize them, having taken a natural logarithm of each term of the series:

$$x_i(t_n) = \ln X_i(t_n), t_n = \Delta t_{min} \cdot n; n = 0, 1, 2, \dots, N - 1; i = 1, 2, \dots, M. \quad (2)$$

Let us consider that every new series $x_i(t_n)$ is a one-dimensional trajectory of a certain fictitious or abstract particle numbered i , while its coordinate is registered after every time span Δt_{min} , and evaluate mean square deviations of its coordinate and speed in some time window ΔT .

The «immediate» speed of i particle at the moment t_n is defined by the ratio:

$$v_i(t_n) = \frac{x_i(t_{n+1}) - x_i(t_n)}{\Delta t_{min}} = \frac{1}{\Delta t_{min}} \ln \frac{X_i(t_{n+1})}{X_i(t_n)} \quad (3)$$

with variance D_{v_i} and mean square deviation Δv_i .

The chaotic nature of real time series allows to $x_i(t_n)$ as the trajectory of a certain abstract quantum particle (observed at Δt_{min} time spans). We can write an uncertainty ratio for this

trajectory [3]:

$$\Delta x_i \cdot \Delta \nu_i \sim \frac{h}{m_i}, \text{ or } : \frac{1}{\Delta t_{min}} \left(\left\langle \ln^2 \frac{X_i(t_{n+1})}{X_i(t_n)} \right\rangle_{n, \Delta N} - \left(\left\langle \ln \frac{X_i(t_{n+1})}{X_i(t_n)} \right\rangle_{n, \Delta N} \right)^2 \right) \sim \frac{h}{m_i}, \quad (4)$$

where m_i – economic “mass” of an i series, h – value which comes as an economic Planck’s constant.

We will test the model for the cryptocurrency market on the example of the bitcoin. Bitcoin is an important electronic and decentralized cryptographic currency system. It is based on a peer-to-peer network architecture and secured by cryptographic protocols and there is no need for a central authority or central bank to control the money supply within the system. Bitcoin attracts considerable attention of researchers of different levels, using modern methods and models of analysis of the peculiarities of the dynamics of the popular digital currency. Thus, the identification of possible trends of the cryptocurrency movement, construction and modeling of indicators of stability and possible crisis states is extremely relevant.

During the entire period (16.07.2010 – 01.04.2018) of verifiably fixed daily values of the bitcoin price (BTC) (<https://finance.yahoo.com/cryptocurrencies>) in relative units, five crisis phenomena were recorded and marked with arrows on fig. 1. Calculations were carried out within the framework of the algorithm of a moving window. For this purpose, the part of the time series (window), for which there were measures M was selected, then the window was displaced along the time series in a one-day increment and the procedure repeated until all the studied series had exhausted. Further, comparing the dynamics of the actual time series and the corresponding measures M , we can judge the characteristic changes in the dynamics of the behavior of M with changes in the cryptocurrency.

Obviously, there is a dynamic characteristic values M depending on the internal dynamics of the market. In times of crisis known (marked by arrows in the fig. 1) “mass” is significantly reduced in the pre-crisis period. In fig. 2 mass calculations were proved only for the last fifth crisis and with the window not in 500 days, but only 50. Obviously that the value of M remains a good indicator and in this case.

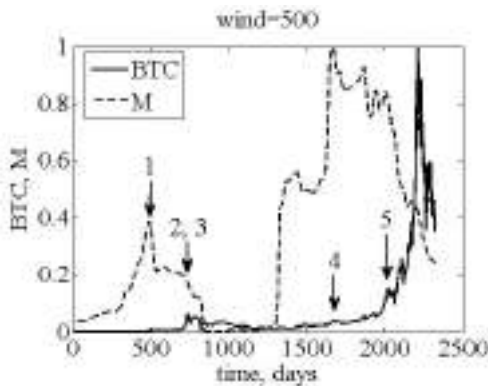


Figure 1

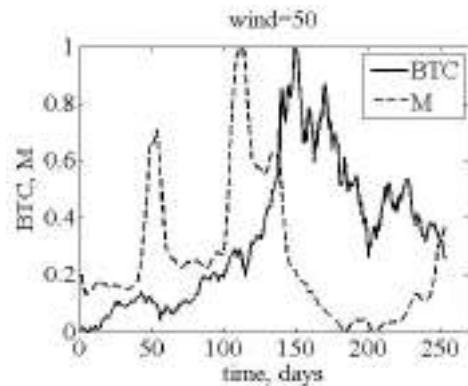


Figure 2

This fact can be used as an indicator-precursor crisis.

References. 1. R.N. Mantegna, H.E. Stanley, “An Introduction to Econophysics”. – Cambridge: Cambridge University Press, 2000. – 144 p. 2. V. Sapsin, V. Soloviev (2009, Jul 7) Relativistic quantum econophysics – new paradigms in complex systems modelling [Online]. Available: <https://arxiv.org/0907.1142>. 3. V. Soloviev, V. Sapsin (2011, Nov 10) Heisenberg uncertainty principle and economic analogues of basic physical quantities [Online]. Available: <https://arxiv.org/1111.5289>.

Ved O.V.

National Technical University "Kharkiv Polytechnic Institute", Kharkiv, Ukraine

Basic Concepts of Catalytic Exhaust Gas Aftertreatment

Atmospheric pollutants, such as nitrogen oxides, carbon monoxide and volatile organic compounds, mainly emanating from industrial plants and vehicles, are harmful to human health [1]. Catalysis is one of the most effective and economical technologies for dealing with serious air pollution problems.

Reducing pollutant emission of vehicles and meeting stricter exhaust gas standards are major challenges when developing catalytic converters. Catalytic neutralization of the organic gas mixture, monoxide and nitrogen oxide is the most forward-looking method. Thus make it possible process multicomponent gases with small initial concentrations of substances, achieve high purification rates, conduct the process continuously, and also avoid in most cases the formation of secondary pollutants.

The efficiency of the heterogeneous catalytic process as a whole is significantly influenced by the parameters of the catalyst converter, such as porosity, specific surface area, heat resistance, mechanical strength. The development of effective converters of harmful gas components is based on the need to create hydrodynamic and kinetic models for neutralization of exhaust gases with both carrier characteristics and the multicomponent nature of engine exhaust gases. Increasing the requirements for the efficiency and efficiency of gas emission purification devices from the economic and environmental point of view, the complexity and cost of technological processes require more universal and accurate methods for calculating the parameters and properties of catalytic blocks using mathematical modeling methods [2, 3].

A new three-level block model of a heterogeneous catalytic reaction has been done. The model has a fairly universal character and can be applied to describe a wide class of experiments. It is studied using the example of the CO oxidation reaction on the catalyst surface. The main distinguishing feature of the model is the simultaneous consideration of all the basic physicochemical processes occurring during the reaction in the catalyst layer: reactions on the metal surface; internal diffusion in the pores of the catalyst grains; limited reagent feed rate; the thermal effect of the reaction; heat and mass transfer through the layer.

The model is create on a hierarchical principle and includes several levels of description, corresponding to a certain space-time scale. It defines all the main relationships between the processes occurring in the reaction at the micro, meso, and macro levels; it is a kind of minimal model for an adequate description of the reaction in the catalyst layer. Each level can be detailed, taking into account additional factors and naturally integrate them into the full model, or to study certain approximations of the complete model.

The original numerical algorithm is also presented, which allows calculating the full model. Using the proposed algorithm, the complete distributed model was investigated on the basis of different variants of the point model, in a wide range of external parameters and characteristics of the catalyst layer.

References. **1.** Dai H. Environmental catalysts: a solution for the removal of atmospheric pollutants. *Sci. Bull.*, Vol. 60, No.19, pp. 1708–1710, 2015. **2.** Deutschmann O. Modeling of the Interactions Between Catalytic Surfaces and Gas-Phase. *Catalysis Letters*, Vol.145, No.1, pp. 272–289, 2015. **3.** Gossler H., Deutschmann O. Numerical optimization and reaction flow analysis of syngas production via partial oxidation of natural gas in internal combustion engines. *International Journal of Hydrogen Energy*, Vol.40, No.34, pp. 11046–11058, 2015.

Бахрушин В.Є., Подковаліхіна О.О., Логвіненко В.О.

Запорізький національний технічний університет, Запоріжжя, Україна

Задача розподілу інвестицій в умовах статистичної невизначеності

Моделі динамічного програмування застосовують при розв'язуванні таких задач, як розробка правил управління запасами, що встановлюють момент поповнення запасів і розмір замовлення; розробка принципів календарного планування виробництва і вирівнювання зайнятості в умовах мінливого попиту на продукцію; розподіл дефіцитних капітальних вкладень між можливими новими напрямками їх використання; складання календарних планів поточного і капітального ремонту складного обладнання та його заміни і т.п. [1]

Метод динамічного програмування дає можливість отримати оптимальний розв'язок n -вимірної задачі шляхом її декомпозиції на n етапів, кожен з яких є підзадачею з однією змінною. Таким чином, розв'язування n -вимірної задачі зводиться до розв'язування одновимірних оптимізаційних підзадач. Обчислювальні аспекти розв'язування оптимізаційних підзадач на кожному етапі проектується і реалізуються окремо, що не виключає застосування єдиного алгоритму для всіх етапів.

В динамічному програмуванні обчислювання виконують рекурентно, оптимальний розв'язок однієї підзадачі використовують як вихідні дані для наступної. Розв'язавши останню підзадачу, отримуємо оптимальний розв'язок вихідної задачі. Спосіб виконання рекурентних обчислень залежить від того, як здійснюють декомпозицію вихідної задачі. Підзадачі зазвичай пов'язані між собою деякими загальними обмеженнями.

Розглянемо задачу розподілу інвестицій (задача розподілу коштів між підприємствами) [1]. Планується модернізація n підприємств. Для цієї мети виділено кошти обсягом y_1 ум. од. Для кожного підприємства j розроблено декілька альтернативних проектів. Кожен проект характеризується витратами c_j та майбутніми прибутками R_j . На кожному підприємстві можна реалізувати лише один проект. Необхідно обрати такі проекти для кожного підприємства, щоб фірма отримала максимальний річний прибуток.

Математична постановка задачі розподілу інвестицій:

$$f_{n+1}(y_{n+1}) = 0, \\ f_j(y_j) = \max_{x_j | c_j(x_j) \leq y_j} (R_j(x_j) + f_{j+1}(y_j - c_j(x_j))), j = \overline{n, 1},$$

де x_j – номер проекту, обраного на підприємстві j ; y_j – кількість грошей для підприємств $j, \dots, n$; $c_j(x_j)$ – витрати на j -ом підприємстві, коли обрано проект з номером x_j ; $R_j(x_j)$ – річний прибуток, що буде отримано від реалізації проекту x_j ; $f_j(y_j)$ – максимальний річний прибуток, що буде отримано від реалізації проектів $x_j, x_{j+1}, \dots, x_n$ при заданому обсязі інвестицій y_1 .

Відомий розв'язок [1, 2] цієї задачі для наведеної вище детермінованої постановки. Але на практиці витрати і майбутні прибутки точно не відомі. Тому актуальним є її розв'язання в умовах невизначеності. Найпростішим видом невизначеності є статистична невизначеність, яку можна задати функцією розподілу, або типом і параметрами розподілу. У цьому випадку критерієм оптимальності є математичне сподівання річного прибутку фірми.

Задача розподілу інвестицій була розв'язана за допомогою пакета MATLAB двома методами: методом динамічного програмування (табличний метод) і методом перебору. Невизначеність вхідних параметрів задавали як випадкові послідовності, що підпорядковуються нормальним розподілам із заданими параметрами. Програма дозволяє отримати оптимальні розв'язки для задачі розподілу інвестицій в умовах статистичної невизначеності.

Література. 1. Исследование операций в экономике: Учеб. пособие для вузов / Н.Ш. Кремер, Б.А. Путько, И.М. Тришин, М.Н. Фридман; Под. ред. проф. Н.Ш. Кремера. – М.: ЮНИТИ, 2005. – 407 с. 2. Таха, Хемди А. Введение в исследование операций, 7-е издание.: Пер. с англ. – М.: Издательский дом “Вильямс”, 2005. – 912 с.

Беседовский М.О.

КПИ им. Игоря Сикорского, ИПСА, Киев, Украина

Модернизация маркетинговой стратегии Котлера на этапе рыночного спада в Украине

Вступление. Маркетинговая стратегия – это развернутый план ведения и организации рабочего процесса. Из-за затянувшегося ухудшения экономической ситуации в Украине — это актуальный вопрос для предприятия любого масштаба от малого бизнеса до крупного. Рассмотренной является одна из самых распространенных на данный момент стратегий работы, стратегия «Удержания», разработанная Филипом Котлером. Цели стратегии заключаются в максимально длительном периоде удержания существующей доли рынка, основным методом развития которой является распространение товаров через посредников [1]. В период кризиса стратегия «Удержания» используется как стратегия «Сокращения номенклатуры», которая заключается в сохранении на рынке товара с наибольшей нормой прибыльности, концентрации на масштабе [2].

Цель исследования. Модернизировать стратегию Филипа Котлера в условиях рыночного спада и усиления конкуренции. Также проверить результаты модернизации на практическом примере компании по установке видеонаблюдения.

Ситуация конкуренции. При выборе стратегий руководители предприятий и маркетингологи зачастую не ориентируются на конкуренцию рынка, считая, что всегда найдется свободная доля рынка, но на деле оказывается, что за нее нужно бороться, так как все потребители уже пользуются определенным товаром.

Методика исследования. Проведен PEST анализ рынка товаров дешевого сегмента, среднего и luxury продукции. Данный анализ позволяет определить макросреду предприятия и рыночные тенденции отрасли. Также был проведен SWOT анализ компаний этих ценовых сегментов, для определения внутренних и внешних факторов компании, оценке рисков и конкурентоспособности товара в отрасли.

Для проведения исследования была выбрана компания по установке видеонаблюдения из среднего ценового сегмента с низкой нормой прибыли, основными клиентами которой является сфера B2B.

Для определения результатов было проведено изменение прибыльности компании за трехмесячный период после изменения маркетинговой стратегии.

Модернизация маркетинговой стратегии. Стандартная стратегия «Сокращения номенклатуры» не предусматривает борьбу с конкуренцией, и поэтому был модернизирован наиболее подходящей стратегией Майкла Портера, одного из основных исследователей и практиков в сфере конкуренции, стратегией «Дифференциации», идея которой заключается в предоставлении уникальной характеристики в товаре [3].

Литература. 1. Филип Катлер, Основы маркетинга. Краткий курс, Пер. с англ. — М.: Издательский дом “Вильямс”, 2007. 2. Power branding (23.03.2018) <http://powerbranding.ru/marketing-strategy/spad/> 3. Портер Е. Майкл, Конкурентная стратегия: Методика анализа отраслей и конкурентов, М.: Альпина Бизнес Букс, 2005.

Болдак А.О.¹, Єфремов К.В.², Пишинограєв І.О.^{2,3}, Горох О.С.¹

¹КПІ ім. Ігоря Сікорського, ФІОТ, Київ, Україна; ²Світовий центр даних з геоінформатики та сталого розвитку КПІ ім. Ігоря Сікорського, Київ, Україна; ³КПІ ім. Ігоря Сікорського, ФММ, Київ, Україна

Використання результатів SWOT-аналізу в сценарному моделюванні

У роботі [1] наведені основні ідеї організації сценарного моделювання для соціо-економічних систем, що передбачає побудову сценаріїв та їх тестування на основі використання баєсівської мережі довіри (БМД) в якості моделі об'єкта управління. Згідно [1, 2] відповідальність за побудову такої БМД покладена на групу експертів.

Побудова сценарію, тобто послідовності подій, що відображають керуючі впливи та реакції об'єкта управління, яка відповідає критеріям правдоподібності та внутрішньої узгодженості [1], є фундаментальним аспектом в контексті побудови оптимізованої стратегії управління, тобто послідовності кроків (дій, що розгортаються в часі), які з певною ймовірністю можуть привести до бажаного або запланованого кінцевого стану об'єкта управління [2], та повинна виконуватися таким чином, щоб усунути суб'єктивне бачення експертами моделі об'єкта управління. Виконання цього завдання викликає труднощі, пов'язані з необхідністю вибору великої кількості факторів та змінних, а також областей їх визначення, та побудови на цій множині причинно-наслідкової моделі поведінки об'єкта управління.

З іншого боку, експертами, що займаються стратегічним плануванням, відомий інструментарій SWOT-аналізу [3]. Цей інструментарій, головною перевагою якого є його простота, з успіхом був використаний, наприклад корпораціями Shell [4], Google [5], для планування стратегічного розвитку соціо-економічних систем різного рівня на різних часових горизонтах [6].

В процесі виконання SWOT-аналізу здійснюється формування факторів чотирьох категорій: сильних сторін (S), слабких сторін (W), можливостей (O) та загроз (T). Внутрішні фактори S та W визначають властивості об'єкта управління, тобто його стан можна задати в просторі, а зовнішні фактори O та T впливи, що змінюють його стан. У свою чергу можливості відображають дії, які залежать від особи, що приймає рішення (ОПР), а загрози – дії, як є незалежними від ОПР, тобто можуть інтерпретуватися як події-керуючі впливи E_x в термінології, запропонованій в роботах [1, 2].

Розвитком метода SWOT-аналізу, запропонованим та апробованим в роботі [6], є методика побудови матриць впливів на основі результатів експертних опитувань. Якщо побудувати згадані матриці впливів та задати певний поріг значень ступеня впливу можна отримати граф причинно-наслідкових відношень на множині вершин, що можна використати в якості основи БМД. Після параметричного налаштування БМД на основі експертних даних та(або) результатів експериментальних досліджень, а також визначення функцій розподілу ймовірностей настання подій в часі, ця модель може бути використана для сценарного моделювання, згідно з методикою, описаною в роботах [1, 2].

Таким чином, запропонована методика застосування результатів SWOT-аналізу для проведення сценарного моделювання, дозволяє спростити процес побудови формальної моделі об'єкта управління.

Література. 1. Л.О. Болдак, О. Горох Моделювання сценаріїв за допомогою баєсівських мереж довіри: Матеріали 19-ї Міжнародної науково-технічної конференції SAIT 2017, Київ, 2017 р. / ННК "ІПСА" НТУУ "КПІ ім. Ігоря Сікорського", 2017. – с. 42. 2. Л.О. Болдак, О. Горох Система підтримки прийняття рішень на основі сценарного планування: Матеріали 6-ї Міжнародної наукової конференції ICS-2017, Львів, 2017 р. / Видавництво ЛП, 2017. – с. 194. 3. Sarsby, A. (2016). SWOT Analysis. Spectaris Ltd. 4. SWOT Analysis of Shell, marketing91.com/swot-analysis-shell. 5. Alphabet (Google) SWOT analysis 2018, strategicmanagementinsight.com/swot-analyses/google-swot-analysis.html. 6. Форсайт та побудова стратегії соціально-економічного розвитку України на середньостроковому (до 2020 року) і довгостроковому (до 2030 року) часових горизонтах / наук. кер. акад. М.З. Згуровський. – Київ: Вид-во "Політехніка", 2016. – 184 с.

Васильев В.И., Вишталъ Д.М., Любашенко Н.Д.

КПИ им. Игоря Сикорского, ФПМ, Киев, Украина

Алгоритм синтеза топологии сети по критерию доступности сетевой услуги

В качестве модели сети рассматривается стохастический граф $G(Y, X)$, где Y – множество узлов сети, X – множество каналов, элементы которого взвешены весовыми параметрами.

Эволюция каналов сети моделируется соответствующими независимыми альтернирующими процессами восстановления (Θ_{kr}, η_{kr}) , $r, k = 1, |Y|$, $r \neq k$, где Θ_{kr} – случайное время доступности канала, η_{kr} – случайное время недоступности канала [1].

Предполагается, что существуют стационарные вероятности $P\{x_{kr} = 1\} = P_{kr}$.

Потребность в связи между узлами сети задается матрицей тяготения $H = (h_{ij})$ на множестве узлов графа [2]. Элемент $h_{ij} = 1$, если сеть должна обеспечить связь узлов y_i и y_j . В противном случае, $h_{ij} = 0$. Предполагается, что потребность в связи между узлами y_i и y_j удовлетворяется, если между ними существует по крайней мере один маршрут определенного качества. Вероятность выполнения этого условия служит характеристикой доступности сетевой услуги для соответствующей пары узлов. Множество вероятностей для всех пар узлов характеризует доступность сетевой услуги, в которой связь отдельной пары узлов обеспечивается независимо от остальных пар. Это множество удобно представлять в виде матрицы $P_{ij}(G)$.

Для сети одновременной связи требуется, чтобы связь была обеспечена всем тяготеющим узлам одновременно, доступность характеризуется вероятностью $P(G)$ связности графа G .

В предположении абсолютной надежности узлов, задача синтеза конфигурации сети формулируется следующим образом. Заданы:

1. Начальная конфигурация графа $G_0(Y, X_0)$.
2. Дополняющий резервный граф $\bar{G}(Y, \bar{X})$, каждое из ребер $x_{ij} \in \bar{X}$ характеризуется приведенными затратами c_{ij} и вероятностью безотказной работы P_{ij} .
3. Множество корреспондирующих пар узлов W и условия удовлетворения потребностей в связи для каждой пары из W .
4. Совокупность p_{rs} для $x_{rs} \in W$ сети независимой связи и требования к вероятности доступности одновременной связи $P(G)$ для сети.

Требуется найти такую конфигурацию G дотройкой G_0 ребрами из дополняющего графа $\bar{G}$, при которой удовлетворяется требование к вероятности доступности $P(G) \geq P$ для сети одновременной связи при минимальных приведенных затратах

$$\sum_{x_{ij} \in \bar{G} \wedge G} C_{ij} \rightarrow \min.$$

Для решения этой задачи используется метод слабейшего сечения, состоящий в реализации монотонной пошаговой процедуры, на каждом шаге которой имеющаяся конфигурация сети связи наращивается одним и только одним каналом. В предлагаемом алгоритме решения поставленной задачи важную роль играет представление условия работоспособности текущего графа G в виде конъюнктивной нормальной формы (КНФ) [3]. Для получения КНФ-представления работоспособности текущего графа надо выполнить следующие шаги:

1) Найти все маршруты $M_{y_i y_j}$, соединяющие множество корреспондирующих пар узлов из W , при этом одна из вершин, например, y_i фиксирована. Каждому маршруту соответствует конъюнктивная неповторная цепочка из числа ребер, составляющих простой маршрут в графе G .

2) Осуществить логическое умножение всех подмножеств маршрутов, соединяющих корреспондирующие вершины, и минимизировать представление, то есть

$$\bigcup_{k=2}^{IWI} M_{y_i y_j} = \bigcap_{i \in I} D_i,$$

где D_i ($i \in I$) – множество всех минимальных остовных деревьев, покрывающих соответствующие узлы.

3) Осуществить логическое умножение и минимизировать представление в классе дизъюнктивной нормальной формы (ДНФ), то есть

$$\bigcap_{i \in I} D_i^d = \bigcup_{j \in J} C_j,$$

где D_i^d – двойственное представление i -го остовного дерева; C_j – конъюнктивная цепочка ребер, образующих минимальное сечение.

4) Записать представление условия полносвязности сети в виде КНФ, то есть в виде:

$$\bigcup_{j \in J} C_j^d,$$

где C_j^d – двойственное представление j -го минимального сечения.

5) Для всех C_j^d вычислить $P\{C_j^d\} = 1$ и найти слабое сечение, лимитирующее надежность сети, то есть, найти

$$\min_j \{P(C_j^d) = 1\}$$

и упорядочить по величине.

6) Вычислить вероятность связности текущего графа G

$$P(G) = P\left\{\bigcap_{j \in J} C_j^d = 1\right\}.$$

Если $P(G) \geq P_{given}$, то конец алгоритма, иначе перейти к шагу 7.

7) Расшить слабое место одним соответствующим ребром из дополняющего графа $\bar{G}$, если таковое в графе $\bar{G}$ имеется. В противном случае перейти к ближайшему, следующему по величине вероятности связности сечению.

8) Для полученной на шаге 7 новой конфигурации графа G_1 , найти $P(G_1)$ и сравнить с требуемым нормативом P_{given} . Если $P(G_1) \geq P_{given}$ и других кандидатов на расшивку сечения нет, то конец алгоритма, в противном случае переходим к шагу 9.

9) Оценить прирост вероятности доступности сетевой услуги на единицу затрат при вводе в конфигурацию минимального сечения соответствующего ребра $(y_i y_j)$ из графа $\bar{G}$,

$$\Delta_1 = \frac{P(G_1) - P(G)}{C_{ij}},$$

где C_{ij} – стоимость введения ребра $(y_i y_j)$.

Если в дополняющем графе $\bar{G}$ имеется несколько кандидатов для расшивки минимального сечения, то выбирается ребро, дающее наибольший прирост вероятности доступности сетевой услуги на единицу затрат Δ_1 . Если при этом $P(G_1) \geq P_{given}$, то конец алгоритма. В противном случае конфигурация графа становится текущей и процедура повторяется.

Процедура синтеза, представленная описанным алгоритмом, в общем случае не гарантирует получение структуры минимальной избыточности. Однако очевидная эффективность целенаправленного перебора позволяет говорить о приближенном методе синтеза структуры сети. Наличие эффективных методов оценки вероятности связности текущего графа дает возможность практического использования метода на стадии проектирования сети.

Литература. 1. Ф. Байхельт, П. Франкен «Надежность и техническое обслуживание. Математический подход» М.: «Радио и связь», 1988 г. – 392 с. 2. Зайченко Ю.П., Васильев В.И., Вишталъ Д.М., Хотячук Р.Ф. «Анализ живучести СПД на стадии проектирования». Матеріали Міжнародної науково-практичної конференції «Інтелектуальні системи прийняття рішень та інформаційні технології» (19–21 травня 2004 р., м. Чернівці). – 2004. – С. 48–49. 3. Г. Биргхоф, Т. Барти «Современная прикладная алгебра» М.: «Мир» 1976 г. – 400 с.

Власов М.Д.¹, Болдак А.О.^{1,2}

¹КПІ ім. Ігоря Сікорського, ФІОТ, Київ, Україна; ²Світовий центр даних з геоінформатики та сталого розвитку КПІ ім. Ігоря Сікорського, Київ, Україна

Оптимізація конфігурацій розподілених інформаційних систем

Вступ. Сучасна організація бізнес процесу розробки програмного забезпечення, такого як Rational Unified Process (скорочено – RUP) [1] передбачає наявність фаз що відповідають за ключові етапи створення програмного продукту. Згідно RUP, існує 4 фази. На початковій фазі відбувається оцінка системи, визначення вартості розробки та необхідного бюджету. На фазі уточнення відбувається пом'якшення ключових ризиків, робиться аналіз предметної області і проектується базова архітектура. На фазі конструювання відбувається розробка системи та її компонентів, можлива зміна архітектури для забезпечення кращих характеристик розроблюваної системи. Фаза впровадження є орієнтованою на покращення системи для користувачів і навчання останніх нею користуватись. Недоліком такого підходу є занадто пізнє визначення архітектурного рішення яке буде лежати в основі реалізації. Зменшення часу на підготовку проектного рішення (артефакту) дозволяє отримати наступні бенефіти:

- Зниження технологічних ризиків (чим раніше – тим більше часу на нівелювання ризиків).
- Зниження вартості реалізації.
- Повторне використання підсистем.

Отже, дослідження, які дозволяють змістити на більш ранню стадію створення артефакту – є актуальними. Мета даної роботи — показати можливість отримання архітектурного рішення на фазі уточнення, виходячи з вимог, сформованих протягом початкової фази і фази уточнення.

Для досягнення мети необхідно вирішити такі завдання:

- Визначити що собою являє такий артефакт як “архітектурне рішення” (що з себе представляє конфігурація).
- Визначити перелік і семантику вимог достатніх для створення архітектурних рішень.
- Сформулювати завдання оптимізації, конфігурації розподілених інформаційних систем (PIC) на безлічі доступних альтернатив і обмежень.

Архітектура програмного забезпечення в циклі розробки. Архітектура програмного забезпечення має багато визначень, таких як: “Найвищий рівень розбиття системи на частини; рішення, які важко змінити; наявність декількох архітектур у системі; те що може значущо змінити архітектуру протягом життя системи; і, врешті-решт, архітектура зводиться до важливих речей.” [2]

“Архітектура програмного забезпечення для програмної або обчислювальної системи – це структура або структури системи, що складаються з елементів програмного забезпечення, видимих зовнішніх властивостей цих елементів та відносин між ними. Архітектура має відношення до видимої сторони інтерфейсів; приховані деталі елементів – деталі, що пов'язані виключно з внутрішньою реалізацією – не є архітектурними.” [3]

Проте вони усі зводяться до наступного: архітектура програмного забезпечення — це процес визначення структурованого рішення, що відповідає всім технічним та операційним вимогам, при оптимізації спільних атрибутів якості, таких як продуктивність, безпека та керованість. Вона включає низку рішень, що базуються на різноманітних факторах, і кожне з цих рішень може мати значний вплив на якість, продуктивність, ремонтпридатність та загальний успіх програми.

Архітектурним проектуванням називають перший етап процесу проектування, на якому визначаються підсистеми, а також структура управління і взаємодії систем.

Підсистема — це система, операції якої не залежать від сервісів, що надаються іншими підсистемами. Підсистеми складаються з модулів — системних компонентів, що надають один або кілька сервісів для інших модулів. Кожна підсистема може представляти чорний ящик з відомою функціональністю і інтерфейсом. Для опису і налаштувань модулів використовуються конфігурації — форми інструкцій для засобів розгортання архітектури.

Вимоги, достатні для визначення архітектури програмного забезпечення. Моделі архітектури можуть залежати від функціональних, експлуатаційних та інших вимог що разом звуться нефункціональні вимоги [4]. Найважливіші з них:

- Ефективність. За критичні операції відповідає якомога менше підсистем, тобто використовується крупномодульна архітектура.
- Захищеність. Багаторівнева архітектура системи, найбільш критичні елементи захищені на нижньому рівні.
- Надійність. Включаються явно зайві компоненти, які можна змінювати не перериваючи роботу системи.
- Зручність супроводу. Архітектура складається з дрібних компонентів, які можна легко адаптувати під вимоги предметної області.
- Безпека. За всі операції, що впливають на безпеку, має відповідати якнайменше підсистем.
- Вартість. Сумарна ціна розробки не має перевищувати виділений на неї бюджет.
- Вартість супроводу. Ціна заміни елементу чи його поліпшення не має бути критичною.

Проблема оптимізації. Завдання оптимізації конфігурації ставиться на безлічі альтернативних конфігурацій, отриманих на безлічі доступних підсистем що покривають функціональні вимоги. f – критерій оптимізації, що обчислюється на основі оцінок нефункціональних вимог.

$$f \subseteq \text{Ресурси} \times \text{Вимоги}$$

При цьому нефункціональні вимоги є аргументами критеріїв оптимізації.

Завдання оптимізації базуються на переліку встановлюваних вимог до системи і їх пріоритетності.

У кожного Ресурса є конфігураційні метадані, що описують усі можливі характеристики і вимоги до використання даного Ресурсу. За допомогою них можна вираховувати усі можливі композиції, відкинути рішення що не задовольняють вимоги і надати кінцевому користувачу найбільш підходящі конфігурації систем, тим самим виконуючи задачу оптимізації і зменшуючи час на підготовку проектного рішення. Дану концепцію можна представити у вигляді таблиці.

Висновки. Сформульовано підхід до визначення архітектури на ранніх фазах життєвого циклу ПЗ виходячи з функціональних і не функціональних вимог.

Поставлені завдання оптимізації конфігурації які можуть вирішуватися відомими математичними методами, такими як “метод гілок і меж” і “динамічного програмування”, що є частинами дискретної оптимізації. Задачами зі схожою проблематикою у дискретній оптимізації є задача пакування рюкзака [5] і задача про призначення [6].

Одним з необхідних умов успішної реалізації відповідних програмних засобів є формування реєстру метаданих доступних для використання підсистем, на основі яких повинні визначатися альтернативні конфігурації.

Література. 1. William Cottrell “Standards, compliance, and Rational Unified Process”, IBM, 2004, <https://www.ibm.com/developerworks/rational/library/4763.html>. 2. Martin Fowler “Patterns of Enterprise Application Architecture”, 2003, <http://www.pearsonhighered.com/educator/academic/product/0,3110,0321127420,00.html>. 3. Len Bass, Paul Clements, Rick Kazman, “Software Architecture in Practice, 3rd Edition”, 2013, <https://www.pearson.com/us/higher-education/program/Bass-Software-Architecture-in-Practice-3rd-Edition/PGM317124.html>. 4. van Leeuwen, J., Muscholl, A., Peleg, D., Pokorný, J., Rumpe, B. “SOFSEM 2010: Theory and Practice of Computer Science” 36th Conference on Current Trends in Theory and Practice of Computer Science, Špindleruv Mlýn, Czech Republic, January 23–29, 2010. Proceedings, <http://www.springer.com/us/book/9783642112652#>. 5. Kellerer H., Pferschy U., Pisinger D. Knapsack Problems — Springer-Verlag Berlin Heidelberg, 2004. — 548 p. — ISBN 978-3-642-07311-3 — doi:10.1007/978-3-540-24777-7. 6. Brualdi, Richard A. (2006). Combinatorial matrix classes. Encyclopedia of Mathematics and Its Applications. 108. Cambridge: Cambridge University Press. ISBN 0-521-86565-4. Zbl 1106.05001.

Волкова В.Н.¹, Логинова А.В.¹, Черный Ю.Ю.², Леонова А.Е.³

¹Санкт-Петербургский политехнический университет Петра Великого, Санкт-Петербург, Россия;

²Институт научной информации по общественным наукам Российской академии наук, г.

Санкт-Петербург, Российская Федерация; ³ОАО «НИЦЭВТ», г. Санкт-Петербург, Российская Федерация

Методика сравнительного анализа и выбора технологических инноваций четвертой промышленной революции

Предлагается методика сравнительного анализ и выбора технологических инноваций четвертой промышленных революций для конкретного предприятия (организации).

Ключевые слова: косвенные количественные оценки; методы системного анализа; подход; промышленная революция; технологическая инновация; теория систем.

Введение. В настоящее время становится все более очевидным, что активное развитие технологий четвертой промышленной революции в ближайшее время изменят не только промышленность, транспорт, экономику в целом, но и кардинально преобразят условия жизни человека. При этом технологические инновации могут оказывать как положительное, так и отрицательное воздействие на все сферы социально-экономических систем, и необходимо проводить анализ и принимать решения о целесообразности выбора новых технологий для конкретного предприятия (организации) с учетом полезности и последствий их внедрения.

Концепция четвертой промышленной революции. Прогнозируемая *четвёртая промышленная революция* означает появление полностью цифровой промышленности, основанной на взаимном проникновении индустриальных и информационных технологий.

Для развития технологий четвертой промышленной революции немецкий экономист, основатель и бессменный президент с 1971 г. Всемирного экономического форума (международной организации в области частно-государственного сотрудничества в Давосе) Клаус Шваб ориентируется на идеи Индустрии 4.0, в соответствии с которой для крупной промышленности планируется широкое внедрение в заводские процессы киберфизических систем (CPS).

К.Шваб предлагает рассматривать три блока – физический, цифровой и биологический [1]. *Физический блок* (беспилотные транспортные средства, 3D-печать, передовая робототехника, новые материалы); *цифровой блок* (интернет вещей и его приложения, удалённый мониторинг, цепочка блоков транзакций – блокчейн, экономика по требованию и др.); *биологический блок* (управление генетикой человека, животных и растений, а также создание клеток взрослых организмов, включая людей, 3D-производство живых тканей – биопечать).

Представленный в книге К. Шваба [1] перечень технологий, полученный в результате опроса 800 руководителей высшего звена (отчёт “Глубинное изменение – технологические переломные моменты и социальное воздействие”, Прогноз до 2025) включает 23 технологии.

Стратегия развития инновационного общества в Российской Федерации на 2017–2030 гг.

В Стратегия развития информационного общества в Российской Федерации на 2017–2030 годы в качестве основных предлагаются следующие технологии [2, 3]:

- а) конвергенция сетей связи и создание сетей связи нового поколения;
- б) обработка больших объемов данных;
- в) искусственный интеллект;
- г) доверенные технологии электронной идентификации и аутентификации, в том числе в кредитно-финансовой сфере;
- д) облачные технологии и туманные вычисления;
- е) интернет вещей и индустриальный интернет;
- ж) робототехника и биотехнологии;
- з) радиотехника и электронная компонентная база;
- и) информационная безопасность.

Таким образом, технологические инновации и их трактовки уже сейчас образуют достаточно обширное неупорядоченное пространство, которое к тому же непрерывно расширяется, и перед

руководителями предприятий (организаций) возникает задача сравнительного анализа и выбора новых технологий с учетом их особенностей и возможностей.

Начинать анализ этого пространства предлагается с изучения закономерностей их взаимовлияния и взаимодействия между собой и с окружающей средой, на которую они оказывают влияние, закономерностей устойчивого развития социально-экономических систем, в условиях активного внедрения технологических инноваций.

Оценка значимости технологических инноваций для конкретного предприятия (организации). Для проведения исследования предлагается применить одну из основных идей системного анализа – упрощение сложного путем постепенного сужения области допустимых решений при организации процесса принятия решений.

Для реализации этой идеи предлагается методика, последовательность этапов которой приведена на рис. 1:



Рис. 1. Методика сужения области допустимых решений при выборе технологических инноваций

Для выполнения этапа 4 применяются методы организации сложных экспертиз и реализующие их автоматизированные диалоговые процедуры [4–8] (рис. 2):

- методы оценки, предложенные в методике ПАТТЕРН (относительной важности, взаимной полезности, “состояние–срок”);



Рис. 2. Методы организации сложных экспертиз

- метод иерархий Т. Саати;
- методы, основанные на информационных оценках А.А. Денисова:

1) информационная оценка потенциал (значимость) H_i компоненты:

$$H_i = -q_i \log(1 - p'_i),$$

где p'_i – вероятность достижения целей организации при использовании инновационной технологии; q_i – вероятность использования конкретной технологии при реализации, достижении соответствующей подцели.

Оценки p' и q_i производится единичными экспертами, для которых определяются сферы компетентности.

Для получения значения относительной значимости α_j необходимо выполнить процедуру нормирования, т.е.

$$H_{\Sigma_\alpha} = \sum H_\alpha, \alpha_j = \frac{H_{\alpha_j}}{H_{\Sigma_\alpha}};$$

- 2) оценка внедрения инновационных технологий во времени и принятие решений о целесообразности продолжения их дальнейшем инвестировании и т.п.;
- 3) оценки сравнительного анализа технологий с учетом их взаимовлияния.

Заключение. При принятии решений результаты экспертных и косвенных количественных оценок инноваций, получаемые на основе предлагаемой методики, удобно представлять виде гистограмм. Можно также исследовать оптимизационную модель с целевой функцией, в которой используется информационная оценка значимости $\sum_{i=1}^n \sum_{j=1}^m H_{ij} x_{ij} \rightarrow \max$ и ограничения по ресурсам.

Литература. 1. Шваб К. Четвертая промышленная революция. Перевод с англ. М.: Изд-во «Э», 2017. – 208 с. 2. Губанов С.С. Державный прорыв. Неоиндустриализация России и вертикальная интеграция. М.: Книжный Мир, 2012. – 224 с. 3. Стратегия развития информационного общества в Российской Федерации на 2017–2030 годы. 9.05.2017. 4. Волкова В.Н., Денисов А.А. Теория систем и системный анализ: учебник / Москва, 2017. – Сер. 58 Бакалавр. Академический курс (2-е изд., пер. и доп). – 462 с. 5. Денисов А.А. Современные проблемы системного анализа. – СПб.: Изд-во Политехн. университета, 2005. – 296 с. 6. Классификация методов и моделей в системном анализе / В.Н. Волкова, В.Н. Козлов, В.Е., Магер, Л.В. Черненькая. // Международная конференция по мягким вычислениям и измерениям. – 2017. – Т.1. – С.223–226. 7. Моделирование систем и процессов. учебник для академического бакалавриата / Под ред. В.Н. Волковой и В.Н. Козлова. М.: Издательство Юрайт, 2015. – 449 с. 8. Моделирование систем и процессов. Практикум: учеб. пособие для академического бакалавриата / Под ред. В.Н. Волковой. М.: Издательство Юрайт, 2016. – 295 с.

Волошин О.Ф.¹, Маляр М.М.²

¹Київський національний університет ім. Тараса Шевченка, Київ, Україна; ²Ужгородський національний університет, Ужгород, Україна

Нечітке математичне моделювання фінансових ризиків інноваційних проектів

Нечітке математичне моделювання є одним із найбільш активних і перспективних напрямів прикладних досліджень в сфері управління і прийняття рішень у слабо структурованих системах. Діапазон застосування нечітких методів з кожним роком розширюється, охоплюючи різні нові області. Нечітку математичне моделювання це коли елементами дослідження є не числа, а деякі нечіткі множини або їх поєднання. В основі такого підходу лежить не традиційна логіка, а логіка з нечіткою істинністю, нечіткими зв'язками і нечіткими правилами виводу. Основними характеристиками такого підходу є використання лінгвістичних змінних замість числових змінних, відношення між змінними описуються за допомогою нечітких висловлювань, складні відношення описуються нечіткими алгоритмами.

При проектуванні і управлінні складних соціо-економічних систем виникає проблема, коли людина не здатна дати точні і в той же час практичні значення суджень про їх поведінку.

Пропонуються дослідження актуальної задачі розроблення математичної моделі інформаційної технології оцінювання ризику проектів відносно рівня безпеки їх фінансування, з використанням нечіткої математики, для різних інвестиційних суб'єктів. Розробка такої технології дасть можливість адекватно підійти до розгляду проектів, підвищити ступінь обґрунтованості прийняття рішень щодо інвестування і основне підвищити економічну та управлінську безпеку.

Сучасні умови розвитку економіки України характеризуються значною кількістю чинників, яким притаманні невизначеність та дестабілізуючий характер дії. Інвестиційна діяльність – важлива складова розвитку економіки. Недостатність інвестицій – болюче питання розвитку економіки будь-якої країни світу. Поряд з тим, інвестиційних проектів існує дуже багато, таких як оновлення матеріально-технічної бази, нарощення об'ємів виробництва, освоєння нових видів діяльності і т.д. Для реалізації таких проектів потрібні ресурси, яких на сьогоднішній день мало. Тому, інвестори дуже обережно підходять до прийняття рішень, щодо інвестування у той чи інший проект, враховуючи різні види ризиків. У такому разі, задача вибору інвестиційних проектів постає дуже актуальною.

У вітчизняній економічній науці відсутні загальновизнані теоретичні положення про ризик інноваційних проектів, недостатньо розроблені методи його оцінки, відсутні рекомендації про шляхи і способи зменшення і запобігання ризику. Останнім часом з'явилися наукові праці, в яких при дослідженні питань планування, економічної діяльності комерційних організацій, співвідношення попиту і пропозиції, розробки проектів розглядаються питання ризику. Що розуміють під ризиком проекту – це ступінь небезпеки для успішного здійснення проекту. Основною характеристикою є невизначеність, яка пов'язана з можливістю виникнення у ході реалізації проекту несприятливих ситуацій і наслідків. Невизначеність передбачає наявність чинників, при яких результати дій не є детермінованими, а ступінь можливого впливу цих чинників на результати невідомий; це неповнота або неточність інформації про умови реалізації проекту.

Ризик в інноваційному підприємстві – вірогідність втрати ресурсів або недоотримання доходів порівняно з варіантом, який розрахований на раціональне використання ресурсів у даному виді діяльності. На сьогоднішній день універсальної класифікації ризиків поки не створено, а тим більше відносно ризику інноваційних проектів.

Приймати ефективні рішення з управління ризиками – значить реалізувати процедури аналізу ризику в ході підготовки стратегічних, тактичних або оперативних рішень і оцінити ризикованість рішення.

Годна А.В., Жданова О.Г., Сперкач М.О.

КПІ ім. Ігоря Сікорського, ФІОТ, Київ, Україна

Про один підхід складання розкладу виконання завдань паралельними пристроями з метою мінімізації сумарного відхилення моментів завершення від директивних строків

Постановка задачі. Задано дві множини завдань $I_1 = \{1, 2, \dots, i_1, \dots, n_1\}$ та $I_2 = \{1, 2, \dots, i_2, \dots, n_2\}$ та кількість пристроїв m ($m > 1$). Множина завдань I_1 має спільний директивний строк d_1 , а множина I_2 – директивний строк d_2 . Пристрої працюють паралельно, є взаємозамінними і відрізняються один від одного продуктивністю виконання завдань. Для пристрою j , $j = \overline{1, m}$ задано коефіцієнт продуктивності h_j такий, що тривалість виконання завдання i на пристрої j дорівнює $h_j p_i$. “Еталонним” будемо називати пристрій з коефіцієнтом продуктивності $h = 1$. У цьому випадку величина p_i є тривалістю виконання завдання i на еталонному пристрої. Передбачається, що всі завдання множин I_1 та I_2 надходять одночасно в нульовий момент часу, процес обслуговування кожного завдання проходить без переривань до завершення обслуговування. Пристрої можуть починати свою роботу в ненульовий момент часу. Необхідно побудувати розклад, що мінімізує сумарне відхилення моментів закінчення виконання завдань від директивних строків.

Поставлена задача відрізняється від задач, розглянутих у роботах [1, 2] тим, що має два директивних строки, тобто є частковим випадком цих задач.

Запропоновано алгоритм розв’язання задачі, що складається з наступних етапів.

ЕТАП I. *Попереднє закріплення пристроїв між підмножинами завдань*

I.1 Розв’язання допоміжної оптимізаційної задачі – отримання підмножин M_1 (пристрої, на які будуть призначатися завдання множини I_1) та M_2 (пристрої, на які будуть призначатися завдання множини I_2).

ЕТАП II. *Побудова початкового допустимого розкладу*

II.1 Побудова розкладу S_1 виконання завдань множини I_1 пристроями множини M_1 .

II.2 Побудова розкладу S_2 виконання завдань множини I_2 пристроями множини M_2 .

ЕТАП III. *Покращення розкладу, отриманого на ЕТАПІ II*

Для реалізації ЕТАПУ I сформульована оптимізаційна задача, яка полягає у тому, щоб розділити множину пристроїв M на дві підмножини M_1 та M_2 , де M_1 – множина пристроїв, на які будуть призначатися завдання множини I_1 , M_2 – множина пристроїв, на які будуть призначатися завдання множини I_2 . З визначення коефіцієнту продуктивності h_j слідує, що $\frac{1}{h_j}$, $j = \overline{1, m}$ – швидкість роботи пристрою j , тоді: $\sum_{j \in M_1} \frac{1}{h_j}$ – сумарна швидкість пристроїв

множини M_1 ; $\sum_{j \in M_2} \frac{1}{h_j}$ – сумарна швидкість пристроїв множини M_2 . Введемо позначення:

$P_1 = \sum_{i \in I_1} p_i$, $P_2 = \sum_{i \in I_2} p_i$ – сумарні обсяги робіт із виконання завдань з множин I_1 та I_2

відповідно. Представляється логічним розбивати множину M на дві підмножини за наступним правилом: сумарна швидкість відповідної підмножини пристроїв повинна бути прямо пропорційна сумарному обсягу робіт або (та) обернено пропорційна до відповідного директивного строку. Твердження «Сумарна швидкість відповідної підмножини пристроїв є прямо пропорційною сумарному обсягу робіт» означає, що:

$$\frac{P_1}{\sum_{j \in M_1} \frac{1}{h_j}} = \frac{P_2}{\sum_{j \in M_2} \frac{1}{h_j}}. \quad (1)$$

Твердження «Сумарна швидкість відповідної підмножини пристроїв є обернено пропорційною директивному строку відповідної множини завдань» означає, що:

$$\frac{d_2}{\sum_{j \in M_1} \frac{1}{h_j}} = \frac{d_1}{\sum_{j \in M_2} \frac{1}{h_j}}. \quad (2)$$

Твердження «Сумарна швидкість відповідної підмножини пристроїв є прямо пропорційною сумарному обсягу робіт та обернено пропорційною до директивного строку відповідної множини завдань» означає, що:

$$\frac{P_1 d_2}{\sum_{j \in M_1} \frac{1}{h_j}} = \frac{P_2 d_1}{\sum_{j \in M_2} \frac{1}{h_j}}. \quad (3)$$

З урахуванням того, що P_1 та P_2 є відомими цілими числами, а $\frac{1}{h_j}$ $j = \overline{1, m}$ є заданими дійсними числами, тому рівності (1), (2) та (3) можуть бути недосяжними, отже постає задача таким чином розбити множину M на дві підмножини, щоб різниця між правими та лівими частинами цих рівнянь була б мінімальною. Таким чином, можуть бути сформульовані наступні критерії розподілу:

$$\left| \frac{\sum_{j \in M_1} \frac{1}{h_j}}{\sum_{j \in M_2} \frac{1}{h_j}} - \frac{\Phi_1}{\Phi_2} \right| \rightarrow \min, \quad (4)$$

де $\frac{\Phi_1}{\Phi_2}$ відповідно (1), (2) та (3) дорівнює $\frac{P_1}{P_2}$, $\frac{d_2}{d_1}$ та $\frac{P_1 d_2}{P_2 d_1}$.

Математична модель допоміжної оптимізаційної задачі. Пристрою j поставимо у відповідність змінну:

$$x_j = \begin{cases} 1, & \text{якщо } j \in M_1, \\ 0, & \text{якщо } j \in M_2. \end{cases} \quad j = \overline{1, m} \quad (5)$$

З урахуванням змістовного значення змінних сумарна швидкість пристроїв множини M_1 становить $\sum_{j \in M_1} \frac{1}{h_j} = \sum_{j=1}^m \frac{1}{h_j} x_j$, множини M_2 становить $\sum_{j \in M_2} \frac{1}{h_j} = \sum_{j=1}^m \frac{1}{h_j} (1 - x_j)$. Тоді критерій оптимізації буде мати вигляд:

$$\left| \frac{\sum_{j \in M_1} \frac{1}{h_j} x_j}{\sum_{j \in M_2} \frac{1}{h_j} (1 - x_j)} - \frac{\Phi_1}{\Phi_2} \right| \rightarrow \min. \quad (6)$$

Обмежень, окрім $x_j \in \{0; 1\}$, $j = \overline{1, m}$, немає.

Експериментальні дослідження показують, що кращий розподіл досягається при використанні критерію оптимізації, що відповідає (3).

На етапі II допустимі розклади будуються за алгоритмом, що є модифікацією алгоритму побудови розкладу для поліноміально розв'язуваної задачі, який наведений в [3]. Модифікація алгоритму полягає у введенні додаткової умови призначення завдання на виконання до множини випереджаючих: $h_j p_i \leq d - h_j \sum_{i \in W_j} p_i$, де W_j – множина завдань, які на пристрої j на поточний момент виконуються з випередженням.

На етапі III здійснюється покращення розкладу, отриманого на етапі II, шляхом виконання перестановок завдань між двома множинами пристроїв M_1 та M_2 .

Література. 1. Mensendiek, A., Gupta, J., Herrmann, J. (2015). Scheduling identical parallel machines with fixed delivery dates to minimize total tardiness. *European Journal of Operational Research*, 243(2), 514–522. 2. Kayvanfar, V., Komaki, G., Aalaei, A., Zandieh, M. (2014). Minimizing total tardiness and earliness on unrelated parallel machines with controllable processing times. *Computers & Operations Research*, 41, 31–43. 3. Годна А.В. Задача мінімізації сумарного відхилення від спільного директивного строку при виконанні завдань паралельними пристроями / А.В. Годна, О.Г. Жданова, А.О. Маленко, М.О. Сперкач // Науковий огляд. – 2017. – №9(14). – С. 14–32.

Горелова Г.В.

*Институт управления в экономических, экологических и социальных системах,
Инженерно-технологическая Академия Южного федерального университета, Таганрог, Россия*

О разработке инструментария когнитивного моделирования социально-экономических систем

Применение когнитивного подхода к исследованию социально-экономических, геополитических, социотехнических и других сложных систем, развитие соответствующей теории когнитивного моделирования [1–4] потребовало создания специального инструментария, включающего программные средств (ПС) поддержки принятия решений. Такие ПС необходимы как в процессе исследования, так и при функционировании интеллектуальной системы поддержки принятия решений в соответствующей предметной области. В докладе представлен ряд результатов разработки программных систем когнитивного моделирования сложных систем и иллюстрационные примеры.

Инструментарий когнитивного моделирования сложных систем – это комплекс методов сбора, обработки, моделирования и анализа информации об объекте, полученной как на реальном объекте, так и в результате имитационного когнитивного моделирования его структуры и поведения. Конечной целью применения этого инструментария к изучению сложных систем разной природы является решение задач: описания и объяснения системы в ее взаимодействии с внешней средой, предвидения возможных путей развития системы (научное предвидение будущего), управления системой или адаптации к ее возможностям, разработки стратегий развития сложной системы.

Комплекс теоретических методов исследования позволяет решать взаимосвязанные задачи когнитивного моделирования: идентификации объекта (применение экспертных, статистических и др. методов), анализа путей и циклов когнитивной модели (методы теории графов), анализа наблюдаемости, управляемости, устойчивости, чувствительности, адаптируемости, катастроф (методы теории управления, теории катастроф); композиции – декомпозиции; анализа различных аспектов сложности, анализа связности (методы теории графов, топологический анализ q -связности); прогнозирования (статистические методы); научного предвидения (сценарный анализ, анализ развития ситуаций – импульсное моделирование); решение задач оптимизации (решение обратной задачи, методы математического программирования); принятия решений в условиях различного рода неопределенности (методы теории принятия решений), сопутствующей существованию и изучению сложной системы. При этом принятие решений происходит как по отношению к самому изучаемому объекту, так и по отношению к процессу исследования. Таким образом происходит когнитивный процесс накопления, преобразования, порождения новых знаний.

Когнитивное моделирование сложных систем (КМСС, акцент в названии сделан на понятии “сложные системы” в целях обозначить отличие от других методологий когнитивного моделирования в психологии, лингвистике и др. когнитивных науках) поддерживается программными системами ПС КМ и CMSS [1–3].

Рис. 1 иллюстрирует две из основных возможностей КМСС: изображение укрупненной когнитивной модели сложной системы “Межрегиональные социально-экономические взаимодействия” и результаты вычислительного эксперимента по одному из возможных сценариев развития этой системы. На изображении когнитивной модели выделен один из отрицательных циклов, общее нечетное (9) количество которых свидетельствует о структурной устойчивости системы. Когнитивная модель имеет вид знакового взвешенного по вершинам и функциональным дугам ориентированного графа, что отличает ее от “простой” когнитивной карты (знаковый ориентированный граф). Графики импульсного процесса в нижней части рисунка соответствуют модельному сценарию развития, вычисленному в предположении о развитии промышленности регионов А и В в условиях отрицательного воздействия внешней среды.

Предвидится положительная тенденция развития ситуаций уже с 3-го шага моделирования

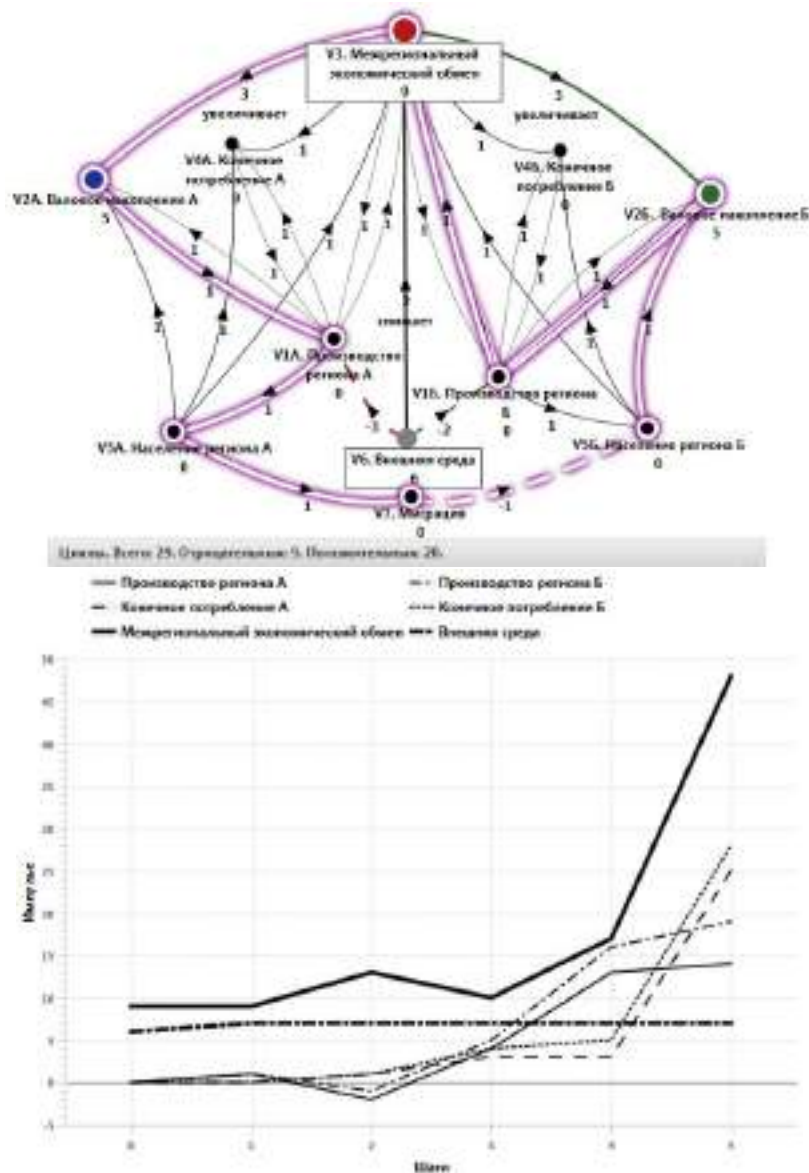


Рис. 1. Когнитивная модель “Межрегиональные социально-экономические взаимодействия” и возможный сценарий развития событий в региональных системах

Литература. 1. Горелова Г.В., Захарова Е.Н, Радченко С.А. Исследование слабоструктурированных проблем социально-экономических систем: когнитивный подход. – Ростов н/Д: Изд-во РГУ, 2006. – 332 с. 2. Инновационное развитие социально-экономических систем на основе методологий предвидения и когнитивного моделирования / Под ред. Г.В. Гореловой, Н.Д. Панкратовой. – Киев: Наукова думка, 2015. – 464 с. 3. Gorelova G.V. Cognitive modeling as the instrument in course of knowledge of large system // International Journal “Information Theories and Applications”, Bulgaria, Vol. 18, N2. – 2011. – P. 172–182. 4. Ginis L.A., Gorelova G.V., Kolodenkova A.E. Cognitive and simulation modeling of development of regional economy system // International Journal of Economics and Financial Issues. – 2016, Vol.6, No 5S, P. 97–103.

Городецкий В.Г., Осадчук Н.П.

КПИ им. Игоря Сикорского, Киев, Украина

Численно-аналитический метод решения обратной задачи для систем одного класса

В данном исследовании рассматривалась задача реконструкции модели в виде системы обыкновенных дифференциальных уравнений (ОДУ) с полиномиальными правыми частями по временному ряду одной наблюдаемой переменной. Так как по одной наблюдаемой переменной невозможно однозначно определить все коэффициенты системы ОДУ [1], выполнялся поиск всех систем, которые способны воспроизвести заданный временной ряд, но при этом имеют разный общий вид уравнений и наименьшее количество ненулевых коэффициентов.

Исследование выполнялось на примере системы Лоренца [2], которая имеет вид

$$\begin{cases} \dot{x}_1 = a_1 x_1 + a_2 x_2, \\ \dot{x}_2 = b_1 x_1 + b_2 x_2 + b_6 x_1 x_3, \\ \dot{x}_3 = c_3 x_3 + c_5 x_1 x_2, \end{cases} \quad (1)$$

где $a_1 = -10$, $a_2 = 10$, $b_1 = 28$, $b_2 = b_6 = -1$, $c_3 = -8/3$, $c_5 = 1$ и $x_1(t)$ — наблюдаемая переменная. В исследовании использовался предложенный в [3] подход, который предполагает вывод аналитических соотношений между коэффициентами неизвестной оригинальной системы (ОС) и так называемой стандартной системы (СС), содержащей сложные выражения только в одном из уравнений. ОС (1) соответствует СС, которая имеет вид

$$\begin{cases} \dot{y}_1 = y_2, \\ \dot{y}_2 = y_3, \\ \dot{y}_3 = (N_4 y_1^2 + N_5 y_1 y_2 + N_6 y_1 y_3 + N_7 y_2^2 + N_8 y_2 y_3 + N_{20} y_1^4 + N_{21} y_1^3 y_2) / D_1 y_1, \end{cases} \quad (2)$$

где $N_4 = 720$, $N_5 = -29.33$, $N_6 = -13.67$, $N_7 = 11$, $N_8 = D_1 = 1$, $N_{20} = -10$, $N_{21} = -1$ и $y_1(t) \equiv x_1(t)$. Анализ соотношений между коэффициентами ОС и СС показал, что СС (2) соответствует системе

$$\begin{cases} \dot{x}_1 = a_1 x_1 + a_2 x_2, \\ \dot{x}_2 = b_1 x_1 + b_2 x_2 + b_6 x_1 x_3, \\ \dot{x}_3 = c_0 + c_3 x_3 + c_4 x_1^2 + c_5 x_1 x_2, \end{cases} \quad (3)$$

одним из частных случаев которой является система Лоренца (1).

Поиск только простейших ОС выполнялся с помощью метода перспективных коэффициентов [4]. Было установлено, что существуют 5 частных случаев системы (3), содержащие 7 ненулевых коэффициентов в правых частях уравнений, в том числе и система Лоренца (1). Все системы содержат коэффициенты a_1 , a_2 , b_6 , c_3 , c_5 и одну из пар коэффициентов b_2 , c_0 ; c_0 , c_4 ; b_1 , b_2 ; b_1 , c_4 ; b_2 , c_4 . Каждой из 5 найденных ОС соответствует СС (2) и, следовательно, они способны воспроизвести временной ряд системы Лоренца (1). Несмотря на то, что по одной наблюдаемой переменной невозможно однозначно определить все коэффициенты ОС, но используя аналитические соотношения и численные значения коэффициентов СС, можно однозначно определить значения коэффициентов ОС a_1 , b_2 , c_3 .

Полученные результаты показывают, что несмотря на неоднозначность, по одной наблюдаемой переменной можно однозначно определить часть коэффициентов системы ОДУ. Если о системе нет дополнительной информации, то целесообразно искать только простейшие ОС, используя метод перспективных коэффициентов.

Литература. 1. C. Lainscsek Nonuniqueness of global modeling and time scaling, Phys. Rev. E 84 (2011) 046205. 2. E. N. Lorenz, Deterministic nonperiodic flow, J. Atmos. Sci. 20 (1963) 130–141. 3. G. Gouesbet, Reconstruction of standard and inverse vector fields equivalent to the Rössler system, Phys. Rev. A 44 (1991) 6264–6280. 4. V. Gorodetskyi, M. Osadchuk, Reconstruction of chaotic systems of a certain class, Int. J. Dynam. Control 3 (2015) 341–353.

Гришин И.Ю., Тимиргалеева Р.Р.

Институт компьютерных систем и информационной безопасности ФГБОУ ВО Кубанский государственный технологический университет, Краснодар, Россия

Метод автоматизации управления процессом обнаружения движущихся объектов радиолокационным комплексом

Работа посвящена анализу и разработке методов и алгоритмов управления при обнаружении движущихся объектов в зоне обзора многофункционального радиолокационного комплекса.

Постановка задачи. В процессе работы радиолокационного комплекса условия его функционирования могут изменяться в достаточно широком диапазоне. Количество ситуаций, с которыми может столкнуться радиолокационный комплекс (РЛК), довольно велико и при его проектировании невозможно предусмотреть заранее оптимальное (рациональное) поведение в каждой ситуации.

РЛК обладает ограниченными возможностями для получения и обработки информации о потоках объектов, находящихся в его зоне обзора. Эти ограничения определяются в основном конечными значениями его энергетических и информационных ресурсов.

Современные РЛК, оснащенные фазированными антенными решетками (ФАР), обладают повышенным энергетическим и информационным потенциалом по сравнению с радиолокаторами предыдущих поколений. Обслуживание объектов начинается с их обнаружения в зоне обзора. Чем раньше будет обнаружен объект, тем больше времени может быть выделено на его обслуживание в других режимах (сопровождение, наведение). Указанное обстоятельство является существенным, поскольку объект будет обслужен более эффективно.

Необходимо разработать метод оптимизации управления современными РЛК с ФАР, позволяющий сократить время обнаружения движущихся объектов в зоне обзора.

Метод решения. Использован метод, связанный с дискретизацией процесса обнаружения объектов, как по времени, так и по состоянию. На фиксированном множестве состояний задается закон преобразования функции распределения состояния в виде рекуррентного алгоритма формирования апостериорной плотности. Для оптимизации процесса поиска используется дискретный аналог принципа максимума. Показано, его применение дает не только необходимые, но и достаточные условия оптимальности.

При этом задача управления формулируется следующим образом. На заданном временном интервале организовать поиск таким образом, чтобы в конце этого интервала было минимальным среднее время поиска объекта. Указанная задача сводится к задаче оптимального управления со свободным правым концом, где элементами фазового пространства являются экстраполированные вероятности наличия объекта в ячейке зоны обзора.

Следует отметить, что получить аналитическое решение при имеющихся ограничениях, характерных для реальных РЛК, не удастся. В связи с этим задача решается численно с использованием метода последовательных приближений.

На основе разработанного метода синтезирован алгоритм оптимального управления поиском неизвестного количества движущихся объектов.

Полученные результаты. Для оценки эффективности предложенного метода была разработана имитационная модель процесса функционирования обзорного РЛК с ФАР, в котором реализован предложенный алгоритм поиска неизвестного количества движущихся объектов. Результаты моделирования показали, что при использовании предложенного алгоритма оптимального управления среднее время поиска объектов может быть существенно (в 1,5...3,5 раза) снижено.

Вывод. Разработан метод и алгоритм оптимального управления многофункциональным РЛК, позволяющий существенно снизить время обнаружения объектов в зоне обзора.

Работа выполняется при поддержке РФФИ (грант 18-29-03258).

Денисенко А.И.

Запорожский национальный технический университет, Запорожье, Украина

Оптимизация геометрических параметров теплообменных элементов газовых котлов

Теплообменники газовых котлов обычно представляют собой резервуары с теплоносителем и встроенными вертикальными каналами, через которые проходит высокотемпературная газозвдушная смесь. Для интенсификации теплообмена в каналы помещают турбулизаторы в виде винтообразных вставок. Геометрия турбулизаторов существенно образом влияет на эффективность работы котла. Так, невысокая степень закрутки ленточных вставок уменьшает аэродинамического сопротивления в теплообменнике, что ведет к увеличению расхода газозвдушной смеси и как результат – увеличению доли подмешиваемого холодного воздуха и снижению температуры продуктов горения. Слишком высокая степень закрутки ленточных вставок повышает эффективность теплоотдачи, но с другой стороны увеличенное аэродинамическое сопротивление препятствует эвакуации продуктов горения через дымоход и может привести к отравлению угарным газом.

В работе рассматривалась сопряженная задача конвективного теплообмена в вертикальных каналах с винтообразными ленточными вставками. Численное моделирование выполнялось методом конечных элементов при помощи программного комплекса COMSOL Multiphysics. Проанализировано влияние степени закрутки турбулизаторов на эффективность теплообмена в котле и на расходные характеристики газозвдушной смеси. Как показывают численные эксперименты, наиболее интенсивная теплоотдача происходит в нижней части канала, где присутствует большие температурные градиенты. Эффективность теплообмена в верхней области теплообменника существенно ниже, в силу снижения температурных градиентов. Повысить теплоотдачу в верхней области предлагается путем использования винтообразных вставок с неравномерным шагом закрутки, который уменьшается по высоте канала. Использование таких вставок позволяет перераспределить теплоотдачу по высоте канала, повысить эффективность теплообмена не повышая при этом аэродинамическое сопротивление турбулизаторов. Геометрия винтообразной вставки определяется двухпараметрической зависимостью, параметры которой регулируют степень неравномерности и количество витков турбулизатора.

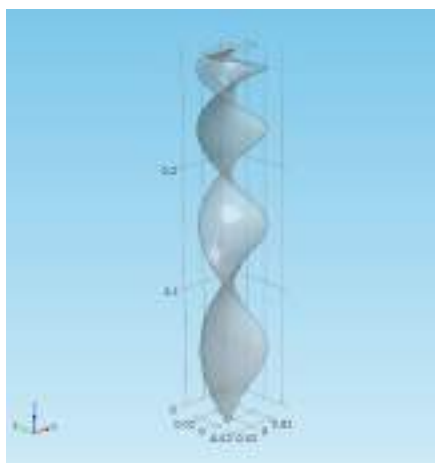


Рис. 1. Турбулизатор с переменным шагом

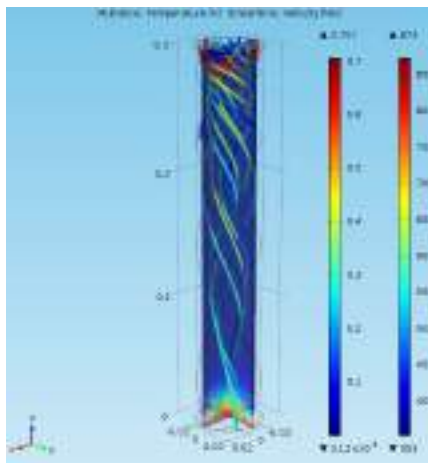


Рис. 2. Распределение температуры и линии тока

Встроенный в COMSOL модуль оптимизации позволяет определять оптимальные геометрические параметры турбулизатора в зависимости от внешних условий и характерных температурных режимов работы котла. На рис. 1, 2 представлены примеры расчета теплообменника.

Духновская К.К., Ковтун О.И.

Факультет информационных технологий КНУ им. Тараса Шевченко, Киев, Украина

Динамическая мера TF-IDF

Актуальность. Согласно статистике международной аналитической компании IDC каждые полтора года, количество данных в сети Internet удваивается. Поэтому важным аспектом исследования информационного поиска является динамика информации, которая публикуется в Internet, в частности динамические модели текстового документа.

Модели текстовых документов. На сегодня широко используется модель текстового документа, основанная на статической мере TF-IDF. Пусть имеем коллекцию текстовых документов и словарь W – упорядоченный набор терминов, мощность которого M . Мощность словаря – это количество терминов, которые в нем содержатся. Тогда документ можно представить в виде вектора [1, 2]:

$$D_i = \langle w_{1i}, w_{2i}, \dots, w_{Mi} \rangle, \quad (1)$$

где w_{ik} – частота k -го термина в i -м текстовом документе из коллекции ($i = 1, M$).

Частота термина рассчитывается по формуле TF-IDF:

$$TF_{ik} = \frac{m_{ki}}{M_{ki}}, \quad (2)$$

m_{ki} – количество вхождений k -го термина в i -м текстовом документе из коллекции, M_i – общее количество терминов в i -м документе;

$$IDF_k = \ln \frac{N}{n_k}, \quad (3)$$

N – общее количество документов в коллекции, n_k – количество документов в коллекции, в которых встречается k -й термин. Величина IDF_k характеризует важность k -го термина в коллекции документов. Тогда мерой TF-IDF будет:

$$w_{ki} = TF_{ki} * IDF_{ki}. \quad (4)$$

Динамические модели. Чтобы получить количественную оценку скорости старения научных публикаций Р. Бартон и Р. Кеблер использовали модель Мальтуса [3]. Данная модель может иметь место при следующих допущениях: а) пусть N – количество текстовых документов по некоторой тематике; б) предполагается, что скорость роста количества текстовых документов этой тематики прямо пропорциональна N .

Последнее предположение вытекает из статистических исследований международной аналитической компании IDC. Если рассматривать модель документа, которая задается формулами (1-4), то одна из составляющих этой модели TF не зависит от времени, при условии, что количество слов в словаре не изменяется со временем. Действительно, TF вычисляется согласно формуле (2), но ни количество вхождений k -го термина в i -й документ, ни общее количество терминов в i -ом документе не изменяются со временем. Динамической составляющей этой модели является величина IDF, которая вычисляется согласно формулы (3).

Пусть $IDF_k(t)$ – важность k -го термина в коллекции документов в момент времени t по данной тематике. Относительное изменение этой величины за отрезок времени t определим, как $R(t)$:

$$R(t) = \frac{IDF_k(t + \Delta t) - IDF_k(t)}{IDF_k * \Delta t}. \quad (5)$$

Тогда:

$$R(t) * IDF_k(t) = \frac{IDF_k(t + \Delta t) - IDF_k(t)}{\Delta t}, \quad (6)$$

при переходе к пределам будем иметь:

$$R(t) * IDF_k(t) = \frac{d(IDF_k(t))}{dt}. \quad (7)$$

Получаем дифференциальное уравнение Бернулли:

$$\frac{dIDF_k(t)}{IDF_k} = R(t)dt. \quad (8)$$

Проинтегрируем левую и правую часть уравнения (8):

$$\ln IDF_k = \int_0^T R(t)dt, \quad (9)$$

или

$$IDF_k = e^{\int_0^T R(t)dt}. \quad (10)$$

Решение уравнения (10) зависит от функции $R(t)$ – относительного изменения важности k -го термина в коллекции документов. Если $R(t)$ – не зависит от времени, т.е. является константой, то получаем модель Мальтуса:

$$IDF_k = IDF_{k0}e^{\alpha T - \beta T}, \quad (11)$$

где IDF_{k0} – важность k -го термина в момент времени $t = 0$; β – коэффициент полураспада актуальности тематики коллекции, определяемый экспертным путем, для каждой тематики отдельно; α – коэффициент роста количества текстовых документов рассматриваемой коллекции. Рассмотрим модель документа (1), где для термина вес w_{ik} определяется формулой (4). Тогда, опираясь на модель Мальтуса, можно получить динамические координаты вектора (1):

$$w_{ik} = TF_{ik}IDF_{k0}e^{\alpha T - \beta T}, \quad (12)$$

где TF_{ik} – локальная частота k -го термина в i -м текстовом документе коллекции, которая определяется формулой (2), IDF_k – инверсия частоты, с которой некоторый термин встречается в коллекции документов, определяется формулой (3).

Результаты. Построена динамическая модель текстового документа на основании TF-IDF, которая может использоваться в информационно-поисковых системах.

Литература. 1. К.К. Духновська Формування пошукового динамічного векторного простору // «Штучний інтелект». – 2016. – №3,4. 2. Пошук інформації // Анісімов А.В, Глибовець А.М., Глибовець М.М., Шабінський А.С. – К.: Національний університет “Києво-Могилянська академія”, 2015. – 283 с. 3. Д.В. Ландэ. Основы интеграции информационных потоков. – К.: ООО «Инжиниринг», 2006. – 235 с.

Євграфова К.Л., Бідюк П.І.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Модифікація підходу до формування ціни у роздрібній торгівлі

Вступ. Загальний обсяг світової роздрібною торгівлі дорівнює приблизно 10 трлн. доларів на рік. Підприємства цієї галузі не зважаючи на свою складність та багато ієрархічність, ще не є високотехнологічними. Багато процесів у цій сфері ще не є автоматизованими, а деякі навіть комп'ютеризованими. Тому через людський фактор і брак економічних та статистичних інструментів допускається безліч помилок.

Аналіз існуючих підходів до ціноутворення. Одним з процесів, що потребують модифікації є ціноутворення. На сьогодні розроблено багато цінових політик і стратегій для підприємств роздрібною торгівлі, але досі не існує ефективної чи універсальної методики ціноутворення. Ефективність сформованої ціни визначається: її адекватністю економічної і фінансової стратегії підприємства, ступенем реалізації завдань цінової політики, успішністю реалізації товару, позитивною зміною рівня рентабельності підприємства та її відповідності якості товару.

Широке поширення мають методики: засновані на витратах, орієнтовані на конкурентів або на попит, і параметричні методи [1, 2]. Розглянувши ці методики детальніше можна зробити висновок, що, навіть усі разом вони відображають лише невеликий діапазон факторів, які впливають на формування ціни (собівартість товару, обсяги продажу товару за попередньою ціною, деякі змінні та постійні витрати, можливу норму рентабельності або властивості продукту). А кожна окрема методика описує вплив тільки одного або пари факторів, не комбінуючи їх та не враховуючи фінансовий стан компанії, її цінову політику і т.д. Усі інші фактори, на яких не концентрується обраний підприємством метод, враховуються «на око» переважно експертами з маркетингу. Це призводить до помилок, що в свою чергу ведуть до втрат компанії, як фінансових так і репутаційних, як одноразових так і в перспективі.

Формування повного списку факторів впливу на ціноутворення. Під ціноутворюючими факторами розуміється сукупність різних змінних умов, які впливають на формування рівня, структури і динаміки цін, визначаючи їх тенденцію. Всі ці умови можна поділити на внутрішні, фактори на які підприємство може впливати, і зовнішні, зовнішні чинники впливу [1–3].

До внутрішніх відносяться наступні фактори

- Цілі підприємства (вижити, утримати або зайняти лідируючу позицію і т.п.);
- Витрати на зберігання та транспортування;
- Управління витратами (заробітні плати персоналу, оренда, дослідження та ін.);
- Маркетингові дослідження ринку (вивчення конкурентного середовища, рекламної діяльності, підвищення іміджу фірми та ін.);
- Організація механізму роботи (вдосконалення роботи з постачальниками і споживачами, вибір системи оплати продукції, оптимізація організаційної структури та ін.);
- Стратегія реалізації певної групи товарів (методи просування продукції);
- Фінансові можливості підприємства;
- Організація ціноутворення (можливість оперативного і еластичного регулювання цін);
- Супутні послуги, що надаються разом з товаром або протягом його життєвого циклу (технічне обслуговування, гарантійний ремонт);
- Позиціонування (якщо товар вважається загальнодоступним, то і ціна буде не високою, але якщо товар позиціонується, як клас «Люкс», ціна відповідно зростає).

До зовнішніх факторів належать наступні

- Собівартість товару;
- Життєвий цикл товару (новий товар в момент надходження може мати дуже високу вартість, а на заході стане дешевою, те саме відноситься і до обслуговування товару);
- Державне регулювання (податки, мито, пільги, антимонопольні заходи, граничні ціни);
- Споживча поведінка, її сезонність (враховуючи сезонність можна проводити акції на не

- сезонні товари, стимулюючи продажі, та піднімати ціни на сезонні, збільшуючи маржу);
- Залежність попиту від специфічних умов (окрім сезонності попиту має значення наявність точок ремонту та обслуговування, величина експлуатаційних витрат споживача та ін.);
 - Рівень конкуренції на ринку, цінова політика конкурентів (на основі власного розуміння положення на ринку та конкурентної поведінка формується цінова політика компанії);
 - Макроекономічні (фаза економічного циклу, стан платоспроможності та ін.);
 - Рівень попиту і пропозиції товару на ринку, наявність товарів аналогів, їх ціна та якість;
 - Споживчі якості товару (як і позиціонування може мати прямий вплив на ціну);
 - Суб'єктивні вподобання цільової аудиторії;
 - Сприйняття покупцями ціни та цінності товару (висока ціна може бути виправдана такою ж високою цінністю товару, тривалим циклом життя або вірним позиціонуванням);
 - Еластичність попиту (крива попиту показує залежність між можливою кількістю продажів товару і його ціною, іншими словами, показує кількість товару, яке бажає купити цільова аудиторія при різних рівнях роздрібною ціни);
 - Стан економіки (в період економічних криз зростає попит на товар економ-сегмента, і зростає чутливість споживача до ціни);
 - Правові норми на ринку (в країні можуть існувати певні закони, що забороняють цінову дискримінацію або встановлюють максимальний поріг цін на певні види товарів).

Модифікований підхід. З метою уникнути негативних наслідків методів, що використовуються на даний момент, та створити оптимальну ціну, що максимально відповідає потребам підприємства і бажанням користувача запропонована система підтримки прийняття рішень по формування ціни. Вона складається з наступних етапів.

Перший етап. Перший етап полягає у зборі даних: стосовно власного фінансового становища, кількості конкурентів і їх поведінки на ринку, також аналізуються продажі за останні декілька місяців або рік, якщо є можливість, на сезонність (вплив таких факторів, як погода, зміна курсу валют, важливі громадські події, свята та ін.). На основі зібраних даних визначається мета компанії, цінова політика та стратегія. Параметрично задається гранична ціна (наприклад, при створенні власної ціни не перевищувати максимум серед конкурентів при політиці з орієнтацією на максимальний поточний прибуток). Також виходячи з прогнозу і сезонності обирається позиціонування товару, наприклад, як акційного.

Другий етап. Другий етап – формування базової ціни продукту з урахуванням собівартості, витрат на зберігання і транспортування (з перерахунком на одиницю товару), усі можливі державні врегулювання (ПДВ, мито і т.д.). Тобто першим і другим етапом створюється діапазон граничних цін.

Третій етап. На третьому етапі будується багатокритеріальна модель і враховуються усі інші фактори. Кожен зводиться до параметричного вигляду та має свою вагу, що встановлюється підприємством виходячи з власних цілей. Таким чином враховується обслуговування товару з боку підприємства протягом життєвого циклу, витрати на управління та дослідження, а з іншого боку загальний стан економіки та споживчої спроможності, еластичність попиту, прогнози продажу та сприйняття його цінності споживачем. На виході пропонується ціна, або вузький діапазон цін, що відповідає усім факторам впливу та вимогам підприємства.

Висновки. Запропонована система підтримки прийняття рішень потребує навчання персоналу для її використання, проте вона може значно зменшити людський фактор і помилковість сформованих цін, пришвидшити процес ціноутворення, зробити його більш гнучким і відповідним реаліям ринку, за рахунок чого збільшити загальну рентабельність підприємств роздрібною торгівлі.

Література. 1. Помитов С.А. Особенности установления цен предприятиями розничной торговли. М.: Экономическая школа, 2014. – С. 103–152. 2. Шкварчук Л.О. Ціни і ціноутворення: навч. посібник. К.: Кондор, 214 с. 3. Heidhues P., Koszegi B. Regular prices and sales. Forthcoming: Theoretical Economics.

Жданова О.Г., Маленко А.О., Сперкач М.О.

КПІ ім. Ігоря Сікорського, ФІОТ, Київ, Україна

Складання розкладу виконання завдань паралельними пропорційними пристроями з метою мінімізації сумарного відхилення від спільного директивного строку

Постановка задачі. Задано множину завдань $I = \{1, 2, \dots, i, \dots, n\}$ та кількість пристроїв m . Пристрої працюють паралельно, є взаємозамінними і відрізняються один від одного продуктивністю виконання завдань. Пристрої можна впорядкувати за швидкістю виконання завдання і цей порядок однаковий для всіх завдань. Для пристрою j , $j = \overline{1, m}$ задано коефіцієнт продуктивності h_j такий, що тривалість виконання завдання i на пристрої j дорівнює $h_j p_i$. «Еталонним» будемо називати пристрій з коефіцієнтом продуктивності $h = 1$. У цьому випадку величина p_i є тривалістю виконання завдання i на еталонному пристрої. Передбачається, що всі завдання множини I надходять одночасно в нульовий момент часу та мають спільний директивний строк d , процес обслуговування кожного завдання проходить без переривань до завершення обслуговування. Пристрої можуть починати свою роботу в ненульовий момент часу. Необхідно побудувати розклад, що мінімізує сумарне відхилення виконання завдань від директивного строку.

Введемо позначення: C_i – момент закінчення завдання i ; $Z_i = \max\{0, C_i - d\}$ – запізнення завдання i ; $E_i = \max\{0, d - C_i\}$ – випередження завдання i ; $Z = \sum_{i=1}^n Z_i$, – сумарне запізнення; $E = \sum_{i=1}^n E_i$, – сумарне випередження; u – момент запуску розкладу (момент початку виконання найпершого із завдань).

З урахуванням позначень, цільова функція задачі (мінімізація сумарного відхилення моментів закінчення завдань від директивного строку) має вигляд: $E + Z \rightarrow \min$.

Розроблений для декількох пропорційних пристроїв алгоритм базується на алгоритмі складання розкладу для декількох ідентичних паралельних пристроїв, який, в свою чергу, базується на алгоритмі складання розкладу для одного пристрою. Тому доцільно почати зі складання розкладу для одного пристрою.

Складання розкладу для одного пристрою. В цьому випадку сумарне випередження завдань у розкладі таке: $E = \sum_{i \in W} E_i = \sum_{i=1}^{w-1} (w-i)p_{[i]}$, де W – множина завдань, що виконуються з випередженням, $|W| = w$; $[i]$ – номер завдання, що є i -тим в послідовності завдань множини W , якщо рахувати від директивного строку справа наліво по часовій осі. Сумарне запізнення завдань: $Z = \sum_{i \in Q} Z_i = \sum_{i=1}^q (q-i+1)p_{[i]}$, де Q – множина завдань, що запізнюються, $|Q| = q$; $[i]$ – номер завдання, що є i -тим за порядком в послідовності завдань, що запізнюються. Отже, цільова функція (ЦФ) визначається так: $E + Z = \sum_{i=1}^{w-1} (w-i)p_{[i]} + \sum_{i=1}^q (q-i+1)p_{[i]}$ [1].

1. Умови поліноміальної розв'язуваності задачі з одним пристроєм

а) Достатньо велике значення директивного строку

Згідно [2] існує процедура поліноміального розв'язання за умови:

$$d \geq \Delta = p_n + p_{n-2} + \dots + p_{n-2\lceil(n-2)/2\rceil} \quad (1)$$

При виконанні умови поліноміальної розв'язуваності (1), для мінімізації сумарного відхилення необхідно розподіляти завдання на пристрої за правилом «найдовшому завданню повинен відповідати найменший коефіцієнт ЦФ».

• Алгоритм №1

Крок 1. Впорядкувати завдання за незростанням тривалостей ($p_1 \geq p_2 \geq \dots \geq p_n$); $l_w = 1$ – ініціалізація порядкового номеру позиції для завдань, що виконуються з випередженням; $l_q = n$ – ініціалізація порядкового номеру позиції для робіт, що запізнюються.

Крок 2. Цикл по завданнях i від 1 до n

ЯКЩО i – непарне ТО включити завдання i до множини W на позицію l_w , $l_w = l_w + 1$;
ІНАКШЕ включити завдання i до множини Q на позицію l_q , $l_q = l_q - 1$.

- б) Достатньо мале значення директивного строку

Якщо директивний строк достатньо малий, а саме: $d < \min_i p_i$, то у будь-якому розкладі усі завдання запізняються, тоді SPT-упорядкування дає оптимальний розклад [3].

2. Умови не поліноміальної розв'язуваності задачі

За відсутності достатнього люфту при складанні розкладу, тобто при не виконанні умови (1), задача відноситься до класу NP.

В основу розробленого алгоритму покладено алгоритм розв'язання задачі при виконанні умови (1). Попередньо упорядковуються завдання за незростанням. Потім за Алгоритмом №1, при визначенні позиції завдання у розкладі, враховується яким чином це вплине на допустимість розкладу. Алгоритм №1 виконується до тих пір, доки тривалість поточного завдання не почне перевищувати проміжок часу, що залишився між можливим моментом початку виконання завдання та директивним строком. В такому випадку крок 2 Алгоритму №1 модифікується таким чином: при визначенні місця поточного завдання виконується порівняння значення його тривалості з проміжком часу, що залишився між можливим моментом початку виконання завдання та директивним строком.

Складання розкладу для декількох ідентичних пристроїв. Розглянемо тепер випадок, коли в системі є $m > 1$ ідентичних пристроїв. Введемо позначення: Q_j – множина завдань, що запізняються на пристрої j , $|Q_j| = q_j, j = \overline{1, m}; j[i]$ – номер завдання, що виконується на пристрої j та є i -тим за порядком в послідовності завдань, що запізняються на цьому пристрої; W_j – множина завдань, що на пристрої j , виконуються з випередженням, $|W_j| = w_j, j = \overline{1, m}; j[i]$ – номер завдання, що є i -тим з кінця в упорядкуванні завдань з множини W_j .

Сумарне випередження: $E = \sum_{j=1}^m \sum_{i=1}^{w_j-1} (w_j - i) p_{j[j[i]]}$; сумарне запізнення: $Z = \sum_{j=1}^m \sum_{i=1}^{q_j} (q_j - i + 1) p_{j[j[i]]}$; сумарне відхилення: $E + Z = \sum_{j=1}^m \sum_{i=1}^{w_j-1} (w_j - i) p_{j[j[i]]} + \sum_{j=1}^m \sum_{i=1}^{q_j} (q_j - i + 1) p_{j[j[i]]}$.

1. Умови поліноміальної розв'язуваності задачі з ідентичними пристроями

- а) Достатньо велике значення директивного строку

Узагальнимо умову (1) поліноміальної розв'язуваності для цього випадку:

$$d \geq \Delta = p_n + p_{n-2m} + \dots + p_{n-2m[(n-2m)/2m]}. \quad (2)$$

При виконанні умови поліноміальної розв'язуваності (2), для мінімізації сумарного відхилення необхідно розподіляти завдання між пристроями за правилом «найдовшому завданню повинен відповідати найменший коефіцієнт ЦФ».

• Алгоритм №2

Крок 1. Впорядкувати завдання за незростанням тривалостей ($p_1 \geq p_2 \geq \dots \geq p_n$)

Крок 2. Впорядкувати значення коефіцієнтів ЦФ за зростанням

Крок 3. ЦИКЛ по завданням i від 1 до n

Завдання i призначити на пристрій з найменшим вільним коефіцієнтом ЦФ (який визначає позицію завдання в упорядкуванні відповідного пристрою).

Згідно Алгоритму № 2, спочатку заповнюються перші позиції всіх пристроїв (ім відповідають коефіцієнти 0), потім останні позиції всіх пристроїв (ім відповідають коефіцієнти 1), після цього – другі позиції всіх пристроїв (ім відповідають коефіцієнти 1) і т.д.

- б) Достатньо мале значення директивного строку

Якщо директивний строк достатньо малий, а саме: $d < \min_{1 \leq i \leq m} p_i$, то у будь-якому розкладі усі завдання запізняються, тоді SPT-упорядкування дає оптимальний розклад.

2. Умови не поліноміальної розв'язуваності задачі

При невиконанні умови (2), задача відноситься до класу NP.

В основу розробленого алгоритму покладено алгоритм розв'язання задачі при виконанні умови (2). Попередньо упорядковуються завдання за незростанням. Потім за Алгоритмом

№2, при визначенні позиції завдання у розкладі, враховується яким чином це вплине на допустимість розкладу. Алгоритм №2 виконується до тих пір, доки тривалість поточного завдання не почне перевищувати проміжок часу, що залишився між можливим моментом початку виконання завдання та директивним строком. В такому випадку крок 3 Алгоритму №2 модифікується таким чином: при визначенні місця поточного завдання виконується порівняння значення його тривалості з проміжком часу, що залишився між можливим моментом початку виконання завдання та директивним строком.

Складання розкладу для декількох пропорційних пристроїв.

1. Випадок поліноміальної розв'язуваності задачі

У цьому випадку оцінити в закритій формі (подібній формулам (1), (2)) максимальний сумарний час тривалостей завдань, що йдуть з випередженням на пристроях неможливо. Спочатку розглянемо поліноміальний алгоритм побудови оптимального розкладу завдань за умови, коли d є достатньо великим числом, тобто таким, що у будь-якому розкладі $u \geq 0$.

У випадку паралельних пропорційних пристроїв [1]: сумарне випередження: $E = \sum_{j=1}^m h_j \sum_{i=1}^{w_j-1} (w_j - i) p_{j[i]}$; сумарне запізнення: $Z = \sum_{j=1}^m h_j \sum_{i=1}^{q_j} (q_j - i + 1) p_{j[i]}$; сумарне відхилення: $E + Z = \sum_{j=1}^m h_j \sum_{i=1}^{w_j-1} (w_j - i) p_{j[i]} + \sum_{j=1}^m h_j \sum_{i=1}^{q_j} (q_j - i + 1) p_{j[i]}$.

Назвемо складеним коефіцієнтом ЦФ добуток коефіцієнта продуктивності пристрою на ціле число, яке відповідає позиції завдання у розкладі. Для мінімізації сумарного відхилення необхідно розподіляти завдання між пристроями за правилом «найдовшому завданню повинен відповідати найменший складений коефіцієнт ЦФ».

• Алгоритм №3

Крок 1. Впорядкувати завдання за незростанням тривалостей ($p_1 \geq p_2 \geq \dots \geq p_n$)

Крок 2. Сформувані можливі значення складених коефіцієнтів ЦФ та впорядкувати їх за зростанням

Крок 3. ЦИКЛ по завданням

Завдання i (завдання з поточної множини непризначених завдань, яке має максимальну тривалість) призначити на ту позицію у розкладі, якій відповідає найменший вільний складений коефіцієнт.

Якщо розклад, побудований за Алгоритмом №3, є допустимим (момент запуску розкладу невід'ємний) – то він є і оптимальним.

2. Загальний випадок

Розроблений алгоритм є модифікацією Алгоритму №3. Кроки 1 та 2 залишаються без змін. На кроці 3 при визначенні позиції випереджаючого завдання у розкладі враховується яким чином це вплине на допустимість розкладу (для кожного пристрою сумарна тривалість випереджаючих завдань не повинна перевищувати значення директивного строку): поточне завдання включається у множину W_j , якщо його тривалість з урахуванням коефіцієнту продуктивності пристрою j не перевищує проміжок часу, що залишився між можливим моментом початку виконання завдання та директивним строком. Якщо поточне завдання не може бути включене до множини W , то воно включається до множини Q на позицію з найменшим вільним складеним коефіцієнтом.

Література. 1. Задача мінімізації сумарного відхилення від спільного директивного строку при виконанні завдань паралельними пристроями / А.В. Годна, А.О. Маленко, О.Г. Жданова, М.О. Сперкач. // Науковий огляд. – 2017. – №9(14). – С. 14–32. 2. Ventura J.A. An improved dynamic programming algorithm for the single-machine mean absolute deviation problem with a restrictive common due date / J.A. Ventura, M.X. Weng. // Operations Research Letters. – 1995. – №17(3). – Р. 149–152. 3. Конвей Р.В. Теория расписаний / Р.В. Конвей, В.Л. Максвелл, Л.В. Миллер. – Москва: Наука. Главная редакция физ.-мат. литературы, 1975. – 360 с.

Заводник В.В., Костах Т.В.

КПИ им. Игоря Сикорского, ИПСА, Киев, Украина

Оценивание состояния демографических процессов на Украине

Демографической статистике, отражающей количественные показатели демографических процессов происходящих в социуме (государстве), в частности на Украине, уделено ничтожно малое внимание в работе Государственной службы статистики Украины [1]. А эти процессы, наравне с мощностью потоков инвестиций являются первообразными всех процессов протекающих в социуме: общественно-политических; социально-экономических; культурно-бытовых. Более того, как показано в [2], будущее страны и нации определяется двумя параметрами: – текущим количеством молодежи (в процентах от общего населения), особенно количество мальчиков и юношей; – текущим состоянием рынка труда и его динамикой.

Начиная с 2014 года, Государственная служба статистики публикует данные по численности населения Украины без учета «временно оккупированной территории республики Крым и г. Севастополя», а данные по природному и миграционному движениям населения без учета «временно оккупированной территории республики Крым и г. Севастополя и временно оккупированной территории Донецкой и Луганской области». Европейская комиссия в своих ежегодных статистических сборниках [3] сообщает, что за прошедшие 5 лет Европейский Союз выдал вид на жительство (работа, учеба, воссоединение с семьей и т.д.), в первый раз, для более чем 2,3 млн. граждан Украины, в том числе Польша около 2 млн. Количество постоянно работающих граждан Украины в России и Белоруссии неизвестно. Возникает вопрос, а какова истинная, реальная численность населения, проживающая на территории Украины, контролируемой центральной властью?

В представленном докладе приводятся результаты работы по оцениванию реального количества постоянного населения Украины, начиная с 2014 г. по 2017 г.

В качестве исходных данных исследования взяты следующие данные (показатели) [1, 3]: численность населения, на конец года; рождаемость; смертность; продолжительность жизни; возрастной состав населения; экономически активное население; количество полных ставок; стоимость полной ставки (одного рабочего места) по затратам основного капитала (рассчитывается по таблице данных межотраслевого баланса); потоки внутренней (в пределах Украины) и внешней (трансграничной) миграций; введения в эксплуатацию жилья (площадь, количество квартир); рынок мобильной связи; количество официальных (законных) мигрантов в странах ЕС с Украины.

В работе не применялись методы статистического анализа. По экономической динамике выше перечисленных показателей решалась задача восстановления функциональных зависимостей между этими показателями и количеством населения Украины.

В результате проведенного анализа и моделирования получены следующие оценки:

- за последние 4 года население страны уменьшилось на 3,4 млн. чел. и составляет 42,387 млн. чел. (официальная статистика), а интервальная оценка показывает – от 34 до 37,6 млн. чел.;
- по динамике индексов рождаемости и смертности (рис. 1) население оценивается в 37,086 млн. чел., на конец 2017 г. (рис. 2);
- население территории подконтрольной украинской власти Донецкой и Луганской областей уменьшилось на 3,7 млн. чел.;
- население семи западных областей Украины (Волинской, Ровенской, Львовской, Ивано-Франковской, Тернопольской, Закарпатской и Черновицкой) уменьшилось на 1,5 млн. чел.;
- рост населения происходит только в двух регионах, в столице и в Киевской области, что обусловлено двумя факторами: жилищным строительством, около 40 процентов жилья от общеукраинского уровня: непрерывно увеличивающемся рынке труда, тогда как общеукраинский рынок сокращается;

- количество мигрантов (работа, учеба) в ЕС непрерывно увеличивается, в 2017 г. составило более 0,65 млн. чел.

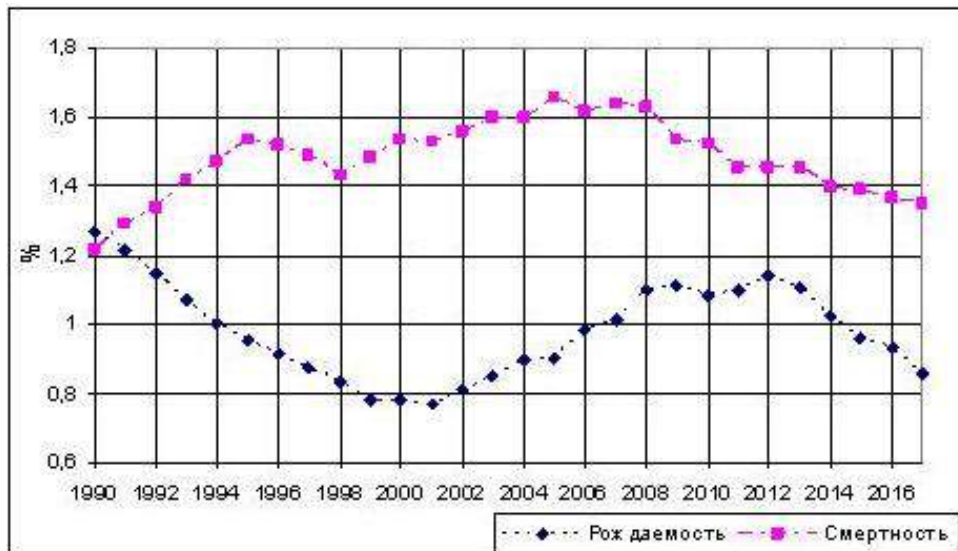


Рис. 1. Динамика изменения коэффициентов рождаемости и смертности на Украине

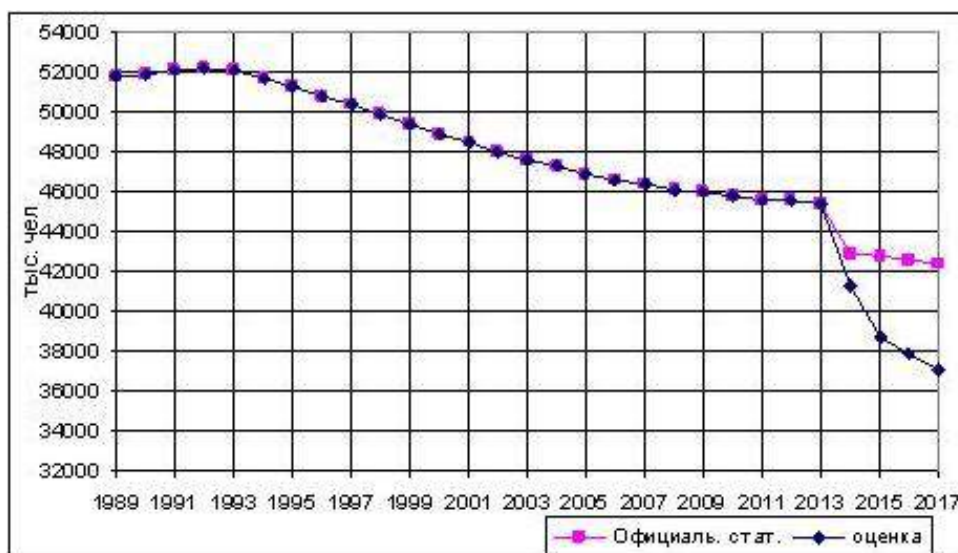


Рис. 2. Количество населения на Украине, на конец года

Общий вывод, создается урбанизированная мега зона (Киев и окружающие города спутники), остальные 23 административно-территориальные региона, в демографическом смысле, кто медленно, кто быстро – угасают.

Доклад проиллюстрирован большим количеством таблиц, графиков, рисунков, выполненных по результатам решения поставленной задачи.

Литература. 1. Государственная служба статистики Украины. – Электронный ресурс, <http://ukrstat.gov.ua/>. 2. Gunnar Heinsohn Söhne und Weltmacht. Terror im Aufstieg und Fall der Nationen. – Zürich: Orell Füssli Verlag AG, 2003. – 191 p. 3. Европейская комиссия: европейская статистика. – Электронный ресурс, <http://ec.europa.eu/eurostat>.

Зайченко А.Є., Савченко І.О.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Система підтримки прийняття рішень на основі мережі морфологічних таблиць для аналізу лісової пожежі

У наш час проблема безпеки та усунення наслідків природних катастроф стає все більш актуальною, адже стихійне лихо — це надзвичайне природне явище, що діє з величезною руйнівною силою, завдає вагомої шкоди місцю, в якому відбувається, порушує нормальну діяльність населення, знищує будівлі, споруди, загрожує життю і призводить до загибелі людей, тварин, знищення матеріальних цінностей.

Ризик виникнення катастроф на території України наразі є достатньо високим. У зв'язку з техногенними факторами і урбанізацією масштабність наслідків і руйнівний потенціал стихійних лих постійно зростає, тому задача запобігання виникненню надзвичайних ситуацій, ліквідації або мінімізації їх наслідків є пріоритетною.

Стихійні лиха як явища характеризуються суттєвою невизначеністю і великою кількістю ключових факторів, які описують їх протікання. Для дослідження таких подій і прийняття управлінських рішень щодо них використовуються спеціальні методи якісного аналізу, одним з яких є модифікований морфологічний аналіз [1]. Однак для задачі оцінювання надзвичайних ситуацій може бути недостатньо описаної в роботі [1] двохетапної процедури методу, оскільки в якості об'єктів морфологічного моделювання доцільно розглядати як саму ситуацію, так і потенційні її наслідки. Між цими об'єктами існує односторонній зв'язок, і параметри рішення повинні враховувати можливі конфігурації обох цих об'єктів. Для такої задачі доцільно застосувати підхід, оснований на використанні мереж морфологічних таблиць [2].

Ціллю даної роботи є розробка мережі морфологічних таблиць для такого стихійного лиха як лісова пожежа, за допомогою якої можливо як оцінити загальну підготовленість до цього типу лих в межах певних регіонів або країни в цілому, так і визначити наслідки пожежі з фіксованими параметрами та запропонувати пом'якшувальні міри.

Мережа морфологічних таблиць складається з трьох рівнів (рис. 1):

- на першому рівні знаходиться таблиця, що складається з параметрів катастроф, що є спільними для всіх та групи таблиць зі спеціальними параметрами для кожного лиха;
- на другому рівні знаходиться таблиця, яка складається з можливих наслідків цих катастроф;
- на останньому рівні – таблиця, що містить заходи, що можуть пом'якшити наслідки стихійних лих.

На підготовчому етапі оцінюються взаємозв'язки всередині таблиць і зв'язки між таблицями.

Далі отримуємо оцінки від експертів для пожеж чи конкретної пожежі для складених таблиць.

На наступному кроці потрібно перерахувати оцінки альтернатив за допомогою процедур модифікованого методу морфологічного аналізу [1] для таблиць першого рівня, і на основі цих оцінок розрахувати оцінки альтернатив наступного.

Аналогічно знаходяться оцінки третього рівня за допомогою попередніх.

Для конкретного стихійного лиха – лісової пожежі розроблено наступну мережу морфологічних таблиць.

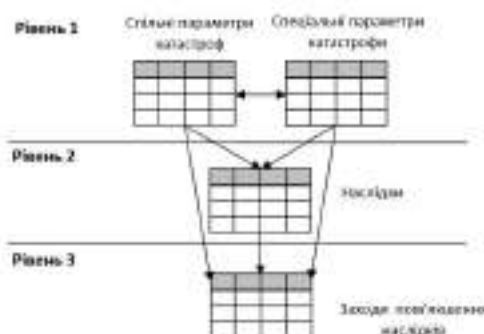


Рис. 1. Мережа морфологічних таблиць для катастроф

Спільні параметри катастроф:

1. Масштаб (окремий об'єкт, невелика площа/квартал/декілька об'єктів, середня площа/район/мале місто, велика площа/місто, область, декілька областей/країна).
2. Місцевість (густа забудова, середня забудова, легка забудова/окраїни, сільська, незаселена).
3. Тривалість дії (більше тижня, 1 тиждень – 3 дні, 3 дні – 1 день, 1 день – 10 годин, 5 – 10 годин, 1 – 5 годин, менше години, одномоментна).

Спеціальні параметри катастрофи [3]:

1. По швидкості пересування вогню (слабкі, середні, сильні).
2. По елементу лісу (низові, верхові, підземні).

Наслідки:

1. Тип загрози (прямий, біологічний, хімічний, радіаційний, психологічний).
2. Ступінь загрози (немає, низький, середній, високий, дуже високий).
3. Будівлі (цілі, переважно цілі, суттєва шкода, зруйновані).
4. Запаси питної води (достатні, частково достатні, недостатні, відсутні/зруйновані).
5. Енергомережа (ціла, частково пошкоджена, сильно пошкоджена).
6. Транспортна доступність (без обмежень, пошкоджені або погані дороги, недоступно для звичайного транспорту).
7. Потенціал заворушень (немає, хвилювання, безпорядок/непокоря/вандалізм, масові заворушення, повстання).

Заходи пом'якшення наслідків:

1. Запобігання дії катастрофи на людей (евакуація, тимчасові укриття, тимчасові інженерні споруди).
2. Необхідна допомога (медична, фізіологічна, їжа/питна вода).
3. Залучені підрозділи (волонтери, місцеві підрозділи МНС, державні підрозділи МНС, національна гвардія).
4. Необхідний транспорт, обладнання (не потребується, автомобілі/автобуси, спеціальні інженерні, гелікоптери, військові).

Створена мережа морфологічних таблиць має суттєві переваги:

- вона допомагає структурувати стихійне лихо за певними характеристиками;
- дозволяє розглянути одночасно велику кількість варіантів;
- будь-який параметр може бути зафіксовано для того, щоб прослідкувати поведінку інших параметрів;
- допомагає визначити наслідки лиха та пропонує пом'якшувальні міри;
- результати дослідження отримуються у зручному для прийняття рішень вигляді ваг альтернатив і можуть бути наглядно представлені за допомогою діаграм або графіків.

Варто зазначити, що розроблена мережа морфологічних таблиць може бути легко поширена на випадок інших стихійних лих шляхом вибору іншої відповідної таблиці спеціальних параметрів катастроф. Таким чином, при належній експертній оцінці вхідних даних створюється потужний аналітичний інструмент для надання підтримки прийняття рішень у організаціях і відомствах, пов'язаних із запобіганням і пом'якшенням наслідків від стихійних лих.

Література. 1. І.О. Савченко. Методологічне і математичне забезпечення розв'язання задачі передбачення на основі модифікованого методу морфологічного аналізу // Системні дослідження та інформаційні технології. 2011. – № 3. – С. 18–28. 2. I. Savchenko. Modeling disasters using networks of morphological tablets. MFOI'2017. November 9–11. Chisinau, Republic of Moldova, 2017 – pp. 162–165. 3. Під ред. Л.А. Михайлова. Надзвичайні ситуації природного, техногенного та соціального характеру і захист від них. / Підручник для вузів. — СПб.: Пітер. 2009. — 240с. — С. 45–50

Касумова Т.А.

Научно-исследовательский институт аэрокосмической информатики НАКА, Баку, Азербайджан

Роль и место космической информации в формировании общества и культуры информационного мира

Сегодня страны мира живут в переходный период от промышленного общества в информационное. В связи с появлением современных информационных технологий, развитием рыночной экономики и формированием новых рыночных отношений в культурных, экономических и социальных областях общества происходят коренные перемены. Современные технологии информации и связи, всестороннее развития страны, повышение интеллектуального уровня населения и активное привлечение республики в интеграционные процессы открывают новые возможности формирования информационного общества.

Благодаря информационно-коммуникационным технологиям создание и развитие космической промышленности имеет высокие достижения. Развитие космической промышленности и вывод на орбиту телекоммуникационных спутников, многоцелевых спутников устраняет зависимость информационного обмена от других стран и обеспечивает информационную безопасность.

Развитие космической промышленности Азербайджана сделало возможным сотрудничество с 10-ю развивающимися странами. Азербайджан является членом космического клуба, одной из первых стран СНГ, имеющей телекоммуникационные спутники в Кавказском регионе. Главной целью государственной программы является, используя достижения космической промышленности, обеспечить потребности государственных структур в спутниковой связи, потребности населения регионов в радиовещании, увеличить каналы международной связи страны и полезное использование космического пространства, а также развивать социальные, экономические, культурные, научные и другие области, а также безопасность страны.

В настоящее время развитие многих областей общества связано с ростом информационных ресурсов и возможностей компьютеров и интернета. Будущее любой страны определяется не только богатством природных ресурсов, а так же людскими ресурсами, имеющими высокий научно-культурный и информационный интеллектуальный потенциал [1].

Сегодня граждане информационного общества должны стать носителями информационной, технологической культуры и стратегически мыслить для обнаружения и решения проблем. Одной из главных задач является подготовка кадров, имеющих информационную культуру для создания и развития космической промышленности. Для благополучного исполнения государственной программы имеет особое значение подготовка кадров по эксплуатации и управлению спутником, производство космического оборудования, получение и обработка космической информации. Имея в виду, что космическая промышленность является новой областью, необходимо за короткий срок подготовить в отечественные кадры по этим направлениям.

Быстрое развитие ИКТ и телекоммуникации в мире требует принять меры в направлении повышения конкурентной способности спутниковых сетей. Кроме этого, создание супер многофункционального компьютерного центра по принятию и обработки космической информации, требует подготовку специалистов в области космической промышленности и спутниковых систем. Экономика государства и вооруженные силы требуют расширения фундаментальных научных знаний. Освоение космического пространства является одним из определяющих признаков уровня развития экономической и национальной безопасности [2].

Культура информационного мира. Наша жизнь постоянно изменяется и наполняется новыми технологиями. Но каждый человек, живущий в век информационных технологий должен владеть информационной культурой.

В рамках культурологического подхода информационная культура рассматривается как способ жизнедеятельности человека в информационном обществе, как составляющая процесса формирования культуры человечества. В зависимости от субъекта, который выступает носителем информационной культуры, последнюю можно рассматривать на трех уровнях: информационная культура личности; информационная культура отдельных групп сообщества (определенного социума, нации, возрастной или профессиональной группы и т.д.); информационная культура общества в целом. В современных исследованиях информационной культуры преобладает информационный подход, поскольку данная проблематика пришла в науку из

информационной сферы [3].

Вопрос о воздействии техники на культуру приобрел особую актуальность в условиях формирования информационного общества, когда радио и телефоны, телевизоры и магнитофоны, принтеры и сканеры, калькуляторы и компьютеры, международная информационная система интернет и специализированные сети оказывают все возрастающее и сложное воздействие не только на сферу материального производства, социально-политическую жизнь общества, но и на ее духовную жизнь.

Не трудно заметить в этих радикальных рассуждениях необоснованное отрицание всей прошлой культуры, благодаря которой совершалась социальная эволюция человека, все его достижения, в том числе и научно-технические. В действительности, речь должна идти о диалектическом снятии из прошлой культуры всего ценного и построении на этом богатом фундаменте новой культуры, соответствующей реалиям сегодняшнего развития человека. В процессе этой преемственности большая роль принадлежит технике. Техника является материальным носителем социальной наследственности. В истории человеческого общества связь поколений осуществляется через технику.

В первую очередь безбумажная технология находит себе применение в обработке оперативного получения информации, содержание которой быстро меняется и устаревает, что делает просто ненужным ее длительное хранение. Многие предсказывают замену ежедневной газеты круглосуточным выпуском телевизионных новостей. Электронные издательства переживают настоящий бум. Архивные материалы, редкие книги, перенесенные на лазерные диски с помощью каналов спутниковой связи, интернет, телефаксов становятся общедоступными. Наука, и как специфическая форма деятельности людей, и как система используемых в ходе общественной практики знаний, являясь феноменом культуры оказывает на ней огромное воздействие [4].

В конечном счете, информационная культура – есть умение соблюдать должное равновесие между формализуемой и не формализуемой составляющими человеческого знания. С учетом рассмотренного определения необходимо разработать определенные системные требования к подготовке специалистов практически всех отраслей социально-культурной сферы (библиотекарей, музееведов, социальных и информационных работников и др.). На наш взгляд, компьютерное обучение в комплексе с культурологическими и другими дисциплинами дает возможность указанным специалистам выступать в качестве проводников национальной культуры вообще и компьютерной культуры в частности.

Поэтому возникает необходимость в наличии профессионально подготовленных информационных посредников, выполняющих в информационных системах функций, во многом аналогичных функциям библиографов-консультантов при библиотечных каталогах и других традиционных способах поиска информации в библиотечных, музейных, архивных и других документоведческих структурах. Кроме выше названных традиционных методов поиска, эти посредники должны уметь работать с информацией, размещенной на web-серверах интернета. Информационная культура принадлежит национальной культуре и одновременно является приобретением международного опыта. Поэтому, информационная культура – это самостоятельный элемент национальной культуры.

Только люди с высокой информационной культурой способны осуществить точный учет природных, людских, экологических и других ресурсов, что позволяет повысить оперативность и эффективность управленческих решений во всех сферах жизнедеятельности нашей страны. Наряду с другими факторами это способствует процветанию Азербайджана [5].

Литература. 1. А.Ə. Şirin-zadə., T.A. Qasimova., İnformasiya cəmiyyətinin və informasiya mədəniyyəti nin formalaşmasında “Azərbaycanın kosmik sənayesinin əsas funksiyaları”, MAKA-nın Xəbərləri №3(20), 2017, S. 71–74. 2. Т.А. Касумова, М.М. Рашаева. “Роль информационной культуры в создании и развитии информационного космического общества.” – Sistem Analysis and Information Technologies 19-th International Conference SAIT 2017 Kiye, Ukraine, May 22–25, 2017 Proceedings P.281–282. 3. Т.А. Qasimova, S.İ. Hüseynova, M.M. Paşayeva. “Azərbaycan kosmik sənayesinin yaradılmasında və inkişafında informasiya mədəniyyətinin rolu”. Memarlıq, İnşaat və Nəqliyyat sahələrində progressiv texnoloqiyalar mövzusunda elmi-praktik konfrans, Bakı-2016, S. 161–163. 4. В.Н. Сорокина, Культура информационного общества, Антропол., 2008, S. 2–7. 5. www.qooqle.az.

Кірік О.Є.

КПІ ім. Ігоря Сікорського, Київ, Україна

Системний підхід до формування управлінських рішень для великих ієрархічно структурованих систем

Системний підхід до формування управлінських рішень для складних систем різної природи включає процес дослідження та аналізу основних властивостей і об'єктивних тенденцій розвитку; побудову моделей і процедур прийняття рішень; створення інформаційних систем і баз даних. Зміст системного підходу стосовно енергетики конкретизується шляхом виділення таких складових [1]: врахування зовнішніх зв'язків; ієрархічне уявлення про внутрішню структуру об'єктів і процесів управління ними; врахування невизначеностей, обумовлених неповнотою вихідної інформації, багатокритеріальністю та іншими факторами; застосування інформаційних технологій. Системні дослідження енергетичних комплексів проводяться для аналізу тенденцій і закономірностей їх розвитку, прогнозування на різних рівнях – від регіональних, міжрегіональних і національних до міждержавних. Якщо на найвищих рівнях ієрархії енергетика виступає як одна з галузей економіки, то при переході до більш низьких рівней все більш істотним стає облік схемно-структурних особливостей конкретних систем, їх фізико-технічних характеристик і умов функціонування. В кінцевому підсумку йдеться про створення ефективної системи виробництва, транспорту і розподілу палива і енергії як відповідної інфраструктури, що забезпечує надійність і якість енергопостачання споживачів.

Процес створення точних і одночасно прийнятних за складністю моделей призводить до ієрархічних принципів їх побудови, зокрема для послідовного уточнення прийнятих рішень. При цьому декомпозиція загальної задачі на деяку ієрархію взаємопов'язаних підзадач, так само, як і агрегування, може носити двоякий характер. В одному випадку вона відповідає реальній послідовності задач, які можна виділити вже на етапі змістовної постановки вихідної задачі відповідно до реальної просторової структури системи або виходячи з окремих інтервалів часу. В іншому випадку представляє більш формальний поділ задач по будь-яких групах змінних або рівнянь, що враховує їх фізичний зміст в значно меншій мірі.

У роботі розглядаються нелінійні задачі оптимізації функціонування великих просторово або функціонально розподілених систем за наявності в них незалежно функціонуючих підсистем та обмежень на деякі спільні ресурси. Метою роботи є побудова ефективних обчислювальних процедур на основі застосування комбінації апроксимуючих методів нелінійної оптимізації та прийомів декомпозиції. Запропоновано триступеневу ітераційну схему [2], де на верхньому рівні відбувається заміна вихідної задачі послідовністю апроксимуючих задач з адитивно-сепарабельними цільовими функціями. На другому рівні розв'язуються координуючі задачі, що формуються за рахунок інформації, отриманої на найнижчому рівні при розв'язанні локальних блочних підзадач. У випадку задач зі зв'язуючими обмеженнями [2] ітераційний процес базується на корекції двоїстих оцінок ресурсів, що робить для підсистем вигідними плани, що все більше наближаються до оптимального плану всієї системи. В випадку задач зі зв'язуючими змінними [3] відбувається корекція лімітів загальносистемних ресурсів, що надаються підсистемам, тобто показники загальних ресурсів цілеспрямовано змінюються від кроку до кроку таким чином, аби підсистеми працювали все ефективніше з точки зору ефективності роботи всієї системи. Запропонований підхід та побудовані на його основі алгоритми можуть застосовуватися для розв'язання широкого кола проблем, пов'язаних з оптимальним розподілом ресурсів у великих ієрархічно-структурованих системах.

Література. 1. Системные исследования в энергетике /под ред. Н.И.Воропая. Новосибирск.: "Наука", 2010. – 686 с. 2. Кірік О.Є. Підхід до розв'язання нелінійних оптимізаційних задач блочної структури зі зв'язуючими обмеженнями // Наукові вісті НТУУ "КПІ". – Київ, 2015. – № 5. – С. 32–38. 3. Кірік О.Є. Розв'язання нелінійних оптимізаційних задач розподілу ресурсів у великих блочно-структурованих системах зі зв'язуючими параметрами // Системні дослідження та інформаційні технології. – Київ, 2016. – № 3. – С. 72–85.

Кісельова О.М., Притоманова О.М., Падалко В.Г.

Дніпровський національний університет ім. Олеся Гончара, Дніпро, Україна

Про оптимізацію параметрів нейронечіткої моделі експортних відносин між Україною та Китаєм

Актуальність задачі моделювання експортних відносин між Україною та Китаєм обумовлена, перш за все, високим рівнем вимог, які пред'являються до рішень, що приймаються у процесі макроекономічного планування на основі побудованих моделей. У роботі [1] з позицій системного аналізу запропоновано і обґрунтовано підхід до побудови моделей аналізу та прогнозу тенденцій українського експорту товарів до Китаю із застосуванням нейронечітких технологій. Нейронечіткі технології дозволяють розширити можливості моделювання складних об'єктів, процесів, що є дуже актуальним у реальних умовах розвитку економіки України при відсутності достовірних даних, неповної і нечіткої статистичної інформації про об'єкт, складних нелінійних залежностей виходів від входів системи.

Побудова математичної моделі ідентифікації на основі нейронечітких технологій складається з двох етапів [2]: побудови нечіткої моделі ідентифікації за наявними експертно-експериментальними даними і настройки її параметрів з метою мінімізації відхилення між результатами моделювання та експериментальними даними.

Структура нечіткої моделі досліджуваного об'єкта або процесу, яка отримана після першого етапу, відповідає нечіткій базі знань про об'єкт або процес та може грубо описувати шукану залежність. Якщо побудована нечітка модель недостатньо точно описує модельований об'єкт або процес, то необхідно настроїти її, тобто знайти такі параметри нечіткої моделі (параметри функцій належності термів її лінгвістичних змінних, ваги правил нечіткої бази знань та ін.), які мінімізують відхилення між модельними (теоретичними, отриманими за моделлю результатами) та експериментальними даними.

Етап настройки нечіткої моделі в термінах математичного програмування може бути сформульований як задача оптимізації таким чином: знайти вектор параметрів нечіткої моделі, який забезпечує мінімум середньоквадратичного відхилення між теоретичними та експериментальними даними. Зауважимо, що, як правило, для розв'язання задачі оптимізації застосовують класичні градієнтні методи. Однак, ці методи мають такі недоліки: вимагають безперервної диференційовності функції, що мінімізується, погано працюють (повільно сходяться або зациклюються на околі точок мінімуму) у випадках сильно витягнутих яружних цільових функцій.

У даній роботі для настройки нечіткої моделі застосовано метод мінімізації з розтягуванням простору в напрямку різниці двох послідовних узагальнених антиградієнтів (г-алгоритм Н.З. Шора). Перевага г-алгоритму Н.З. Шора, крім того, що він добре працює в задачах оптимізації великих розмірностей для гладких яружних функцій з сильно витягнутими лініями рівня, полягає в тому, що він може застосовуватися і для негладких цільових функцій. У цих випадках застосування класичних градієнтних методів, що вимагають гладкості функцій, може викликати проблеми в збіжності цих методів.

Розроблений підхід на основі нейронечітких технологій для моделювання складних нелінійних залежностей може бути застосований для прийняття науково-обґрунтованих економічних рішень в умовах невизначеності, нечіткої, неповної, неточної, недостовірної інформації.

Література. 1. Кісельова О.М. Моделювання експортних відносин між Україною та Китаєм на основі нейронечітких технологій / О.М. Кісельова, О.М. Притоманова, В.Г. Падалко // XV Міжнародна науково-практична конференція "Математичне та програмне забезпечення інтелектуальних систем". Матеріали конференції, 22–24 листопада 2017 року, Дніпро, Україна, 2017. – С. 94–95. 2. Киселева Е.М. Оценка инвестиционной привлекательности стартапов на основе нейронечетких технологий / Е.М. Киселева, О.М. Притоманова, С.В. Журавель // Проблемы управления и информатики, 2016, №5. – С. 123–143.

Комлев О.О.

Київський національний університет ім. Тараса Шевченка, Київ, Україна

«Земна поверхня»: структура і планетарна роль

“Земна поверхня” нині є досить широке поняття. Нерідко, воно не зовсім відповідає і характеру об’єкту, до якого від початку відносилось. Двохвимірною землею, де живуть люди, давно знайшла відображення в практичній геометрії (планіметрії) Піфагора, Ератосфена. Але тільки відкритий відносно недавно картографічний спосіб горизонталей (ізогипс) дозволив створювати звичні нам топографічні карти і проводити кількісні виміри земної поверхні. Виникла морфометрія, на науковий рівень піднялись картографія і географія, вдосконалювались методики проведення вимірних робіт у різних прикладних галузях (гідрогеологія, гірничо-геологія, видобуток корисних копалин).

Постали перед людством сьогоденні проблеми навколишнього середовища намагаються вирішувати, використовуючи все більше сучасні технологічні можливості. На рівні теорії, цьому сприяє утвердження в науках про Землю морфодинамічної парадигми, основна роль в розвитку якої належить геоморфології – науці про «рельєф Землі». Цифрові моделі рельєфу (ЦМР) земної поверхні візуалізували її як трьохвимірну (додалась координата Н – висота). Вони відкрили новий напрямок – геоінформаційне моделювання геосистем, зосереджених на земній поверхні, що дозволяє здійснювати оперативне, варіативне, більш точне прогнозування в них і між ними переміщень речовини, енергетичних перетворень, інформаційних і ентропійних обмінів. Нині, використання звичних топографічних карт земної поверхні – формалізованої, матеріальної, але неуречовиненої, недостатньо. Морфодинамічна парадигма все більше при складанні прогнозів вимагає враховувати реальну фізичну земну поверхню, яка складена тілами літосфери, гідросфери, біосфери, техносфери, що мають різні консистенції, проникність ззовні, оточені повітрям, водою, льодом. Співвідношення поверхонь різної природи все більше починає враховуватися в прогнозах, які складаються для окремих територій (місто, країна, материк). Нині, поняття «поверхня» в його первинному значенні, зазвичай, не охоплює реальні об’єкти досліджень і, тому, частіше використовуються поняття «простір», «сфера», «оболонка», «тіло» ін.

Попри різноманітну фізичну природу, основною інтегральною ознакою земної поверхні є її морфологія. Земна поверхня – це морфологічна система, в структурі якої проявляються спільно і окремо чинники її формування – космічні, планетарні, ендемічні, екзогенні. Земна поверхня перетинає гравітаційне поле Землі, яке визначає трансформаційні тренди на поверхні речовини, енергії, інформації, ентропії. Для морфологічної системи земної поверхні характерна, зумовлена діючими на неї чинниками, поліструктурність. Це: «висотні рівні», «відцентрові» і «доцентрові» угруповання морфоеlementів, «басейнові геоморфосистеми». «Корені» їх розташовані в літосфері, проектується на земну поверхню, далі проявляються в атмосфері (атмосферні вихори), океані (траєкторії течій), біосфері (рисунки біогеоценозів), антропосфері, техносфері. Земна поверхня мільярди років перманентно змінюється. Змінюється співвідношення води і суші. Суша розвивається циклічно: то збирається в один материк (Пангеї), то розпадається на частини (як зараз). Динамічні інтерпретації структури земної поверхні повинні охоплювати не тільки реальний час, а сягати значно глибше в часі.

«Земна поверхня» – це окрема геосистема, в якій зведені як елементи всі земні сфери і оболонки. Водночас, вона невід’ємний елемент всієї планетарної мегасистеми Землі. Її власні (внутрішньо системні) функції – це частина загальнопланетарного метаболічного процесу, який підтримує гомеостаз Землі. Розуміння захисної ролі земної поверхні дозволяє вести наскрізне прогнозування (з минулого, через сьогодення, у майбутнє), правильно вирішувати нагальні проблеми навколишнього середовища.

Кравчук Є.С., Сергеев-Горчинський О.О.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Методи оцінки семантичної орієнтації тексту

Вступ. Пошук текстової інформації за семантичною подібністю слів потребує застосування семантичного аналізу та побудови семантичних словників прикладних областей. Семантична подібність розраховується як відстань між словами в семантичному просторі [1] при обробці природних мов (natural language processing, NLP) та інформаційному пошуку (information retrieval, IR) [2].

З поняттям семантичної подібності пов'язана метрика семантичної орієнтації (СО), яка потрібна для класифікації СО текстів за середнім значенням СО слів та фраз. Фраза має позитивну СО, якщо містить слова, що є асоційованими з досліджуваною прикладною областю, та негативну СО, якщо слова асоційовані зі сторонніми прикладними областями [3].

Семантична орієнтація фрази розраховується як різниця значень міри взаємної інформації для слів з позитивною СО та слів з негативною СО. Класифікація СО тексту потребує знаходження відповідей на наступні питання [4]:

- які слова чи фрази є релевантними?
- які речення є релевантними?
- як агрегувати слова чи фрази, що були виділені?

Значущі слова та фрази. Дослідження значущості слів здебільшого включає визначення СО прикметників [4], оскільки вони точніше передають зміст тексту [5]. Одним з підходів є пошук узгоджених асоціацій (coordinated associations), так, наприклад, фраза «цікаві та Х» дозволяє припустити, що слово Х має позитивну СО [5]. Якщо при аналізі навчальних масивів даних слово з'являється в оточенні прикметників з позитивною СО, ймовірно прикметник Х має позитивне значення СО.

Для уточнення значущості слів розраховують силу СО за метрикою точкової взаємної інформації (pointwise mutual information, PMI), яка окрім прикметників, враховує комбінації [5]:

- прикметник-іменник;
- прислівник-іменник;
- прислівник-дієслово;

До іншого підходу з пошуку слів, що інформативно відображають зміст фрази, належить метод «латентно-семантичного аналізу» (latent semantic analysis, LSA) для пошуку синонімів та подібних за семантикою слів [6]. Дослідження показали [7], що PMI краще класифікує синоніми, ніж LSA.

Значущі речення. Очевидно, що не всі частини тексту у рівній мірі впливають на його змістовність. Наприклад, в задачі класифікації теми тексту [8], визначення значущих речень може ускладнитись наявністю суджень щодо речей, непов'язаних з темою.

Деяко інша задача полягає в тому, що текст містить в основному значущу інформацію, але потрібно визначати, яка саме інформація є більш значущою, тобто треба відрізнити судження від фактів, чи суб'єктивну інформацію від об'єктивної. Вирішення задачі потребує класифікації речень, що мають позитивну чи негативну СО, та які далі можуть бути використані для визначення змістовного навантаження тексту [9].

Інший аспект значущості стосується частин тексту, що узагальнюють чи охоплюють тему тексту. В [4], ґрунтуючись на припущенні, що окремі частини тексту відрізняються за інформативністю, наприклад, висновки, які зазвичай розміщуються в кінці тексту, запропоновано прикметникам ставити у відповідність вагові коефіцієнти, залежно від розміщення в першій,

другій чи третій частинах тексту [10].

Об'єднання. Враховуючи, що в різних частинах тексту прикметники мають різний ваговий коефіцієнт, найтипівішим є розрахунок середнього значення СО слів. Дослідження показали [11], що врахування граматичних конструкцій в реченнях покращує точність класифікації СО тексту.

Методи об'єднання аналізують граматичні конструкції, зокрема заперечення, підсилення і припущення, як такі, що впливають на інформативність слова, фрази, речення [8]. Після того, як в тексті визначені значущі слова та застосовані методи для обробки речень, наступним кроком є об'єднання СО речень.

Висновки. Узагальнюючи сказане, можна зробити висновок, що першочерговим індикатором змістовного навантаження тексту є прикметники. Враховуючи різноманітне розміщення прикметників в межах речень, їх поєднання з різними типами мовних конструкції та відносне розміщення в структурі тексту, прикметникам рекомендується ставити у відповідність вагові коефіцієнти. Існує два основних підходи для розрахунку семантичної подібності слів – за допомогою методів, заснованих на семантичних зв'язках з застосуванням побудови семантичних мереж та методів корпусної лінгвістики з застосуванням пошуку значущих слів в масивах текстів. Методи засновані на корпусній лінгвістиці включають метод точкової взаємної інформації та метод латентно-семантичного аналізу. Визначення семантичної орієнтації тексту потребує пошуку речень, що узагальнюють чи охоплюють тему тексту.

Література. 1. Harispe S., Ranwez S., Janaqi S., Montmain J. Semantic measures for the comparison of units of language, concepts or entities from text and knowledge base analysis // arXiv preprint arXiv:1310.1285, pp. 30–44, 2013. 2. Mihalcea R., Corley C., Strapparava C. Corpus-based and knowledge-based measures of text semantic similarity // AAAI, pp. 775–780, 2006. 3. Turney P. Turney P. Thumbs Up or Thumbs Down? Semantic Orientation Applied to Unsupervised Classification of Reviews // Proceedings of the Association for Computational Linguistics, pp. 417–424, 2002. 4. Taboada M., Anthony C., Voll K. Methods for Creating Semantic Orientation Dictionaries // Proceedings of Fifth International Conference on Language Resources and Evaluation, LREC, Genoa, Italy. pp. 427–432, 2006. 5. Hatzivassiloglou V., McKeown K. Predicting the Semantic Orientation of Adjectives // Proceedings of the 35th Annual Meeting of the ACL, New York, pp. 174–181, 1997. 6. Dumais S. Latent semantic analysis // Annual review of information science and technology, vol. 38, no. 1, pp. 188–230, 2004. 7. Turney P. Mining the web for synonyms: PMI-IR versus lsa on toefl // Proceedings of the 12th European Conference on Machine Learning, EMCL '01, London, Springer-Verlag, pp. 491–502, 2001. 8. Polanyi L., Zaenen A. Contextual Valence Shifters. Computing Attitude and Affect in Text: Theory and Applications // The Information Retrieval Series, vol 20. Springer, pp. 1–10, 2004. 9. Nigam K., Hurst M. Towards a robust metric of opinion // AAAI spring symposium on exploring attitude and affect in text, pp. 598–603, 2004. 10. Taboada M., Grieve J. Analyzing Appraisal Automatically // AAAI spring symposium on exploring attitude and affect in text, pp. 158–161, 2004. 11. Kennedy A., Inkpen D. Sentiment classification of movie reviews using contextual valence shifters // Computational intelligence, vol. 22, issue 2, Wiley, pp. 110–125, 2006.

Крамов А.А.

Київський національний університет ім. Тараса Шевченка, Київ, Україна

Екстракція структурованих даних з множини HTML-сторінок

Розглянуто класифікацію та принципи роботи основних методів екстракції структурованих даних з множини HTML-сторінок. Виконано експериментальну перевірку певного методу на колекції HTML-сторінок англomовних статей українських наукових журналів.

Вступ. У зв'язку з високою динамікою приросту інформації в мережі Internet, актуальною є задача автоматизованого пошуку необхідних даних серед множини веб-ресурсів. Задача інформаційного пошуку значно ускладнюється через гетерогенність представлення структурованих даних на веб-сайтах. Пошуковими системами рекомендується використовувати метадані (додаткові елементи розмітки) на веб-сторінці згідно зі стандартизованими форматами представлення даних (JSON-LD, RDF тощо) для опису інформації. Незважаючи на переваги використання метаданих, процес їх впровадження на веб-ресурсах поки що знаходиться на проміжному етапі [1], тому актуальним залишається застосування методів автоматизованої екстракції структурованих даних з множини HTML-сторінок.

Методи екстракції структурованих даних. До вхідного набору документів методів екстракції структурованих даних входять HTML-сторінки, сформовані спільним сценарієм, але різними наборами даних. Завданням методів є формування узагальненого шаблону (регулярного виразу), відповідного всім вхідним документам. Застосування сформованого шаблону до множини HTML-сторінок дозволяє екстрагувати набори даних, за допомогою яких ці веб-сторінки були сформовані сценарієм. На рис. 1 [2] зображено класифікацію методів екстракції структурованих даних залежно від ступеню впливу користувача на процес формування шаблону.



Рис. 1. Класифікація методів екстракції даних з веб-сторінок

Варто звернути увагу на *неконтрольовані* методи екстракції структурованих даних з множини HTML-сторінок, які потребують найменшого втручання користувача в процес формування шаблону серед методів, наведених на рис. 1. Серед неконтрольованих методів доцільно використовувати метод Trinity [3]. Метод Roadrunner [4] дозволяє формувати шаблон одночасно тільки з двох документів, а для коректної роботи методу FiVaTech [5] необхідно здійснювати попереднє вирівнювання розмітки вхідних HTML-сторінок згідно зі специфікацією XHTML. На відміну від двох останніх методів Trinity не потребує попередньої обробки вхідного набору документів та дозволяє здійснювати формування шаблону з декількох HTML-сторінок одночасно. Для мінімізації сформованого шаблону автори методу пропонують послідовно здійснити трансформацію виразу в детермінований скінченний автомат [6], мінімізувати його та виконати перетворення автомату в регулярний вираз.

Експериментальна перевірка методу Trinity. Було створено застосування для перевірки ефективності роботи методу Trinity на множині HTML-сторінок. Для формування вхідного набору HTML-сторінок обрано англomовні версії статей українських наукових журналів. Створене застосування складається з двох сценаріїв, написаних мовою програмування Java, які здійснюють наступні функції.

1. Автоматизоване збирання веб-сторінок статей з веб-сайтів наукових журналів.
2. Екстракція змінних значень з множини HTML-документів та оцінка отриманих результатів.

Автоматизоване збирання HTML-сторінок з веб-сайтів журналів здійснено за допомогою реалізації пошукового роботу. Регулярний вираз, сформований методом Trinity, застосовано до всієї множини HTML-сторінок. Екстраговані дані експортовано до бази даних для подальшого аналізу. В табл. 1 наведено результати порівняння даних про статті, отриманих методом Trinity, з даними, отриманими за допомогою аналізу об'єктної моделі HTML-сторінок власноруч.

Табл. 1. Результати аналізу екстрагованої інформації

Назва журналу	Загальна кількість статей	Відсоток статей, які не відповідають шаблону	Відсоток помилково розпізнаних статей
Реєстрація, зберігання і обробка даних	517	2.5%	1.6%
Системні дослідження та інформаційні технології	478	0%	3.7%

Дані про відсоток статей, які не відповідають шаблону, вказують на якість вибору документів серед множини HTML-сторінок для формування узагальненого шаблону. Ненульові показники цього значення дозволяють зробити висновок про наявність таких документів серед множини HTML-сторінок, які мають унікальні елементи розмітки, але не потрапили до вхідних документів методу. Відсоток помилково розпізнаних статей поєднує інформацію про похибку роботи власне методу Trinity, некоректну розмітку певних вхідних документів та успішність вибору HTML-сторінок для формування узагальненого шаблону.

Висновки. Розглянуто класифікацію основних методів екстракції структурованих даних серед множини HTML-сторінок та обґрунтовано доцільність використання методу Trinity відносно інших неконтрольованих методів. Виконано експериментальну перевірку методу Trinity на множині HTML-сторінок, які описують англomовні статті українських наукових журналів. Обраховано похибку застосування методу на тестовій множині вхідних документів та показано можливість використання методу для автоматизованої екстракції структурованих даних з множини HTML-сторінок.

Література. 1. Usage of structured data formats for websites [Електронний ресурс]. – Режим доступу: URL: https://w3techs.com/technologies/overview/structured_data/all. 2. D. Patel, A. Thakkar, A Survey of Unsupervised Techniques for Web Data Extraction // International Journal Of Computer Science. – 2015. – Vol. 6 (2). – P. 1–3. 3. H.A. Sleiman, R. Corchuelo, Trinity: On Using Trinary Trees for Unsupervised Web Data Extraction // IEEE Transactions on Knowledge and Data Engineering. – 2014. – Vol. 26 (6). – P. 1544–1556. 4. V. Crescenzi, G. Mecca, P. Merialdo, RoadRunner: Towards Automatic Data Extraction from Large Web Sites // VLDB 2001: Proceedings of the 27th International Conference on Very Large Data Bases (Italy, Rome, September 11–14, 2001). – P. 109–118. 5. M. Kaye, C.-H. Chang, FiVaTech: Page-level web data extraction from template pages // IEEE Transactions on Knowledge and Data Engineering. – 2010. – Vol. 22 (2). – P. 249–263. 6. Погорілий С.Д. Програмне конструювання: Підручник / Під ред. академіка АПН України Третяка О.В. – К.: ВПЦ, Київський університет, 2007. – 438 с.

Кулик В.В., Кудін Г.І., Коробова М.В.

Київський національний університет ім. Тараса Шевченка, Київ, Україна

Інтерактивне управління відтворювальними процесами в національній економіці на основі системи спрощених балансів

Соціально-економічна динаміка розвитку національної економіки як складної соціально-економічної системи в сучасних умовах характеризується невизначеністю і ризиками, що пов'язано із всім комплексом відтворювальних (соціально-економічних) процесів. В цих умовах застосування відомих макроекономічних інструментів (системи національних рахунків, платіжного балансу, таблиці «витрати-випуск» й ін.) для аналізу і прийняття рішень ускладнене і не може дати позитивних результатів, що пов'язано із оперуванням великими масивами інформації і у зв'язку з цим складністю прийняття стратегічних рішень щодо відновлення економічного потенціалу і економічного росту.

В умовах існуючої нечіткості – не лише в математичному сенсі, а й у політичному і стратегічному вимірах, соціальному й економічному аспектах – пропонується застосовувати агреговані балансові моделі національної економіки, що ідентифікують процеси виробництва ВВП, його розподілу та формування капіталу [1] щодо конкретної досліджуваної національної економіки.

	Виробництво ВВП	Розподіл ВВП	Утворення капіталу	Зовнішній світ
Виробництво ВВП		Кінцеве споживання	Валові інвестиції	Експорт
Розподіл ВВП	ВВП			...
Утворення капіталу		Валові заощадження		...
Зовнішній світ	Імпорт	...	...	...

Відновлення соціально-економічної стійкості і економічної потуги потребує відновлення рівноваги у процесах відтворення внутрішньої економіки. Процес відновлення рівноваги має покроковий інтерактивний характер і проявляється в такому розподілі ВВП (стовпчик розподіл ВВП – агрегат «Кінцеве споживання»), що дозволяє відновити внутрішні джерела формування інвестиційних ресурсів (строчка утворення капіталу – агрегат «Валові заощадження») й таким чином відновити процеси інвестування (стовпчик утворення капіталу – агрегат «Валові інвестиції»).

Достатність внутрішніх заощаджень для задоволення внутрішніх інвестиційних потреб знижує економічні ризики та сприяє активному залученню інвестицій із-за кордону. Період 1999–2005 рр. характерний саме такими умовами для України і відповідно стійким економічним ростом (рис. 1).

Проте пізніше – після перемоги Помаранчевої революції і припливу інвестицій із-за кордону – було втрачене розуміння важливості внутрішніх заощаджень для забезпечення сталості процесів інвестування і економічного росту.

Висновки. Підтримка економічного росту потребує постійної підтримки внутрішніх макроекономічних пропорцій, сприятливих для відтворення капіталу у внутрішній економіці. Система спрощених балансів національної економіки не лише дозволяє виявити базові проблеми відтворення (диспропорції), але і в інтерактивному режимі відновлювати умови рівноваги для економічного росту через регулювання схильності до заощаджень чи до споживання [2].

Література. 1. Кейнс Дж.М. Избранные произведения: Пер. с англ. / Предисл., коммент., сост. А.Г.Худокормов. – М.: Экономика, 1993. – 543с. 2. Менкью Н.Г. Макроэкономика / Н.Г.Менкью. Пер. с англ. – М. Изд-во МГУ, 1994. – 736.

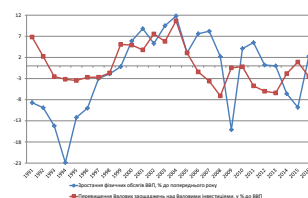


Рис. 1. Зростання реального ВВП та перевищення Валових заощаджень над Валовими інвестиціями

Куценко А.С., Коваленко С.В., Товажнянский В.И.

Национальный технический университет “Харьковский политехнический институт”, Харьков, Украина

Обращение динамических систем для некоторых классов сигналов

Обращение динамических систем, т.е. нахождение входного воздействия по заданному выходу является одной из основных задач теории управления. Решению задачи обращения посвящено множество исследований [1–5], ориентированных на построение обратного оператора, на вход которого подается требуемая функция выхода объекта управления, а на выходе формируется соответствующее управляющее воздействие. В ряде исследований [2, 6] показано, что решение задачи обращения на основе синтеза обратных операторов сопряжено с множеством проблем, среди которых следует выделить проблемы устойчивости, физической реализуемости, робастности и корректности обратных операторов. Перечисленные проблемы не позволяют в общем случае находить практически реализуемое решение задачи нахождения обратного оператора в задаче управления.

В настоящей работе предлагается подход к приближенному решению задачи обращения линейных стационарных динамических систем. Предлагаемый подход основан на аппроксимации входных и выходных сигналов задачи управления непрерывными функциями специального вида, позволяющими использовать методы матричной алгебры для непосредственного определения входных воздействий, соответствующих заданному выходному сигналу.

В качестве пространства Z входных и выходных сигналов предлагается пространство D -функций – линейное пространство непрерывных дифференцируемых функций, производные которых также принадлежат Z . Типичным представителем D -функций является

$$z(t) = \sum_{k=1}^n e^{\alpha_k t} (R_k(t) \sin \omega_k t + Q_k(t) \cos \omega_k t), \quad (1)$$

где α_k и ω_k ($k = \overline{1, n}$) – некоторые постоянные, а $R_k(t)$ и $Q_k(t)$ – векторные многочлены степени не более l . При различных значениях параметров в (1) могут быть получены частные виды D -функций: полиномы, периодические и квазипериодические функции.

Показано, что вынужденные решения линейных дифференциальных уравнений математической модели управляемого процесса при рассмотрении их в среде D -функций можно заменить решением алгебраических линейных матричных уравнений. При этом векторные полиномы $R_k(t)$ и $Q_k(t)$, моделирующие сигналы на входах и выходах системы, представляются в виде матриц коэффициентов при возрастающих степенях t .

Такое представление сигналов позволяет легко решать задачу обращения, а также проводить оценку робастности решения на основе величины числа обусловленности матрицы системы линейных алгебраических уравнений.

Наиболее эффективен предлагаемый метод при обращении SISO систем в среде полиномиальных сигналов. При этом метод достаточно просто реализуется как для математических моделей управляемых процессов в пространстве состояний так и в виде передаточных функций.

Литература. 1. Ильин А.В. Методы робастного обращения динамических систем / А.В. Ильин, С.К. Коровин, В.В. Фомичев. – М. : ФИЗМАТЛИТ, 2009. – 219 с. 2. Костенко Ю.Т. Системы управления с динамическими моделями / Ю.Т. Костенко, Л.М. Любчик. – Х. : Основа, 1996. – 212 с. 3. Пухов Г.Е. Синтез многосвязных систем управления по методу обратных операторов / Г.Е. Пухов, К.Д. Жук. – К. : Наукова думка, 1966. – 218 с. 4. Крутько П.Д. Обратные задачи динамических управляемых систем. Линейные модели / П.Д. Крутько. – М. : Наука, 1987. – 304 с. 5. Критерии обратимости линейных стационарных многомерных систем / В.Т. Борухов // Автоматика и телемеханика. – 1978, вып. 11. – С. 5–11. 6. Гудвин Г.К. Проектирование систем управления / Г.К. Гудвин, С.Ф. Граббе, М.Э. Сальгадо. – М. : БИНОМ. Лаборатория знаний, 2014. – 911 с.

Лабуткина Т.В.

Днепропетровский национальный университет им. Олесь Гончара, Днепр, Украина

Комплекс математических моделей и методов для прогноза и анализа поликонфликтных сближений орбитальных объектов

Число орбитальных объектов в околоземном космосе возрастает высокими темпами. Рост множества орбитальных объектов обусловлен, во-первых, увеличением числа объектов космического мусора (неуправляемых объектов). Во-вторых, оно растет вследствие увеличения числа космических аппаратов. Насыщение околоземного космоса движущимися объектами повышает вероятность возникновения механических конфликтов – орбитальных столкновений. Глобальная задача обеспечения безопасного использования околоземного космического пространства в практических целях включает в себя задачи контроля и анализа текущего состояния механической конфликтности множества орбитальных объектов в околоземном космосе и прогноза его эволюции с учетом тенденций развития спутниковых систем различного назначения. В частности, представляет интерес прогнозирование и анализ ситуаций повышенной опасности механических конфликтов (конфликтных ситуаций) для пар орбитальных объектов. Под конфликтной ситуацией или опасным сближением будем понимать сближение пары объектов на расчетных траекториях на расстояние, опасное с точки зрения возникновения между ними столкновений при отклонении от этих траекторий. Вследствие многочисленности множества орбитальных объектов в околоземном пространстве возрастает не только вероятность опасных (конфликтных) сближений для пар орбитальных объектов. Растет вероятность поликонфликтных сближений (или поликонфликтных ситуаций) – одновременных сближений на опасное расстояние нескольких (более двух) орбитальных объектов. В данной работе представлен комплекс моделей и методов, предназначенных для прогноза и анализа поликонфликтных ситуаций в околоземном космосе.

Исследование задач анализа движения орбитальных объектов друг относительно друга (в том числе – задачи исследования поликонфликтных ситуаций) развиваются в двух направлениях. Первое базируется на моделировании движения орбитальных объектов и анализе текущих расстояний между ними (такой подход может обеспечить высокую точность, но требует существенных затрат времени). Второе направление предполагает моделирование относительно медленной эволюции траекторий, анализ их текущей геометрии и выявление на его основе участков траекторий, отвечающих заданным условиям расположения (в частности – близких участков траекторий двух или нескольких объектов). После определения опасных участков траекторий используются специальные быстрые методы, обеспечивающие с учетом периодичности движения определение интервалов времени одновременного движения объектов по этим участкам. В данной работе предложено развитие методов второго направления для задач прогноза и анализа поликонфликтных ситуаций в совокупном движении многоэлементного множества орбитальных объектов. Пару близко расположенных (опасных) участков траекторий орбитальных объектов будем называть узлом конфликтов. Если участок траектории объекта А, входящий в узел конфликтов, образованный с траекторией объекта В, входит также в узел конфликтов, образованный объектом А с траекторией объекта С, то будем говорить, что участок траектории объект А входит в поликонфликтный узел слабой связности с объектами В и С (у объекта А есть потенциальная возможность одновременного сближения с объектами В и С). Если в составе опасных участков траекторий объектов В и С, входящих в узел поликонфликтов с объектом А, есть участки находящиеся на опасном расстоянии друг от друга, то будем говорить, что траектория объекта А входит в узел полной связности с объектами В и С (возможно одновременное сближение всех трех объектов). Иными словами,

в первом случае можно говорить о слабой связности узлов конфликтности, образованных парами траекторий А, Б и А, С, а во втором – о полной их связности. В общем случае узел поликонфликтов может быть образован несколькими узлами слабой или полной связности.

Методы выявления потенциальной возможности возникновения поликонфликтных сближений орбитальных объектов (выявления узлов поликонфликтов во множестве траекторий) чаще всего базируется на рассмотрении объемов разбиения пространства и проверке наличия в каждом из этих объемов двух и более траекторий. В данной работе представлен другой вариант решения, который базируется на быстрых методах определения близких участков пар траекторий (узлов конфликтов) и быстрых методах определения вхождения узлов конфликтов в узлы поликонфликтов. Под быстрым методом понимается такой, который позволяет получать решение с затратами времени в несколько порядков меньшими, чем на основе моделирования движения орбитальных объектов (методов первого направления), и в несколько раз быстрее по сравнению с методами второго направления, базирующимися на рассмотрении объемов разбиения. Метод выявления узла конфликтов предполагает определение опасных участков траекторий и относительно грубых оценок возможного сближения объектов на этих участках. В том числе, определение узла конфликтов предполагает нахождение значения истинных аномалий, ограничивающих опасные участки траекторий, входящие в узел конфликтов. Метод учитывает эволюцию орбит, определяющую эволюцию узла конфликтов (в том числе, изменение положение опасного участка вдоль орбиты, которой он принадлежит, и его размеров). Предложен метод определения состава узлов поликонфликтов, который базируется на составлении матрицы связности. Для этого каждому участку траектории, входящему в какой-либо узел конфликтов пары траекторий, присваивается сквозной номер. Двумерная матрица кодировки опасных участков, обеспечивающая идентификацию каждого участка и хранение основной информации о нем, имеет структуру, описанную далее. Номер столбца матрицы соответствует сквозному порядковому номеру опасного участка. Элементы столбца содержат следующую информацию: номер объекта, которому принадлежит участок; номер объекта, с опасным участком траектории которого этот участок входит в узел конфликтов; номер узла конфликтов, образованного этими траекториями (две траектории могут образовывать от одного до четырех узлов конфликтов); значения истинных аномалий, определяющие точки начала и конца участка. В двумерной матрице связности элемент содержит информацию о наличии связи между участками, номер одного из которых соответствует номеру строки, а номер другого – номеру столбца. Разработан алгоритм, который позволяет выявлять множество опасных участков, объединенных слабой связностью в узел поликонфликтов. Кроме того, предложены алгоритмы, которые позволяют выявлять в узле поликонфликтов узлы слабой связности заданной кратности (каждый участок такого узла имеет слабую связь с участками, число которых равно значению кратности) и узлы поликонфликтов полной связности, в том числе, – полной связности заданной кратности. Предложен подход к описанию области пространства, охватывающей узел поликонфликтов (назовем ее зоной поликонфликта). С течением времени состав узла поликонфликтов и зона его охвата меняются. Предложена модель, соответствие которой в течение некоторого времени моделирования может быть использовано как критерий проверки вхождения участков траекторий в текущий состав узла поликонфликтов (модель «захвата» узла). Проверка выполняется путем реализации нескольких коротких операций (без выполнения описанных выше действий по выявлению узла поликонфликтов на каждом шаге моделирования). Для отладки предложенных методов разработано несколько модификаций быстрого метода генерации узла поликонфликтов. Разработан метод быстрого определения одновременного прохождения узла поликонфликтов несколькими орбитальными объектами.

Мілявський Ю.Л.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Порівняння методів класичного та нечіткого керування в імпульсних процесах когнітивних карт

Когнітивна карта (КК) – це засіб моделювання складної системи, що використовує математичний апарат зв'язаних орієнтованих графів для відображення взаємозв'язків між основними поняттями (концептами), що входять до складної системи [1]. Серед великого різноманіття типів КК можна умовно виділити дві групи – КК “у дусі Робертса” (“чіткі” КК [2]) та КК “у дусі Коско” (нечіткі КК [3]). У даній доповіді увагу приділено проблемі управління імпульсним процесом у КК. Імпульсний процес – це динамічний перехідний процес у складній системі, ініційований внутрішніми або зовнішніми збуреннями, що описується спеціальними системами рівнянь відносно координат вершин КК залежно від типу КК. У випадку, коли особа, що приймає рішення, має можливість впливати на деякі вершини КК, існує мета функціонування складної системи, описуваної КК, що полягає у досягненні певного стану цієї системи, виникає задача управління її імпульсним процесом. У попередніх роботах автора, наприклад, [4–6], запропоновано ряд методів управління імпульсними процесами КК у дусі Робертса, що базуються на адаптації методів теорії автоматичного управління до поставленої задачі. Однак, у англомовній літературі більш популярними є нечіткі КК, для яких природним є нечітке управління, виражене в термінах лінгвістичних термів, функцій належності, нечітких правил тощо. Відповідно, виникає потреба в синтезі нечітких регуляторів для управління КК, які в свою чергу можуть бути, взагалі кажучи, як нечіткими, так і чіткими. Інструментарій нечітких регуляторів для інженерних систем сам по собі є добре відомим і непогано розробленим [7]. До нечітких КК цей підхід було застосовано в декількох дослідженнях, наприклад, [8], але, на нашу думку, це не набуло достатнього розвитку. Також, наскільки нам відомо, до чітких КК керування за допомогою нечітких регуляторів не застосовувалось взагалі.

У даній доповіді нечіткий регулятор буде застосовано до управління імпульсним процесом КК. Результати буде порівняно з результатами управління тою ж КК за допомогою раніше розроблених методів “чіткого” автоматичного керування. Таким чином, на прикладі буде показано, в чому переваги та недоліки обох груп методів управління складними системами, представленими КК.

Література. 1. Инновационное развитие социально-экономических систем на основе методологий предвидения и когнитивного моделирования / Под ред. Г.В. Гореловой, Н.Д. Панкратовой. – К.: Наукова думка, 2015. – 464 с. 2. Roberts F. Discrete Mathematical Models with Applications to Social, Biological, and Environmental Problems. – Englewood Cliffs, Prentice-Hall, 1976. – 559 p. 3. Kosko B. Fuzzy Cognitive Maps // International Journal of Man-Machine Studies. – 1986. – 24. – P. 65 – 75. 4. М.З. Згуровский, В.Д. Романенко, Ю.Л. Мильавский. Принципы и методы управления импульсными процессами в когнитивных картах сложных систем. Часть 1 // Проблемы управления и информатики. – 2016. – № 2. – С. 21–29. 5. Romanenko V., Milyavsky Y. Systematization of Methods of Automated Control in Cognitive Maps' Impulse Processes for Complex Systems // 2017 IEEE First Ukraine Conference on Electrical and Computer Engineering (UKRCON), Kyiv, Ukraine, May 29 – June 2, 2017. – P. 776 – 782. 6. Romanenko V., Milyavsky Y. Combined Control of Impulse Processes in Complex Systems' Cognitive Maps with Multirate Sampling // The 9th IEEE International Conference on Intelligent Data Acquisition and Advanced Computing Systems: Technology and Applications, Bucharest, Romania, September 21–23, 2017. – P. 8 – 13. 7. Гостев В.И. Нечеткие регуляторы в системах автоматического управления. – К.: Радиоаматор, 2008. – 972 с. 8. C. Stylios, P. Groumpos. Fuzzy Cognitive Maps in modeling supervisory control systems // Journal of Intelligent and Fuzzy Systems. – Vol. 8, no. 1. – 2000. – P. 83–98.

Назарага І.М.

Київський національний університет ім. Тараса Шевченка, Київ, Україна

Матрична множинна регресія та класичні методи біометрії для прогнозування біологічних показників: приклади

Сучасна біологія широко використовує математичний апарат для моделювання та прогнозування процесів у живих системах і формалізації механізмів, що лежать в їх основі. З метою передбачення поведінки таких систем дедалі частіше вдаються до віртуальних експериментів із використанням математичних засобів, що дає змогу контролювати, вимірювати чи прогнозувати ключові змінні. У зв'язку з цим полегшується інтерпретація отриманих результатів і з'являється можливість проводити експерименти без участі живих істот, які в житті неприпустимі з етичних міркувань [1]. Таким чином, розробка, реалізація та апробація різних математичних методів прогнозування біологічних показників є актуальними завданнями для сучасної біологічної науки.

У доповіді розглядаються приклади прогнозування біологічних (вітальних) показників із використанням класичних методів біометрії та із застосуванням алгоритму на основі матричної множинної регресії. Емпіричними даними для розрахунків стали дані експерименту, проведеного в навчально-науковому центрі “Інститут біології та медицини” Київського національного університету імені Тараса Шевченка [2]. Розрахунки проведені у Microsoft Office Excel та у середовищі Wolfram Mathematica.

Зокрема, для розв'язання задач прогнозування вітальних (життєвих) показників біометричними методами використано вбудовані засоби Microsoft Office Excel: лінії тренда (лінійна, експоненційна, логарифмічна апроксимації). Згідно із обчисленими значеннями похибок прогнозу показників за критерієм APE (absolute percentage error) для методів біометрії, що базуються на лінійній та логарифмічній апроксимації, точність прогнозу є задовільною (похибка становить від 0% до 50%), для методу на основі експоненційної апроксимації – незадовільною (похибки для деяких показників більші 50%).

З метою розв'язання задачі оцінювання методом найменших квадратів для множинної матричної регресії використано розроблений у [3, 4] математичний апарат сингулярного (SVD) розкладу і техніку псевдообернення за Муром-Пенроузом у рамках розвитку концепції коротких операторів. Запропонований у [4] алгоритм оцінювання для вектора невідомих параметрів класу матричних функцій реалізовано у середовищі Wolfram Mathematica. Точність прогнозу за критерієм APE (похибки від 0% до 10%) за цим алгоритмом, і, відповідно, якість методу прогнозування вказаних показників є високою.

Отже, проведені розрахунки показують, що результати прогнозування досліджуваних показників за алгоритмом, що базується на матричній множинній регресії є точнішими, ніж із використанням класичних методів біометрії. Таким чином, метод на основі матричної множинної регресії є конкурентоспроможним і може бути використаний для розв'язання задач прогнозування біологічних показників з прийнятною для цього точністю.

Література. 1. Сбальцарини И. Пространственно-временное моделирование в биологии. [Электронный ресурс] / Сбальцарини И. – 2012. Режим доступа: <https://biomolecula.ru/articles/prostranstvenno-vremennoe-modelirovanie-v-biologii> – Название с экрана. 2. Pasichnyk A. The effect of aqueous extract of Phaseolus Vulgaris pods on the some biochemical parameters in the conditions of esophagus burn of second degree in rats / Pasichnyk A., Dmytryk V., Rayetska Ya. // Youth and Progress of Biology: Book of Abstracts of XIII International Scientific Conference for Students and PhD Students. – Ukraine, Lviv, 25 – 27 April 2017. – P. 67 – 68. 3. Donchenko V. “Feature Vectors” in Grouping Information Problem in Applied Mathematics: Vectors and Matrixes / V. Donchenko, T. Zinko, F. Skotarenko // Problems of Computer Intellectualization: international conference, Institute of Cybernetics NASU, ITHEA. – Kyiv, Ukraine, – Sofia, Bulgaria. – 2012. – P. 111 – 124. 4. Донченко В.С. Матрична множинна регресія / В.С. Донченко, О.В. Тарасова // Вісник КНУ імені Тараса Шевченка. Серія: фіз.-мат. науки. – 2015. – № 2. – С. 133 – 138.

Нестеренко В.И.

КПИ им. Игоря Сикорского, ИПСА, Киев, Украина

Разработка стратегии развития группы демографических реестров в системе государственных реестров Украины

На данный момент в Украине существует более 150 реестров, которые принадлежат 40 различным государственным учреждениям, содержание каждого из которых обходится государству в 21 млн. грн. в год. Как показывает мировая практика, для слаженной работы реестров достаточно иметь 7 реестров. Кроме того наблюдается отсутствие единой стратегии развития, каждая институция проходит путь с начала, получая тот же опыт, что другие почерпнули уже давно, а также отсутствие единой базы классификаторов, что не позволяет унифицировать информацию.

Для перехода к более рациональной структуре системы реестров Украины был проведен анализ интероперабельности государственных информационных ресурсов, по результатам которого можно выявить основополагающие для системы реестры и строить связи между реестрами, используя информацию из них [1]. Данное исследование позволяет оценить техническую, семантическую, административную и организационную структуру реестров и выявить слабые места в системе информационных ресурсов, такие как дублирование одинаковой информации до 63-х раз.

Для исследования была выбрана часть реестров – демографические реестры, для которых строится стратегия развития. В дальнейшем тему можно будет масштабировать для покрытия всех реестров в стране.

Для определения рациональной структуры демографических реестров используется метод анализа иерархий, который позволяет выбрать наилучшую стратегию на глобальном уровне – выбор рационального количества реестров, их централизации и уровня защиты информации. Для определения компонент в реестрах, которые влияют на строение системы больше всего и которые будут основополагающими для системы, используется метод морфологического анализа. Преимуществами использования этого метода является возможность обнаружить ранее не видимые или не определенные пути решения [2, 3].

На заключительном этапе построения стратегии используются когнитивные карты, которые позволяют построить сценарий создания демографических реестров как в масштабах системы интероперабельности государственных информационных ресурсов, так и в масштабах доступности данных для физических и юридических лиц в Украине [4].

Результатом работы является получение стратегии построения единого реестра граждан с рациональным количеством полей и связей с реестрами из других областей, что позволит упростить ведение реестра для государственных учреждений, понизить стоимость его содержания и повысить уровень доступности информации для граждан Украины.

Литература. 1. <http://fra.europa.eu/en/publication/2017/fundamental-rights-interoperability>. 2. https://www.nursingcenter.com/journalarticle?Article_ID=3483692. 3. <http://www.dlib.org/dlib/june06/chan/06chan.html>. 4. <https://www.healthit.gov/sites/default/files/ONC10yearInteroperabilityConceptPaper.pdf>.

Панкратов В.А.

КПИ им. Игоря Сикорского, ИПСА, Киев, Украина

Оценивание промышленного сектора Украины на основе методологии когнитивного моделирования

Развитие мировой экономики в период до 2030 года будет отмечаться влиянием ряда факторов и мегатрендов, которые приведут к существенным изменениям общей картины мировой экономики и модификации форм ее организации. Учитывая место Украины на рубеже западноевропейской и восточноевропейской христианской цивилизации, формирование новой глобальной архитектуры ставит Украину перед жесткими геополитическими, геостратегическими и геоэкономическими вызовами, ответы на которые она должна найти именно в период до 2030 года.

В реальных социально-экономических проблемах, которые, как правило, являются слабо-структурированными и неформализованными, невозможен традиционный математический (экономический, социометрический и т.п.) подход к анализу процессов разработки комплексных решений. Для моделирования сложных неформализованных систем используется методология когнитивного моделирования, основанная на когнитивных аспектах. Эти аспекты включают в себя процессы восприятия, мышления, познания, объяснения и понимания. Схематическое, упрощенное описание картины мира, относящееся к проблемной ситуации, представляют в виде когнитивной карты. Технология когнитивного моделирования заключается в том, чтобы на основе когнитивных моделей определять возможные и оптимальные пути управления ситуацией с целью перехода от исходных состояний к желаемым. При когнитивном моделировании имеет место субъективность при предоставлении исходных данных по предметной области, рассматриваемой в процессе декомпозиции и агрегирования вершин графа, субъективность при введении коэффициентов весовых дуг графа. Преимуществом когнитивной модели является то, что она позволяет видеть как всю картину в целом, так и детали, интегрировать логику и фантазию, знания и опыт.

Привлекая основной принцип системного анализа – декомпозицию, с помощью которого сложная проблема сводится до формализованного уровня, выполняется процесс когнитивного моделирования, который реализуется в интерактивно-диалоговом режиме. Под когнитивным моделированием понимается решение взаимосвязанных проблем: построение когнитивной модели (карты), обоснование на каждом этапе моделирования устойчивости по значению и по возмущению, структурной устойчивости, учет многофакторных рисков, неопределенностей различной природы.

В результате проведенного сравнительного анализа систем когнитивного моделирования и системной динамики, надо отметить, что основными функциями таких систем являются: описание ситуации; определение целевых факторов, определение управляющих факторов; определение мер воздействия на ситуацию; определение функциональных связей для построения когнитивной модели в виде закономерностей в статистической информации по исследуемой ситуации, представления значений в виде нечеткого множества, качественных оценок (экспертная оценка), приписывание значений со шкалы силы связей, законов функционирования; определение тенденций развития ситуаций путем проведения имитационного моделирования; разработка стратегий и анализ их перспективности в контексте целей моделирования.

Когнитивное моделирование начинается с разработки когнитивной карты объекта. Когнитивная карта – структурная схема причинно-следственных связей в системе, которая интерпретирует суждения и взгляды ЛПР, строится для того, чтобы понять и проанализировать поведение сложной системы и построить сценарий ее возможного развития.

Модель сложной системы в виде иерархической когнитивной карты можно представить в виде,

$$IG = \langle G_k, G_{k+1}, E_k \rangle, k \geq 2, G_k \langle v_i(k), e_{ij} \{k\} \rangle,$$

где G_k – когнитивная карта k -уровня; $V(1) = \{v_i(1)\}$ – множество вершин нижнего (1-го) уровня, $V(k) = \{v_i(k)\}$ – множество вершин k -уровня. Отношения между вершинами одного уровня – дуги $e_{ij}(k)$, $E(k) = \{e_{ij}(k)\}$; отношения между вершинами разных уровней – $e_{k,k+1}$, $E_k = \{e_{k,k+1}\}$.

Когнитивная карта G_k , кроме графического изображения, может быть представлена квадратной матрицей отношений A_G , строки и столбцы которой соответствуют вершинам графа, а на пересечении i -строки, j -столбца стоят функции принадлежности, отражающие связь между вершинами V_i и V_j .

Построение и анализ когнитивных карт проводится последовательно:

I этап – когнитивный анализ сложной ситуации (погружение в проблему, изучение и идентификация проблемы).

II этап – построение когнитивной модели проблемной ситуации (выделение базовых факторов, описывающих суть проблемы, выделение в совокупности базовых факторов управляющих факторов, которые в модели будут потенциально возможными рычагами влияния на ситуацию, определение факторов-индикаторов, отражающих и объясняющих развитие процессов в проблемной ситуации и их влияние на различные сферы).

III этап – следующая последовательность шагов ее реализации:

1. Определение начальных условий, тенденций, характеризующих развитие ситуации на данном этапе (необходимо для адекватности модельного сценария реальной ситуации, усиливает доверие к результатам моделирования).
2. Задание целевых предпочтительных направлений (увеличение, уменьшение) и силы (слабо, сильно) изменения тенденций процессов в ситуации.
3. Выбор комплекса мероприятий (совокупности связанных факторов), определение их возможной и желаемой силы и направленности действий (мероприятий, факторов) на ситуацию, силу и направленность которых необходимо определить.
4. Выбор наблюдаемых факторов (индикаторов), характеризующих развитие ситуации, осуществляется в зависимости от целей анализа и желания пользователя.

В данной работе исследовался промышленный сектор Украины с целью выявления его приоритетных направлений. По результатам выполненного SWOT-анализа были выделены наиболее влиятельные факторы и установлены связи между ними. Полученные данные послужили исходными данными для когнитивного моделирования. С помощью методологии когнитивного анализа была построена модель когнитивной карты промышленного сектора и исследована ее устойчивость [1–3]. По построенному сценарию в виде устойчивой когнитивной карты выявлены следующие приоритетные отрасли промышленного сектора Украины: оборонная промышленность; машиностроение; композиционная промышленность; металлургия; химическая промышленность.

На основе реальных данных можно построить конкретную стратегию развития каждого из составляющих сектора промышленности Украины. В частности, построенный сценарий развития металлургической промышленности Украины показал целесообразную направленностью металлургии на композитные материалы в процессе производства. Преимуществами композитов является сокращение потерь от коррозии, химическая стойкость, износостойкость, тепло-, термо-, огнестойкость, увеличение ресурса техники, повышение надежности эксплуатации и высокие физико-механические характеристики.

Литература. 1. Zgurovsky M., Pankratov V.A. Strategy of innovative development of region on the basis of foresight methodologies and cognitive modeling synthesis // System research & Information Technologies, №2, – 2014. – P. 7–20. 2. Pankratov V. The creation of strategy for innovation development of socio-economic systems // International Journal. “Information technologies&knowledge”. ITHEA. SOFIA, V.3, №1.– 2014. – P. 84–99. 3. Pankratov V. To the estimation of cognitive maps sustainability // Intern. Conf “X Szkola Geomechaniki 20011”. Materiały Naukowe, Gliwice-Ustron, V.2, 18–21 października 2011. – 2011. – P. 119–128.

Панкратова Н.Д., Маняк Ю.В.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Нечітко-множинний підхід в задачах когнітивного моделювання складних систем

Використання конвенціональних методів і підходів для аналізу процесів в багатьох прикладних задачах суттєво ускладнюється або унеможлиблюється у випадку розгляду слабкоструктурованих принципово неформалізованих систем. Для моделювання та побудови сценаріїв такого класу складних систем застосовується методологія когнітивного моделювання, що спирається на пізнання, мислення, розуміння, сприйняття та інші аспекти когнітивної діяльності людини [1]. Одним із центральних понять когнітивної методології є концепція побудови когнітивної карти, яка описує модель досліджуваної ситуації.

Найбільш відомий тип нечітких когнітивних карт було запропоновано Коско в роботі [2], що отримали назву нечітких когнітивних карт (Fuzzy Cognitive Maps). Проте використання терміну “нечіткі” вказує тільки на те, що причинні зв'язки можуть набувати не тільки значення 0 або 1, а лежати в діапазоні дійсних чисел, що відображає ступінь впливу одного концепту на інший. Після публікації статті Коско було запропоновані численні модифікації нечітких когнітивних карт Коско. Проте значення концептів в переважній більшості даних моделей є чіткими числами, що обмежує їх можливості моделювання складних систем.

В даному дослідженні розробляється модифікація нечітких когнітивних карт із застосуванням апарату нечітких множин та рекомендації щодо їх використання при пошуку сценаріїв побудови складних систем.

Для побудови початкової структури когнітивної карти використовуємо наступну модифікацію методу SWOT-аналізу. Нехай D_i – коефіцієнт впливу внутрішніх сильних та слабких характеристик на реалізацію загроз, H_m – коефіцієнт впливу внутрішніх сильних та слабких характеристик на реалізацію можливостей, F_j – коефіцієнт впливу на загрози та можливості сильних характеристик, G_k – коефіцієнт впливу на загрози та можливості слабких характеристик: $D_i = \sum_j K_{S_j T_i} - \sum_k K_{W_k T_i}$, $H_m = \sum_j K_{S_j O_m} - \sum_k K_{W_k O_m}$, $G_k = \sum_i K_{W_k T_i} + \sum_m K_{W_k O_m}$, $F_j = \sum_i K_{S_j T_i} + \sum_m K_{S_j O_m}$.

Ступені впливу K представимо у вигляді нечітких множин. Об'єднання нечітких множин – сума відповідних показників:

$$\mu_{D_i}(x) = \max\{\max_j \{\mu_{K_{S_j T_i}}(x)\} - \max_k \{\mu_{K_{W_k T_i}}(x)\}; 0\}.$$

Застосувавши дефазифікацію для отримання числового значення отримано, наприклад, вираз для обчислення D_i :

$$D_i^* = \frac{\sum_t \chi^t \cdot \mu_{D_i}(\chi^t)}{\sum_t \mu_{D_i}(\chi^t)}.$$

Результати SWOT-аналізу як одного із методів якісного аналізу використовуються для визначення критичних факторів предметної області і є основою для першого етапу когнітивного моделювання. На основі отриманих значень D_i^* , H_m^* , G_k^* , F_j^* обирається множина концептів, а фазифіковані значення D_i , H_m , G_k , F_j слугують для ініціалізації початкової структури графа зв'язків нечіткої когнітивної карти.

Запропонуємо модифікований підхід до формального визначення когнітивної карти як нечіткого орієнтованого графа $G = \langle V, E \rangle$ [1]. Тут V – множина вершин; $V_i \in V, i = \overline{1..N}$, що представляють концепти (елементи системи). E – множина дуг $E_{ij} \in E; i, j = \overline{1..N}$, які відповідають зв'язкам між вершинами V_i та V_j .

Матриця відношень A_{ω_G} – це нечітка квадратна матриця, рядки й стовпчики якої помічені

вершинами графа G , а на перетині i -го рядка, j -го стовпчика знаходиться функція належності нечіткої величини, яка представляє уявлення експерта про взаємозв'язок між вершинами, якщо існує відношення між елементами V_i та V_j , тобто

$$A_{\omega_G} = [\omega_{ij}(G)],$$

$$\omega_{ij} = \omega_{ij}\langle v_i, v_j \rangle(f) : [-1; 1] \rightarrow [0, 1], i, j = \overrightarrow{1..n},$$

де $\omega_{ij} = \omega_{ij}\langle v_i, v_j \rangle(f) = \omega(v_i, v_j)$ – нечітке число, визначене на $[-1; 1]$. Якщо між вершинами V_i та V_j немає зв'язку, $\omega_{ij} \equiv 0$.

Формалізація імпульсних процесів для запропонованої моделі передбачає природне узагальнення моделі для класичних когнітивних карт:

$$v_i(t+1) = S(q(t+1), v_i(t), \bigwedge_{j=1, j \neq i}^N \{T(\omega_{ij}\langle v_i, v_j \rangle(\cdot), S(v_i(t+1), -v_i(t)))\}),$$

де S -норма – операція агрегування – додавання в термінах нечіткої арифметики (або операція максимуму), T -норма – операція множення в нечіткій арифметиці, $q(t)$ – зовнішній імпульс, $v_i(t)$ – значення концепта i .

При побудові графіків імпульсних процесів використовується центроїдний метод дефазифікації:

$$v_0\langle t \rangle = \frac{\int_{v^{-1}([0;1])} t \cdot v(x)\langle t \rangle}{\int_{v^{-1}([0;1])} v(x)\langle t \rangle}.$$

Зазначимо, що для досягнення структурної стійкості доцільно в першу чергу розглядати концепти, що потенційно є елементом найбільшої кількості циклів з урахуванням циклів, що утворюються після розбиття. Для цього використовуємо схему із застосуванням модифікації алгоритму Гірвана-Ньюмана [3]. При розбитті концептів на некорельовані підконцепти спочатку розглядаємо ребра, що є частиною найбільшої кількості найкоротших шляхів у графі (ступінь опосередкованості). В термінах запропонованої формалізації довжину шляху R обчислюємо як $S_{(i,j) \in R} \{\omega_{ij}\langle v_i, v_j \rangle(\cdot)\}$.

При побудові когнітивної карти використовуємо метод попарних порівнянь для формалізації експертних суджень у вигляді нечітких множин. Оцінки експертів представляються матрицею $M = [m_{ij}], i, j = \overrightarrow{1..n}$. Значення m_{ij} вказує на ступінь переваги $\mu_A(x_i)$ над $\mu_A(x_j)$ і визначається відносно стандартної шкали вербальних суджень [4]. Далі значення в дискретних точках $\mu_A(x_i), \mu_A(x_2) .. \mu_A(x_n)$ визначаємо як розв'язок задачі:

$$MF^T = v_{\max} F.$$

Тут вектор $F = [F_1, F_2, .., F_n], v_{\max}$ – максимальне власне число матриці M . Обґрунтування коректності: $M > 0$ за побудовою $\rightarrow \exists! v_{\max}$. Отже значення ФН в дискретних точках: $\mu_A(x_i) = \frac{F_i}{\sum_{i=1}^n F_i}$, де $\sum_{i=1}^n F_i = 1$.

В дослідженні використовується представлення лінгвістичної змінної, що відповідає кожному з концептів когнітивної карти, у вигляді набору термів, функція належності для кожного з яких будується шляхом апроксимації отриманих дискретних значень однією із стандартних кривих.

Література. 1. Горелова Г.В., Панкратова Н.Д. Инновационное развитие социально-экономических систем на основе методологий предвидения и когнитивного моделирования // Наукова думка. – 2015. – 464 с. 2. Kosko B. Fuzzy cognitive maps / B. Kosko // International Journal of Man-Machine Studies. – 1986. – №24. – P.65–75. 3. Girvan M. and Newman M.E.J., Community structure in social and biological networks, Proc. Natl. Acad. Sci. USA 99, 7821–7826 (2002) 4. Зайченко Ю.П. Нечеткие модели и методы в интеллектуальных системах. – К.: Изд. дом Слово. 2008. – 92 с.

Панкратова Н.Д., Сльота М.Р.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Гарантоване функціонування складних технічних систем з урахуванням ресурсу допустимого ризику

Стратегія супроводження гарантованого функціонування складних технічних систем (СТС) визначається послідовністю системно узгоджених процедур. Вони характеризуються неповнотою та невизначеністю інформації про ситуації, пов'язані зі своєчасністю прийняття рішень та невідворотного порога, що обмежує час переходу до надзвичайної або катастрофічної ситуації [1]. Для складних систем різної природи актуальною проблемою є надійне прогнозування та своєчасне запобігання ситуаціям, що виникають при нормальному режимі під впливом дестабілізуючих факторів на працездатність, безпеку та ефективність процесу функціонування, які виявляються за допомогою показників ступеня ризику, рівня ризику та ресурсу допустимого ризику. Останній показник характеризує джерело, визначене станом працездатності, безпеки та живучості протягом певного часу відповідно до реальних умов роботи. Метою є розробка методики системного оцінювання експлуатації СТС, враховуючи ресурс допустимого ризику. Для досягнення цієї мети було вирішене розробити формалізацію оцінювання ресурсу допустимого ризику в реальному режимі часу функціонування СТС під впливом дестабілізуючих чинників та невизначеностей.

Математична постановка задачі. Дана вибірка значень реальних показників роботи СТС зовнішніх $y_i | i = 1, \bar{m}$ та внутрішніх $x_j | j = 1, \bar{n}$, параметрів, які було отримано протягом проміжку часу $D_0 = \{t_0 | t_0^- \leq t_0 \leq t_0^+\}$ $\langle Y(t_k) = \{Y_i(t_k), i = 1, m\}, X(t_k) = \{X_j(t_k), j = 1, n\}, t_k \in D_0 \rangle$. В процесі визначення критичних ситуацій маємо проміжки часу $T_R = \{[t_{fr}^i, t_{rr}^i] | t_{fr}^i < t_{rr}^i, t_{rr}^i < t_{cr}, t_{fr}^i, t_{rr}^i \in D_0 \cup D_p, r \in [1, n_r], i = 1, m\}$, де t_{fr}^i, t_{rr}^i час початку та завершення формування рішення [2]. В кожному інтервалі $[t_{fr}^i, t_{rr}^i]$ в реальному режимі часу вимірюються показники інформованості. А саме $I_V^i(t_k)$ – показник достовірності, $I_C^i(t_k)$ – показник повноти та $I_T^i(t_k)$ – показник своєчасності інформованості ОПР. Беремо до уваги, що $I_0^i(t_k) = I_V^i(t_k) \cdot I_C^i(t_k) \cdot I_T^i(t_k)$ – інтегрований показник інформованості ОПР [2]. Для поточного значення y_i обчислюється раціональне значення інформованості за інтегрованим показником інформованості, використовуючи формулу $I_0^- = \min_{i \in [1, m]} \{I_0^i(t_k)\}, t_k \in [t_{fr}^i, t_{rr}^i], k = 1, K; t_{rr}^i < t_{cr}\}$, де $I_0^i(t_k) = \frac{1}{K} \sum_{t_k \in [t_{fr}^i, t_{rr}^i]} I_0^i(t_k), i = 1, m$ – середні значення інтегрованого показника інформованості в проміжку $[t_{fr}^i, t_{rr}^i] \in T_R, i \in [1, m]$. Часові проміжки $[t_{s1r}^i, t_{s2r}^i] \subseteq [t_{fr}^i, t_{rr}^i] \subset T_R, i \in [1, m]$ для формування рішення та пошуку додаткової інформації, поки $I_0^i(t_k) \geq I_0^i(t_k)$, використовуються для визначення ресурсу допустимого ризику, щоб уникнути виникнення катастрофічної ситуації. Ресурс допустимого ризику визначений на проміжку часу $[t_{a1r}^i, t_{s2r}^i], t_{s1r}^i \leq t_{a1r}^i < t_{s2r}^i, i \in [1, m]$. А ресурсом допустимого ризику будемо вважати відношення довжини визначеного проміжку до довжини проміжку $[t_{s1r}^i, t_{s2r}^i]$.

Висновки. Таким чином в запропонованій моделі для кожного експлуатаційного індикатора СТС ресурс прийняттого ризику оцінюється часовим інтервалом. Цей час ОПР може використовувати для запобігання виникнення катастрофічних ситуацій. Запропонована методика призначена для оцінки функціонування реальних СТС з урахуванням ресурсу допустимого ризику. Відмінною рисою цього підходу є застосування принципу своєчасного виявлення і усунення причин можливого переходу СТС в неробочий стан, враховуючи показники інформованості щодо ситуацій і непереборне порогове обмеження часу в умовах впливу дестабілізуючих факторів, що забезпечує гарантоване функціонування СТС.

Література. 1. Pankratova N.D. The integrated system of safety and survivability complex technical objects operation in conditions of uncertainty and multifactor risks// Proceedings of conference IEEE (№50). May 29–June 2, 2017, Kyiv, Ukraine. – P. 1135–1140. 2. Pankratova N.D., Kondratova L.P. System evaluation of engineering objects' operating taking into account the margin of permissible risk// Eastern-European Journal of Enterprise Technologies № 3. – 2016. – P. 13–19.

Романкевич О.М., Манілевич Д.Ф.

КПІ ім. Ігоря Сікорського, ФПМ, Київ, Україна

Генератор двійкових векторів постійної ваги для виконання статистичних експериментів

При розробці багатопроцесорних систем важливим аспектом є розрахунок надійності системи. Одним з способів розрахунку є використання моделі поведінки системи у потоці відмов, що потребує подання на неї послідовності тестових двійкових векторів, зокрема векторів з постійною вагою [1], тобто з постійною кількістю одиничних значень у векторі, які формуються генераторами псевдовипадкових чисел. Ці вектори є векторами станів системи $Z = z_1, z_2, \dots, z_n$, де $z_i \in \{0, 1\}$ – стан окремого елемента i , де 0 означає, що елемент справний, 1 – елемент вийшов з ладу. Це дозволяє проводити статистичні тести, щоб побачити, у яких випадках система вийде з ладу цілком, а при яких продовжить працювати.

У доповіді розглядається генератор тестових векторів для моделей неоднорідних багатопроцесорних систем. Неоднорідність системи полягає у тому, що її компоненти мають різну ймовірність виходу з ладу. Ідея побудови генератора полягає в формуванні векторів, у яких ймовірності появи одиничного значення розрядів відповідає впорядкованим по спаданню ймовірностям виходу з ладу елементів. Це дає можливість зменшити час проведення статистичних тестів над моделями без зменшення точності розрахунку надійності. Тобто можна швидше отримати вектори, у яких більшість одиниць знаходяться на позиціях, що відповідають елементам з більшою ймовірністю вийти з ладу, що є більш бажаною ситуацією при проведенні таких тестів.

Генератор будується на основі регістру зсуву з лінійним зворотнім зв'язком, у якому знаходиться вихідна послідовність [2]. На вихідний регістр подаються певні функції збудження $f_i \in \{0, 1\}, i = 1, n$ де n – розрядність регістра, що вказують, які розряди регістру виконують зсув у даному такті роботи (1 – розряд приймає участь у зсуві, 0 – ні). Це дозволяє формувати вектори постійної ваги, що має перевагу над векторами зі змінною вагою, адже у більшості випадків модель поведінки багатопроцесорної тестують на конкретну кількість відмов елементів. Подаючи функції f_i , які приймають значення 1 з ймовірністю тим більше, чим менше значення індексу, дозволить в подальшому збільшити ймовірність появи одиничного значення на позиціях вихідного вектору з відповідними індексами. Також, оскільки $p(f_i), i = 1, n$ – ймовірності, з якою функції f_i приймають одиничне значення для відповідних розрядів, розраховуються на основі ймовірності виходу з ладу елементів системи, вони мають бути нормалізовані, адже значення ймовірностей виходу з ладу елементів надзвичайно малі.

Проте, однієї зміни ймовірностей зсуву елементів недостатньо, щоб відповідним чином налаштувати ймовірності появи одиничного значення на відповідних позиціях вихідного вектора. Тому ще одним важливим елементом даного генератора є перестановка самого правого нульового значення та самого лівого одиничного значення у вихідному векторі. Вплив даної перестановки залежить від частоти її проведення. Не складно побачити, що таку перестановку не можна здійснювати в кожному такті, адже тоді одичне значення ніколи не з'явиться у самому правому розряді вихідного вектора.

У доповіді розглядаються задача формування векторів функцій збудження $F = f_1, f_2, \dots, f_n$, спосіб перестановки самого правого одиничного значення та самого лівого нульового значення у вихідному векторі та залежність роботи генератора від них.

Література. 1. Структурный метод генерации псевдослучайных последовательностей специального вида [Текст] / В.В. Гроль, В.А. Романкевич, Е.Р. Потапова, Мораведж Сейед Милад // Радиоэлектронні і комп'ютерні системи. – 2010. – № 5. – С. 230 – 236. 2. Romankevich A. Generation of Test Sequence's by Means of Shift Register Structures / A. Romankevich, V. Groll, L. Karachun // XI International Conference on Fault Tolerant Systems and Diagnostics. – Berlin. – 1988. – Vol. 2. – P. 129–134.

Слюсар А.В., Данилов В.Я.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Система прийняття рішень на основі вейвлетної ідентифікації хвиль Елліотта

В наш час суттєві економічні зміни відбуваються чи не раз на декілька років і під них буває непросто підлаштуватися. В таких умовах дуже важливо навчитися їх передбачати. Аналіз часових рядів валютного котирування зараз набирає актуальності, оскільки відкриває можливості для заробітку на основі купівлі-продажу валюти та наукової діяльності в економічній сфері. Наразі для аналізу валютного котирування використовується близько двохсот індикаторів і їх кількість постійно збільшується. Така тенденція свідчить про те, що в цій області постійно ведуться дослідження і вона ще повністю не вивчена. Тому кожний новий підхід до аналізу валютного котирування покращує розуміння прихованих процесів, які впливають на курс валютних пар та економіку в цілому.

Хвилі Елліотта. Однією з найбільш ґрунтовних праць в області аналізу економічних тенденцій є так званий хвильовий аналіз Елліота. Хоча сам аналіз був придуманий ще в 30-х роках XX століття, він досі залишається популярним напрямком досліджень. Особливості хвиль Елліотта [1]:

1. фондовий ринок підпорядковується повторюваному ритму 5-3 хвиль зростання(імпульсні), три хвилі падіння(корекційні);
2. залежність між ланками хвиль пов'язана з рівнями Фібоначчі;
3. хвилі Елліотта мають фрактальну структуру;
4. основні моделі 5-3 лишаються постійними, хоча проміжок їх тривалості може змінюватися.

Вейвлет-перетворення. Оскільки ми маємо справу з хвильовим аналізом, тоді логічним кроком було б застосувати математичні інструменти, які працюють з хвилями. Аналіз хвиль виконувався за допомогою вейвлет-перетворень.

Вейвлет-перетворення схоже на перетворення Фур'є. Основна відмінність лежить в наступному: вейвлет-перетворення використовує функції, локалізовані як в реальному, так і в Фур'є-просторі. Усі вейвлет-перетворення розглядають функцію (взяту як функцією від часу) у термінах коливаль, локалізованих за часом (простором) і частотою [2,3]. Для даного дослідження використовувались вейвлети Рікера, Пауля, Морлет, DOG.

Система прийняття рішень. Для того щоб дана теорія була цікава для бізнесу та наукових досліджень процес виявлення хвиль Елліотта за допомогою вейвлет-перетворень потрібно автоматизувати. На цій основі була побудована система прийняття рішень для бірж валютних котирувань.

Система підтримки прийняття рішень — комп'ютеризована система, яка шляхом збору та аналізу великої кількості інформації може впливати на процес прийняття управлінських рішень в бізнесі та підприємстві.

Для проектування цієї системи проаналізована робота трейдерів з торгівлі валютою на біржі [4] та аналогічними системами прийняття рішень наприклад СПІР від Swissquote(онлайн Швейцарський банк). Трейдерам важливо мати перед очима візуалізацію підрахунків, тобто всі індикатори повинні бути відображені в доступному форматі.

В результаті дослідження виявлено такі важливі характеристики системи:

1. можливість швидко переключатися між парами валютних котирувань;

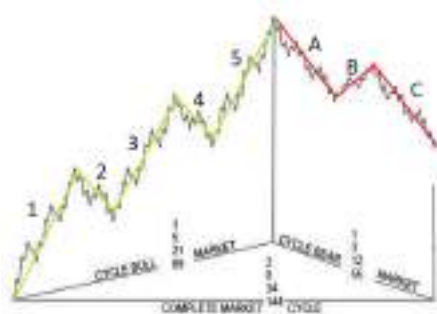


Рис. 1. Приклад повного ринкового циклу хвиль Елліотта. Зеленим кольором (1-5) виділена імпульсна ланка хвилі, червоним (А-С) – корекційна.

2. інструменти для відображення часових рядів повинні мати можливість змінювати масштаб;
3. можливість переглядати технічні індикатори та порівнювати з часовими рядами валютних котирувань;
4. рекомендації щодо купівлі/продажу валюти, яка досліджується.

Саме по цих характеристикам була побудована система прийняття рішень.

Прогнозування на основі аналізу вейвлет-перетворень. Після дослідження вейвлет-перетворень часових рядів валютних котирувань з хвилями Елліотта відкрито, що вейвлет Пауля краще за інших (Морлет, Рікера, DOG) та описує поведінку часових рядів на граничних значеннях. Така поведінка вейвлет-перетворення Пауля аргументована тим, що він чутливий до імпульсної та корекційної хвилі Елліотта. На основі цих досліджень побудовано алгоритм аналізу індикатора вейвлета Пауля, що дозволяє прогнозувати поведінку часових валютних котирувань.

Приклад прийняття рішення. Для того, щоб система прийняла рішення щодо фінансової операції на біржі потрібно ввести валютну пару, часові межі аналізу та частоту даних (кожну годину, день і т.д.). Далі система завантажує дані з відкритого API, виконує попередню обробку, рахує індикатор Херста та вейвлет-перетворення.



Рис. 2. Приклад роботи системи: графік валютних котирувань, ідентифікатор на основі вейвлет-перетворення Пауля та рекомендації щодо фінансової операції.

На рисунку 2 видно, що поведінка на кінцях інтервалів схожа на обох графіках: значення зменшується. За допомогою встановленої кореляції між цим індикатором та часовим рядом, система приймає рішення в сторону продажу цінних паперів. Оскільки тестування проводилось на історичних даних, результати були перевірені на вибірці для валідації. Тестування виявило, що надалі значення ціни валютної пари буде зменшуватися, тож рішення було прийнято вірно.

Література. 1. A.J. Frost Elliott Wave Principle: Key to Market Behavior. / A.J. Frost, R.R. Prechter — Delhi: Elliott Wave International, 2005. — 190 с. 2. Воробьев В.И. Теория и практика вейвлет-преобразования. / Воробьев В.И., Грибунин В.Г. — СПб.: ВУС, 1999. — 364 с. 3. Добеши И. Десять лекций по вейвлетам. / Добеши И. — Ижевск: РХД, 2001. — 464 с. 4. Мэрфи Джон Дж. Технический анализ фьючерсных рынков: теория и практика. / Мэрфи Джон Дж. — Москва: Диаграмма, 2000. — 281 с.

Тимофієва Н.К.

Міжнародний науково-навчальний центр інформаційних технологій та систем НАН та МОН України, Київ, Україна

Задачі комбінаторної оптимізації, цільова функція в яких уведена на нескінченній комбінаторній множині

Показано, що задачі комбінаторної оптимізації розв'язуються як на скінченній множині комбінаторного характеру так і нескінченній. Аргументом цільової функції в них є комбінаторні конфігурації як скінченні, так і нескінченні, як з повтореннями, так і без повторень

Вступ. Розв'язок в описаних в літературі задачах комбінаторної оптимізації, як правило, знаходиться на скінченній множині комбінаторного характеру. Аргументом цільової в них є комбінаторні конфігурації різних типів. Це – перестановки, сполучення та розміщення як з повтореннями, так і без повторень, розбиття натурального числа, графи тощо. Але існують задачі, в яких цільова функція задана на нескінченній множині комбінаторного характеру, а її аргументом є сполучення та розміщення з повтореннями, а також розбиття n -елементної множини на підмножини.

Основна частина. Під комбінаторною конфігурацією розуміємо будь-яку сукупність елементів, яка утворюється з усіх або з деяких елементів заданої базової множини $A = \{a_1, \dots, a_n\}$ [1]. Позначимо її впорядкованою множиною $w = \{w_1^k, \dots, w_\eta^k\}$, де $\eta \in \{1, \dots, n\}$ – кількість елементів у w^k , $W = \{w^k\}_1^q$ – множина комбінаторних конфігурацій. Верхній індекс k ($k \in \{1, \dots, q\}$) у w^k позначає порядковий номер w^k у W , q – кількість w^k у W .

Комбінаторні множини для фіксованого n є скінченними, а для довільного n вони – нескінченні. Базова множина також може бути нескінченною.

За способом обчислення цільової функції виділимо задачі, в яких для певного варіанту розв'язку її значення обчислюється одночасно. Такі задачі назовемо статичними. Задачі, в яких в процесі їхнього розв'язання генерується поточна інформація, за якою оцінюється результат, а пошук оптимального розв'язку проводиться поетапно з обчисленням часткової цільової функції, назовемо динамічними.

Оскільки в статичних задачах результат знаходиться для повністю побудованої комбінаторної конфігурації, то вона, як і її множина – скінченні. Цільова функція в них визначена на скінченній комбінаторній множині. В динамічних задачах аргумент цільової функції будується послідовно. В цьому разі базова множина, з елементів якої він утворюється, може бути нескінченною. Комбінаторна множина також – нескінченна. Відповідно, цільова функція визначена на нескінченній комбінаторній множині.

Розглянемо задачі комбінаторної оптимізації, аргументом цільової функції в яких є розбиття n -елементної множини $A = \{a_1, \dots, a_n\}$ на підмножини. До них відносяться задачі класифікації, кластеризації, покриття заданими об'єктами певної поверхні. Вони полягають в упорядкуванні заданих елементів у порівнянню однорідні групи, тобто за розробленими правилами проводиться розбиття елементів множини A на підмножини. При розв'язанні задачі кластеризації кількість кластерів та їхні характеристики невідомі, зате задано скінченну множину елементів, які необхідно розподілити по кластерах, відповідно відома їхня кількість. У класифікації характеристика класів, як правило, відома, але невідома кількість елементів, які підлягають класифікації. До того ж один і той же об'єкт може належати різним класам. Тому, в цій задачі аргументом цільової функції є розбиття нескінченної базової множини на підмножини з повтореннями.

Розглянемо задачу класифікації. Для неї виділимо такі підзадачі:

а) задано скінченну базову множину A . Класи можуть бути як задано так і не задано. Необхідно розподілити елементи базової множини по класах так, щоб останні не перетиналися.

Ця задача зводиться до задачі кластеризації;

б) задано скінченну базову множину A . Класи можуть бути як задано так і не задано. Елементи множини A розподіляються так, що один елемент може належати різним класам. В даному разі аргументом цільової функції є розбиття n -елементної множини A на η підмножин з повтореннями;

в) задано нескінченну базову множину, частина елементів якої відома, а частина визначається в процесі розв'язання задачі, тобто інформація поступає в процесі розв'язання задачі та змінюється в часі. Аргументом цільової функції в ній є часткове розбиття нескінченної множини $\tilde{A}$ на η підмножин з повтореннями. В цьому разі уводиться часткова цільова функція і часткове розбиття.

Оскільки для перших двох задач розбиття утворюється із елементів скінченної множини, яке характерне для задачі кластеризації, розглянемо аргумент цільової функції для третьої задачі. Розглянемо базову нескінченну множину $\tilde{A}$, в якій елементи $\tilde{a}_s$ для $s = \overline{1, n}$ задано, а для $s = \overline{n+1, \tilde{q}}$ визначаються в процесі розв'язання задачі. З відомих елементів $\tilde{a}_r \in \tilde{A}$, $r = \overline{1, \tilde{q}}$, утворюємо часткове розбиття множини $\tilde{A}$ на η підмножин (блоків) $w = (w_1, \dots, w_\eta)$, $\tilde{q} > n$ – кількість відомих елементів. Тоді множина підмножин $w = (w_1, \dots, w_\eta)$ має такі характеристики: $w_1 \cup \dots \cup w_\eta = A$, $w_p \cap w_l = \emptyset$ або $w_p \cap w_l \neq \emptyset$, $p \neq l$, $w_p \neq \emptyset$, $p, l \in \{1, \dots, \eta\}$. Непуста підмножина $w_p = \{\tilde{a}_1, \dots, \tilde{a}_{\xi_p}\}$ може мати від 1 до $\tilde{q}$ елементів ($\xi_p \in \{1, \dots, \tilde{q}\}$), $\eta \in \{1, \dots, \tilde{q}\}$, $\tilde{a}_r = \tilde{a}_s$ або $\tilde{a}_r \neq \tilde{a}_s$, $\tilde{a}_r, \tilde{a}_s \in w_p$, $r, s \in \{1, \dots, \xi_p\}$.

Як правило, при моделюванні задачі класифікації аргументом цільової функції вважають вхідні дані. Але в цій задачі оцінка результату проводиться за частковими цільовими функціями, аргументом якої є часткове розбиття нескінченної множини на підмножини з повтореннями.

В класифікації характеристика кластерів відома, об'єкти, які необхідно визначити, до якого вони класу відносяться, аналізуються не одночасно, а групами чи окремими елементами. Оскільки результат визначається не одночасно, а за частковою цільовою функцією, то задача класифікації відноситься до динамічних задач комбінаторної оптимізації.

Розглянемо задачу синтезу мовленнєвих сигналів. Вона полягає у їх відтворенні за заданим текстом [2]. Як правило, для розв'язання цієї задачі створюється бібліотека фрагментів, утворених із природних мовленнєвих сигналів, або такі фрагменти створюються штучно. Ця задача з використанням певних правил розв'язується об'єднанням майже періодів або ділянок сигналу, вибраних із бібліотеки, у фонему, що відповідають певним звукам (відповідно буквам заданого тексту). Аргументом цільової функції в цій задачі є розміщення з повтореннями [1].

Задача синтезу мовленнєвих сигналів полягає у знаходженні такого розміщення з повтореннями, для якого одержаний штучний мовленнєвий сигнал відповідав би природному його звучанню. Оскільки утворення штучних сигналів проводиться в динамічному режимі, а для порівняння його з природним сигналом уводяться часткові цільові функції (міри подібності), то ця задача відноситься до динамічних. Відповідно, комбінаторна множина, на якій проводиться пошук оптимального розв'язку – нескінченна.

Висновок. Отже, цільова функція в задачах комбінаторної оптимізації визначена як на скінченній так і нескінченній комбінаторних множинах. Ці задачі, як правило, відносяться до задач штучного інтелекту. Аналіз комбінаторних конфігурацій як аргументу цільової функції дозволяє виявляти характерні їхні властивості, адекватно будувати математичні моделі задач певного класу та розробляти ефективні алгоритми для їхнього розв'язання.

Література. 1. Тимофієва Н.К. Теоретико-числові методи розв'язання задач комбінаторної оптимізації. Автореф. Дис. д-ра. техн. наук: 01.05.02 / Ін-т кібернетики ім. В.М. Глушкова НАН України, – К., – 2007. – 32 с. 2. Винцюк Т.К. Анализ, распознавание и интерпретация речевых сигналов / Т.К. Винцюк. – К.: Наукова думка, 1987. – 262 с.

Фатенко В.В.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Узагальнення задачі Реньї про випадкове паркування

Задача про випадкове паркування (відома також як задача про випадкове заповнення відрізка) вперше була розглянута в роботі А. Реньї [1] в 1958 р. і вважається одним з класичних напрямків досліджень в стохастичній геометрії.

В роботі Реньї ця задача вивчалася в наступній постановці: розглядається одновимірний паркінг довжини $x \gg 1$, який можна уявляти собі як відрізок $[0, x]$. На цьому паркінгу послідовно випадковим чином паркуються автомобілі одиничної довжини. Перший з них розташовується на паркінгу згідно з рівномірним законом розподілу — це означає, що координата ξ_1 його лівого кінця є випадковою величиною, рівномірно розподіленою на відрізку $[0, x - 1]$: $\xi_1 \sim U(0, x - 1)$. Кожний з наступних автомобілів обирає для паркування будь-який (наприклад, перший) з вільних проміжків довжини не менше 1 між вже припаркованими машинами (або між одним з кінців паркінгу та найближчою до нього машиною) та паркується на ньому знову ж таки згідно з рівномірним законом розподілу — $\xi_n \sim U(\theta_-, \theta_+ - 1)$, де через θ_- та θ_+ позначено ліву та праву точку обраного вільного проміжку відповідно. Кількість автомобілів, що можуть припаркуватися на паркінгу фіксованої довжини, є скінченною — в деякий момент часу довжини всіх проміжків між припаркованими машинами стають меншими за 1, і процес паркування закінчується. Нехай $N(x)$ — кількість автомобілів, що вдалося припаркувати на паркінгу довжини x , а $m(x) = \mathbb{E}N(x)$ — її математичне сподівання. В роботі [1] було, зокрема, показано, що має місце співвідношення

$$\lim_{x \rightarrow \infty} \frac{m(x)}{x} = \int_0^\infty \exp \left(-2 \int_0^s \frac{1 - e^{-\tau}}{\tau} d\tau \right) ds \approx 0.7476. \quad (1)$$

В подальшому паркувальна модель Реньї та різні її узагальнення вивчалися багатьма авторами, серед них — А. Дворецкий та Г. Роббінс [2], П. Ней [3], Д. Менніон [4], Г. Соломон та Г. Вейнер [5], М. Пенроуз та Дж. Юкіч [6] та багато інших.

В роботі розглядається узагальнення моделі Реньї, згідно з яким автомобілі обирають місце паркування, керуючись розподілом, що є сумішшю рівномірного та виродженого розподілу в лівому кінці вільного проміжку. Нехай $p \in [0, 1)$. Тоді $\xi_n \sim U_p(\theta_-, \theta_+ - 1)$, де $U_p(a, b)$ має наступну функцію розподілу:

$$F_{p,a,b}(x) = p \cdot \begin{cases} 0, & x < a, \\ 1, & x \geq a, \end{cases} + (1 - p) \cdot \begin{cases} 0, & x < a, \\ \frac{x - a}{b - a}, & x \in [a, b], \\ 1, & x > b. \end{cases} \quad (2)$$

Теорема (Основний результат). В наведеній узагальненій моделі Реньї з законом, що заданий функцією розподілу (2), справедливе наступне асимптотичне співвідношення:

$$\lim_{x \rightarrow \infty} \frac{m_p(x)}{x} = \frac{1}{1 - p} \int_0^\infty \exp \left(-2 \int_0^s \frac{e^\tau - 1}{\tau(e^\tau - 1)} d\tau \right) ds. \quad (3)$$

Теорема допускає наступне узагальнення, дещо аналогічне тому, що було розглянуте в теоремі 2 роботи [7]. Нехай $(\tilde{f}_c, c > 0)$ — сім'я щільностей розподілу таких, що

1. $\tilde{f}_c(t) = 0$ при $t \notin [0, c]$,
2. $\tilde{f}_c(t) + \tilde{f}_c(c - t) = \frac{2}{c}$ при $t \in [0, c]$.

Тоді в моделі з законом, що заданий наступною функцією розподілу

$$F_{p,a,b}(x) = p \cdot \begin{cases} 0, & x < a, \\ 1, & x \geq a, \end{cases} + (1-p) \cdot \int_{-\infty}^{x-a} \tilde{f}_c(t) dt, \quad (4)$$

асимптотика $m_p(x)$ при $x \rightarrow \infty$ залишається тією ж.

Література. **1.** Rényi A. On a one-dimensional problem concerning random space-filling // Publ. Math. Inst. Hung. Acad. Sci. — 1958. — 3. — P. 109–127. **2.** Dvoretzky A., Robbins H. On the parking problem // Publ. Math. Inst. Hung. Acad. Sci. — 1964. — 9. — P. 209–224. **3.** Ney P.E. A random interval filling problem // Ann. Math. Stat. — 1962. — 33. — P. 702–718. **4.** Mannion D. Random space-filling in one dimension // Publ. Math. Inst. Hung. Acad. Sci. — 1964. — 9. — P. 143–154. **5.** Solomon H., Weiner H.J. A review of the packing problem // Comm. Statist. Th. Meth. — 1986. — 15. — P. 2571–2607. **6.** Penrose M.D., Yukich J.E. Limit theory for random sequential packing and deposition // Ann. Appl. Probab. — 2002. — 12, №1. — P. 272–301. doi: 10.1214/aoap/1015961164. **7.** Ананьевский С.М. Некоторые обобщения задачи о «парковке» // Вестник Санкт–Петербургского университета. Серия 1. Математика. Механика. Астрономия. — 2016. — 3(61), вып. 4. — С. 525–532. doi: 10.21638/11701/spbu01.2016.401.

Цегелик Г.Г., Краснюк Р.П.

Львівський національний університет імені Івана Франка, Львів, Україна

Моделювання оптимального розміщення баз даних інформаційних систем за наявності серверів проміжного зберігання даних

В роботі розглянуто задачі оптимізації розміщення баз даних при моделюванні розподілених інформаційних систем за наявності серверів проміжного зберігання даних. Зроблено математичну постановку відповідних задач оптимізації та запропоновано модифікацію алгоритму мурашиної колонії побудови наближеного розв'язку задач.

Проектування баз даних у розподілених інформаційних системах є доволі складною проблемою, що вимагає вирішення низки завдань, серед яких можна виділити такі:

- оптимізація розміщення баз даних за вузлами розподіленої інформаційної системи;
- синхронізація доступу до даних та паралельної обробки запитів для забезпечення необхідної продуктивності запитів;
- підтримка копій даних в декількох вузлах системи для зниження операцій пересилання даних при виконанні запитів.

У відповідності до наявних завдань проектування баз даних у розподілених інформаційних системах, з'явилися нові підходи, серед яких потрібно виділити реплікацію даних (*Data Replication – DR*) – асинхронний процес переносу змін об'єктів вихідної бази даних (*source database*) у бази даних, які належать різним вузлам розподіленої системи. Функції DR виконує спеціальний модуль СКБД – сервер тиражування даних, що носить назву реплікатора (*replicator*). Його завдання – підтримання ідентичності даних у приймаючих базах даних (*target databases*) до даних у вихідній БД. Однак, використання цієї технології вимагає розв'язання низки задач оптимізації, зокрема оптимального розміщення реплікаційних баз даних у розподіленій інформаційній системі за умов обмежень на обчислювальні ресурси вузлів, мінімізації часу синхронізації даних або мінімізації середнього часу, необхідного для пошуку інформації за умови відсутності синхронізації реплікаційних баз даних. В технологічному плані ці задачі можуть бути вирішені за наявності серверів проміжного зберігання даних (СПЗД), що відіграють роль центрів акумуляції даних за використання механізму пакетної передачі. Як наслідок, дослідження питань оптимального розподілу баз даних у інформаційних системах є актуальними і тому формулювання відповідних математичних моделей та розв'язок відповідних задач оптимізації і є предметом цієї роботи.

Формулювання задачі оптимального розміщення баз даних інформаційної системи зі серверами проміжного зберігання даних за умов мінімізації часу синхронізації. Нехай m баз даних необхідно розмістити у n вузлах мережі ($m < n$) за наявності l ($l < m$) серверів проміжного зберігання даних. Пропускна здатність каналу зв'язку за одиницю часу від вузла i до вузла j за використання серверу проміжного зберігання даних k складає ρ_{ijk} . Середній обсяг даних, який необхідно передати від вузла i до вузла j із використанням СПЗД k для синхронізації реплікаційних баз складає α_{ijk} . Тоді, якщо ввести у розгляд коефіцієнти x_{ijk} використання каналу зв'язку у задачі оптимізації, то математична модель набуває вигляду:

$$F = \sum_{i=1}^n \sum_{j=1}^n \sum_{k=1}^l \frac{\alpha_{ijk}}{\rho_{ijk}} x_{ijk} \rightarrow \min, \quad (1)$$

$$x_{ijk} = \{0, 1\}, \sum_{i=1}^n \sum_{j=1}^n \sum_{k=1}^l x_{ijk} \leq 2m, \sum_{i=1}^n \sum_{k=1}^l x_{ijk} \leq m - 1, i, j = 1, \dots, n; k = 1, \dots, l. \quad (2)$$

Умовою оптимальності в цьому випадку є мінімізація загального часу F синхронізації даних між реплікаційними базами даних. Очевидно, якщо для певного вузла i виконується умова

$$\sum_{j=1}^n \sum_{k=1}^l x_{ijk} = 0,$$

то в цьому вузлі база даних має бути відсутня. За математичного формулювання цієї оптимізаційної задачі накладається обмеження (2) щодо розміщення тільки однієї БД у вузлі мережі.

Формулювання задачі оптимального розміщення баз даних інформаційної системи зі серверами проміжного зберігання даних за умови мінімізації середнього часу, необхідного для пошуку інформації. Знову ж таки, m баз даних необхідно розподілити у n вузлах мережі ($m < n$) за наявності l ($l < m$) серверів проміжного зберігання даних, λ_{ijk} – інтенсивність запитів із вузла i до бази j за використання серверу проміжного зберігання даних k , β_{ijk} – величина даних, які мають бути переміщені у відповідності до запиту між вузлом i обчислювальної мережі та базою j , ρ_{ijk} – пропускна здатність каналу зв'язку за одиницю часу від вузла i до бази даних j із використанням СПЗД k . Тоді, якщо ввести коефіцієнти $x_{ijk} = \{0, 1\}$ що визначають, чи база даних j розміщена у вузлі i , то математична модель задачі у цьому випадку є такою:

$$F = \frac{\sum_{i=1}^n \sum_{j=1}^m \sum_{k=1}^l \frac{\lambda_{ijk} \beta_{ijk}}{\rho_{ijk}} (1 - x_{ijk})}{\sum_{i=1}^n \sum_{j=1}^m \sum_{k=1}^l \lambda_{ijk}} \rightarrow \min, \quad (3)$$

$$x_{ijk} = \{0, 1\}, \sum_{i=1}^n \sum_{k=1}^l x_{ijk} \geq 1; j = 1, \dots, m, \sum_{j=1}^m \sum_{k=1}^l (1 - x_{ijk}) = 1; i = 1, \dots, n. \quad (4)$$

Формулювання обчислювального алгоритму побудови розв'язку задач оптимізації. За великої розмірності розглянутих вище математичних моделей ефективним підходом до побудови розв'язку задач оптимізації є використання алгоритму мурашиної колонії [1, 2].

Не зменшуючи загальності, можемо стверджувати, що задачі розподілу баз даних (1), (2) та (3), (4) можуть бути приведені до наступної задачі оптимізації:

$$F = \sum_{i=1}^n \sum_{j=1}^m \sum_{k=1}^l \gamma_{ijk} x_{ijk} \rightarrow \min; i = 1, \dots, n, j = 1, \dots, m, k = 1, \dots, l, \quad (5)$$

та додаткових умов щодо шуканих параметрів задачі x_{ijk} , які не є критичними для ілюстрації загальної схеми алгоритму. У формулі (5) відомі коефіцієнти γ_{ijk} визначають «загальні витрати» щодо розміщення i -ї бази даних у j -му вузлі інформаційної системи.

За алгоритмом MMAS [2], модифікацію якого і розглядаємо у цій роботі, маємо три додаткові умови до загальної схеми мурашиного алгоритму:

- в оновленні феромонів бере участь тільки найкращий за поточну ітерацію мураха;
- значення рівня феромонів обмежено інтервалом $[L_{\min} > 0, L_{\max}]$;
- на початку алгоритму рівень феромонів для усіх вузлів встановлюється на рівні $L_{\max}$, що забезпечує краще дослідження простору розв'язків задачі на початковому етапі.

Наведемо кроки ітераційного алгоритму, коли на кожній ітерації колонія (множина) мурах генерує набір розв'язків у відповідності до рівнів феромонів. Серед отриманих результатів вибирається найкращий, що буде оновлювати ці рівні феромонів.

Крок 0. Вибирається величина похибки обчислень ϵ , значення максимальної кількості ітерацій $t_{\max}$, задається інтервал зміни рівня феромонів $[L_{\min} > 0, L_{\max}]$, значення параметрів «жадібності» алгоритму q та швидкості випаровування феромонів ρ . Для кожного вузла інформаційної системи покладається рівень феромону на рівні $L_{\max}$, встановлюється індекс ітерацій $t = 1$ та початкове значення цільової функції $F_0 = nl(\max_{j,j,k} \gamma_{ijk})$.

Крок 1. Ініціалізація мурашиної колонії одним із можливих способів:

- покриття вузлів – кількість мурах співпадає з кількістю вузлів n інформаційної системи, коли кожен мураха початково розміщується у відповідному вузлі;
- випадкове покриття – початково мурахи у вузлах інформаційної системи розміщуються випадковим чином, коли кількість мурах та вузлів системи можуть не співпадати;
- розміщення колонії у фокусі – уся колонія мурах на кожній ітерації знаходиться в одному вузлі інформаційної системи. Кількість мурах у колонії може бути довільною;
- мігруюча колонія – уся колонія мурах на кожній ітерації переміщується в довільний, випадково вибраний вузол системи. Кількість мурах у колонії також може бути довільною.

Крок 2. Розрахунок розв'язку задачі (5) для кожного мураха з колонії за використання ймовірностей

$$p_{s,t}(i, j, k) = \begin{cases} \frac{[L_t(i, j, k)]^q [\gamma_{ijk}]^{q-1}}{\sum_{p,r \in S_{i,s}} [L_t(i, p, r)]^q [\gamma_{ipr}]^{q-1}}; & j, k \in S_{i,s}, \\ 0; & j, k \notin S_{i,s}, \end{cases} \quad (6)$$

де у виразі (6): $L_t(i, j, k)$ – рівень феромонів, коли t – номер поточної ітерації, i, j та k визначають відповідно індекси бази даних, вузла інформаційної системи та номеру серверу проміжного зберігання даних; $S_{i,s}$ – множина ще не використаних вузлів інформаційної системи, де може бути розміщено базу з індексом i для s -ї мурахи колонії, q – параметр, що визначає «жадібність» алгоритму: за умови $q = 0$ вибір вершини визначатиметься найнижчими «загальними витратами» γ_{ijk} , а умова $q = 1$ визначає вибір вузла лише рівнем феромонів, що зумовлює швидке виродження результату до одного субоптимального розв'язку. Зауважимо, що формування множини $S_{i,s}$ на кожному кроці алгоритму здійснюється з використанням додаткових умов щодо шуканих параметрів задачі x_{ijk} відповідної оптимізаційної задачі.

Крок 3. Вибір найкращого розв'язку x_{ijk}^{best} на поточній ітерації з умови отримання найменшого значення цільової функції з (5): $F_t = F(x_{ijk,t}^{best}) = \sum_{i=1}^n \sum_{j=1}^m \sum_{k=1}^l \gamma_{ijk} x_{ijk,t}^{best}$.

Крок 4. Перевірка умов припинення ітерацій: модуль різниці знайдених значень цільової функції на двох останніх ітераційних кроках алгоритму не перевищує похибки обчислень $|F_t - F_{t-1}| < \epsilon$ або перевищено максимальну кількість ітерацій $t > t_{max}$. За виконання цих умов знайдені на поточному кроці значення $x_{ijk,t}^{best}$ формують розв'язок відповідної оптимізаційної задачі. Інакше переходимо до кроку 5.

Крок 5. Збільшуємо значення індексу ітерації t на одиницю, здійснюємо перерахунок рівня феромонів за формулою

$$L_{t+1}(i, j, k) = \begin{cases} L_{min}; & L^* \leq L_{min}, \\ L^*; & L_{min} \leq L^* \leq L_{max}, L^* = \rho L_t(i, j, k) + \Delta L_t(i, j, k), \\ L_{max}; & L^* \geq L_{max}, \end{cases}$$

де $\rho \in (0, 1)$ – швидкість випаровування феромонів, $\Delta L_t(i, j, k) = x_{ijk,t}^{best}$ – внесок найкращого за ітерацію мурахи у загальний рівень феромонів у вузлах інформаційної системи та повертаємося до кроку 1.

З практичної точки зору у задачах розміщення баз даних не є необхідним отримання глобального мінімуму – достатнім є отримання локального оптимального розв'язку, що забезпечується вибором коефіцієнту q близьким до одиниці та коефіцієнту $\rho > 0.5$.

Висновки. Здійснено математичну постановку задач розміщення баз даних у розподілених інформаційних системах за наявності серверів проміжного зберігання даних та сформульовано модифікацію алгоритму мурашиної колонії, що враховує специфіку розглянутих оптимізаційних моделей.

Література. 1. Dorigo M. Ant Colony System: A Cooperative Learning Approach to the Traveling Salesman Problem / M. Dorigo, L.M. Gambardella // IEEE Transactions on Evolutionary Computation, 1997. – Vol. 1, # 1. P. 53–66. 2. Stützle T. Max-Min Ant System / T. Stützle, H.H. Hoos // Future Generation Computer Systems, 2000. – Vol. 16, Issue 8. – P. 889–914.

Циганок В.В., Роїк П.Д.

Інститут проблем реєстрації інформації НАН України, Київ, Україна

Визначення узгодженості оцінок експертів при підтримці прийняття групових рішень

При підтримці прийняття рішень дуже важливо застосовувати групові експертизи, адже довіра до рекомендацій сформованих на основі знань колективу фахівців, беззаперечно, є значно вищою, ніж до сформованих однією, хоч і дуже кваліфікованою особою. У багатьох випадках оцінки експертів (їхні переваги, судження, тощо) можуть бути між собою недостатньо узгодженими для того, щоб у результаті їхнього подальшого узагальнення (агрегації) можливо було б отримати достовірні результати. Тому, дуже важливим аспектом є визначення ступеня узгодженості суджень експертів, а також визначення рівнів достатньої для агрегації узгодженості. На теперішній час відома досить значна кількість індексів узгодженості експертних оцінок, багато з них базуються на використанні статистичних показників, але такий підхід вбачається не досить доцільним, оскільки, зазвичай, множина оцінок не є репрезентативною вибіркою в статистичному сенсі (тобто, у випадках, що розглядаються, оцінок може бути лише декілька).

Сутність проблеми. У цьому дослідженні пропонується використати спектральний підхід, запропонований В.Г. Тоценком [1] та удосконалений В.В. Циганком [2]. Сутність цього підходу полягає у формуванні на основі експертних матриць парних порівнянь множини оцінок та поданні її у вигляді складових на обмеженій з обох кінців дискретній шкалі. Кожну цю складову можна поставити у відповідність оцінці, наданій деяким експертом при груповій експертизі. Набір таких складових зручно зображати у вигляді, так званого, спектру оцінок, приклад якого показано на рисунку 1.

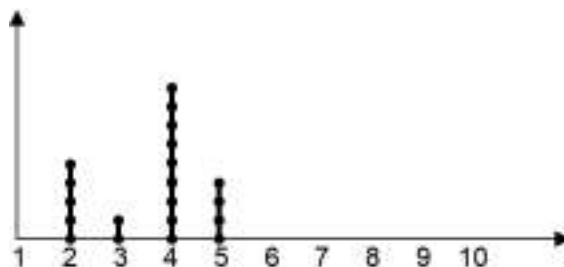


Рис. 1. Приклад зображення спектру, що відповідає набору оцінок різних експертів
 $\{2, 2, 2, 2, 3, 4, 4, 4, 4, 4, 4, 4, 5, 5, 5\}$

Застосування зазначеного підходу для визначення узгодженості шляхом подвійного використання формули ентропії було запропоновано, на ряду із іншими, у роботі [3], де окрім усього, було запропоновано коригування індексу, представленого в [1], з метою обмеження його області значень діапазоном $[0, 1]$, і неможливості попадання в область від'ємних значень.

Крім того, деякі практичні приклади свідчать про те, що функція індексу узгодженості, запропонованого в [1], інколи поводить себе не монотонно при зміні положення (значення) оцінок відносно середнього узагальненого значення.

Тому, виходячи із вище згаданого, задача розробки методу визначення узгодженості, позбавленого таких певних недоліків, вбачається актуальною. Актуальність ще зумовлюється тим, що експертні оцінки іноді без значної втрати точності важко подати на дискретній шкалі, і постає задача визначення узгодженості з використанням оцінок, як дійсних чисел на неперевних шкалах.

Постановка задачі дослідження. Сформулюємо вимоги до індексу узгодженості, які впливають зі зручності його використання та із розуміння здорового глузду щодо узгодженості.

Отже, щоб розробити метод визначення ступеня узгодженості, вирішено взяти до уваги

наступні базові постулати (аксіоми):

1. При експертному оцінюванні завжди існує деяка істинна оцінка, визначення якої є метою проведення групової експертизи.
2. Ця істинна оцінка відповідає деякому усередненому значенню множини індивідуальних оцінок експертів.
3. Множину індивідуальних експертних оцінок можливо подати на обмеженій з обох сторін чисельній неперервній або дискретній шкалі.
4. Максимальний рівень узгодженості (1.0) досягається тоді, і тільки тоді, коли всі експерти вибрали одну й ту саму оцінку.
5. Індекс узгодженості повинен бути незалежним від зсувів оцінок, тобто множина оцінок $\{1, 2, 5\}$ матиме той самий індекс, що й $\{4, 5, 8\}$.
6. При одній, і тій самій множині оцінок, збільшення шкали приводить до підвищення індексу узгодженості, і навпаки.
7. Якщо деяка множина оцінок більш узгоджена, ніж інша в певній шкалі, то вона є більш узгодженою і в будь-якій іншій шкалі.
8. При лінійних змінах (одночасному пропорційному збільшенні/зменшенні) параметрів шкали і значень усіх оцінок індекс узгодженості залишається незмінним. Тобто, індекс має мати властивість масштабованості.

Запропонований варіант вирішення задачі. Пропонується індекс узгодженості подати у вигляді наступного виразу:

$$I = \sum_{i \neq j} f(|x_i - x_j|),$$

де x_i – оцінка i -го експерта.

Отже, індекс узгодженості – це сума попарних порівнянь оцінок експертів. Запропонована функція f має наступну властивість: $f(x) > 0$, якщо $x > 0$ та $f(x) = 0$, якщо $x = 0$. Логічним і правильним було б цей індекс називати індексом *неузгодженості*, адже мінімум функції відповідає повній / найбільшій узгодженості оцінок.

Аналіз результатів. При викладених вище умовах: $I \geq 0$ і мінімум досягається тоді, і тільки тоді, коли всі змінні (оцінки) рівні. Розглядалися дві функції: $f(x) = |x|$ та $f(x) = x^2$. Для обох цих функцій максимум I досягається тоді, коли половина оцінок лежить на одному кінці шкали, інша – на іншому. У випадку, коли кількість оцінок непарна, то “зайва” оцінка може знаходитися на будь-якому кінці інтервалу. Визначення цього максимуму в певному сенсі йде у розріз з підходами тих авторів, які вважають, що максимум повинен досягатися, коли всі оцінки рівномірно розподілені по всій шкалі. Але, на практиці, не дуже важливо, де функція досягає максимуму, оскільки всі випадки високої неузгодженості, вище описані (та інші), не можуть використовуватись для подальшої агрегації оцінок експертів. Узгодженість у цих випадках необхідно “покрашувати”, важливо лише, щоб функція задовольняла базовим аксіомам. Вбачається можливим використання й інших функцій, окрім модуля та квадрата відстані між оцінками, але вони будуть занадто важкі для подальшого аналізу.

Висновки. Побудований індекс задовольняє всім сформульованим аксіомам, але не позбавлений випадків, коли він поводить себе не зовсім очікувано. Наступним етапом дослідження може бути подальший підбір функції f та/або видозміна індексу.

Література. 1. Totsenko V.G. The Agreement Degree of Estimations Set with Regard of Experts Competency // Proceedings of the Fourth International Symposium on the Analytic Hierarchy Process. Simon Fraser University. – Vancouver, Canada, 1996. – P. 229–241. 2. Циганок В.В. Елементи комбінаторного підходу при визначенні спектрального коефіцієнта узгодженості експертних парних порівнянь // Реєстрація, зберігання і обробка даних. 2012. – Т.14, №2. – С. 98–105. 3. Olenko Andriy & Tsyganok Vitaliy Double Entropy Inter-Rater Agreement Indices // Applied Psychological Measurement. 2016. – v. 40(1). – P. 37–55.

Чернуха О.Ю.^{1,2} Білушак Ю.І.^{1,2} Чучвара А.Є.¹

¹Центр математичного моделювання Інституту прикладних проблем механіки і математики ім. Я.С.Підстригача НАН України, Львів, Україна; ²Національний університет "Львівська політехніка", Львів, Україна

Програмний комплекс для моделювання дифузії у тілі з пастками за каскадного розпаду мігруючих частинок

Каскадні реакції достатньо поширені серед хімічних реакцій, де в якості частинок з невикористаними зв'язками виступають вільні атоми або радикали. Поява в середовищі необхідної частинки викликає каскад наступних реакцій, який продовжується до втрати частинки-носія реакції. Наприклад, розпад органічних азотовмісних речовин у ґрунтах в загальному вигляді подається наступною схемою: гумінові речовини, білки $\Rightarrow$ амінокислоти, аміді $\Rightarrow$ аміак $\Rightarrow$ нітрити $\Rightarrow$ нітрати $\Rightarrow$ молекулярний азот. Актуальність розгляду явища дифузії в середовищах з пастками обумовлена його застосуванням в таких галузях як фізика кристалів та напівпровідників, біофізика, наноматеріали, економіка, екологія тощо.

За континуально-термодинамічним підходом побудовано повну нелінійну систему рівнянь взаємозв'язаних теплових, механічних і дифузійних процесів за каскадного розпаду дифундуючих речовин у багатокомпонентному середовищі з пастками. Запропонований варіант лінеаризації рівнянь стану та кінетичних співвідношень і побудована відповідна ключова система рівнянь термомеханодифузії. На основі побудованої моделі сформульовано нові постановки крайових задач каскадного типу, коли концентрація частинок на певному кроці розпаду є джерелом маси розпадної речовини на наступному кроці, яка теж дифундує, сорбується, десорбується і розпадається:

для кроку розпаду $i = 0$

$$\frac{\partial c_1^{(0)}}{\partial t} = D^{(0)} \Delta c_1^{(0)} - k_1^{(0)} c_1^{(0)} + k_2^{(0)} c_2^{(0)} - \lambda^{(1)} c_1^{(0)} - \lambda^{(0N)} c_1^{(0)},$$

$$\frac{\partial c_2^{(0)}}{\partial t} = k_1^{(0)} c_1^{(0)} - k_2^{(0)} c_2^{(0)} - \lambda^{(1)} c_2^{(0)} - \lambda^{(0N)} c_2^{(0)};$$

для кроків розпаду $i = \overline{1, N-1}$

$$\frac{\partial c_1^{(i)}}{\partial t} = D^{(i)} \Delta c_1^{(i)} - k_1^{(i)} c_1^{(i)} + k_2^{(i)} c_2^{(i)} + \lambda^{(i-1)} c_1^{(i-1)} - \lambda^{(i+1)} c_1^{(i)} - \lambda^{(iN)} c_1^{(i)},$$

$$\frac{\partial c_2^{(i)}}{\partial t} = k_1^{(i)} c_1^{(i)} - k_2^{(i)} c_2^{(i)} + \lambda^{(i-1)} c_2^{(i-1)} - \lambda^{(i+1)} c_2^{(i)} - \lambda^{(iN)} c_2^{(i)};$$

для етапу $i = N$

$$\frac{\partial c_1^{(N)}}{\partial t} = D^{(N)} \Delta c_1^{(N)} - k_1^{(N)} c_1^{(N)} + k_2^{(N)} c_2^{(N)} + \sum_{i=0}^{N-1} \lambda^{(iN)} c_1^{(i)},$$

$$\frac{\partial c_2^{(N)}}{\partial t} = k_1^{(N)} c_1^{(N)} - k_2^{(N)} c_2^{(N)} + \sum_{i=0}^{N-1} \lambda^{(iN)} c_2^{(i)}.$$

Тут $c_j^{(i)}$ – концентрація домішкової речовини, яка утворилася на i -му етапі розпаду, у стані j ($i = \overline{0, N}$; $j = 1, 2$); $D^{(i)}$ – коефіцієнт дифузії частинок для етапу $i = \overline{0, N}$; $k_1^{(i)}$ і $k_2^{(i)}$ – коефіцієнти інтенсивності сорбції і десорбції, $\lambda^{(i+1)}$ і $\lambda^{(iN)}$ – коефіцієнти інтенсивності розпаду речовини на i -му кроці, що відповідають утворенню нової речовини на $(i+1)$ -му етапі та нерозпадної речовини для $i = N$.

Прийнято такі крайові умови:

$$c_1^{(i)}(\xi, \tau)|_{\tau=0} = c_2^{(i)}(\xi, \tau)|_{\tau=0} = 0, \quad i = \overline{0, N}; \quad c_1^{(0)}(\xi, \tau)|_{\xi=0} = c_0 \equiv const, \quad c_1^{(0)}(\xi, \tau)|_{\xi=\xi_0} = 0;$$

$$c_1^{(i)}(\xi, \tau)|_{\xi=0} = 0, \quad c_1^{(i)}(\xi, \tau)|_{\xi=\xi_0} = 0, \quad i = \overline{1, N}.$$

Розв'язки крайових задач каскадного типу побудовані за ітераційною процедурою з використанням функцій Гріна. Знайдено концентрації розпадних частинок, потоки маси мігруючих компонент і визначено кількість відповідних речовин, що за певний проміжок часу пройшли через нижню границю шару.

Розроблений програмний комплекс для комп'ютерного моделювання процесів дифузії у тилі з пастками за каскадного розпаду мігруючих частинок, архітектуру якого подано на рис. 1.

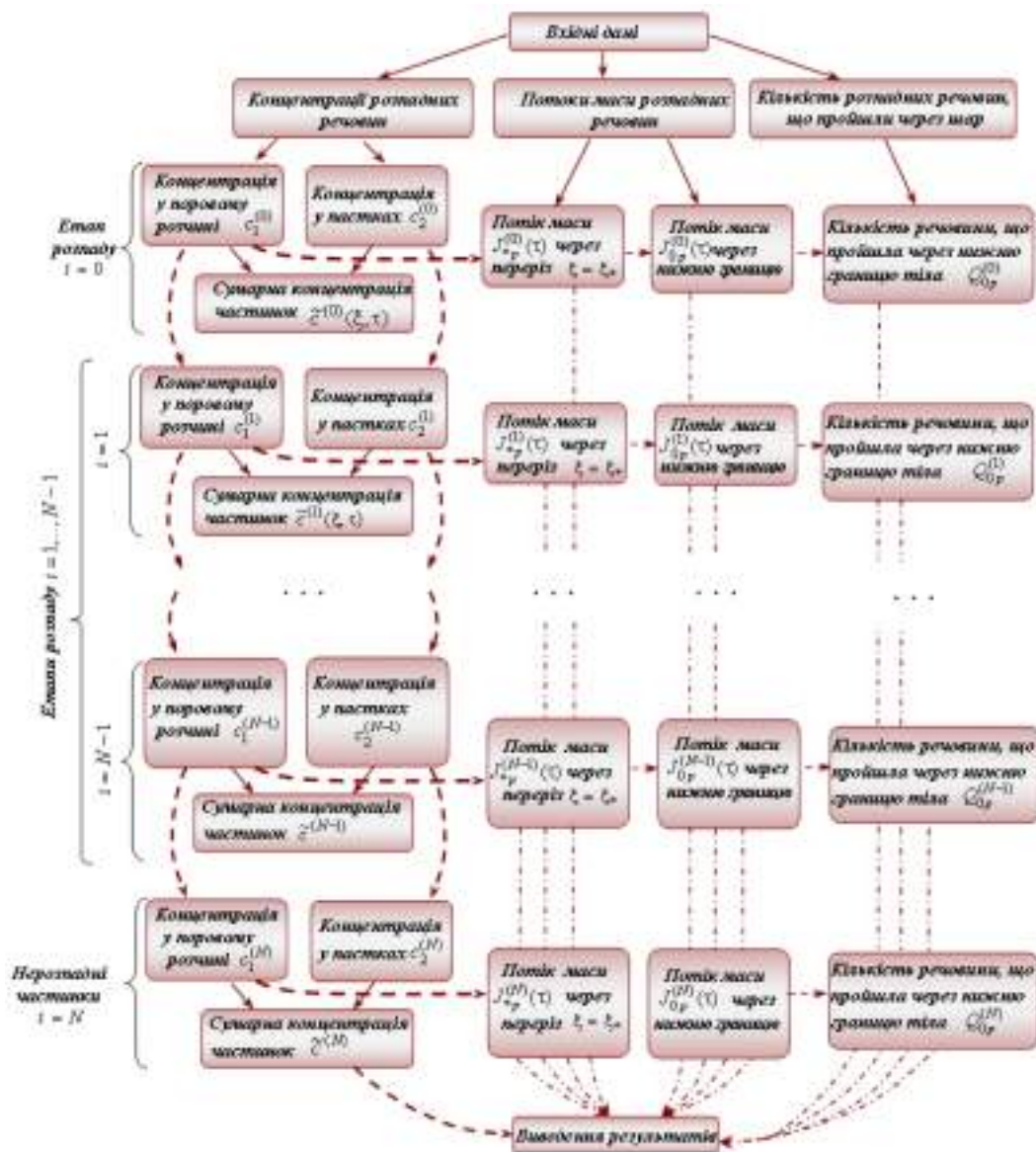


Рис. 1. Архітектура програмного комплексу

Зауважимо, що програмний модуль для розрахунку концентрацій містить два циклічних процеси, тоді як модулі для потоків і кількості речовини – тільки по одному (рис. 1).

Встановлено, що на нульовому етапі характерне суттєве приповерхнєве зменшення сумарної концентрації для малих часів. З часом сумарна концентрація у приповерхневій області зростає і виходить на усталений режим, набуваючи монотонно спадного характеру.

Шибирин А.Р., Савченко І.О.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Система підтримки прийняття рішень на основі двохетапного модифікованого методу морфологічного аналізу

Задачі прийняття рішень зустрічаються в усіх сферах людської діяльності і характеризуються великою різноманітністю. Значимість прийнятих рішень, а тим більше наслідки прийняття помилкових рішень, в деяких випадках можуть нести катастрофічний характер. Тому важливо, щоб продуктом безпосередньої діяльності особи, що приймає рішення (ОПР) було прийняття грамотних рішень. В умовах сучасного світу, коли кількість факторів, що впливають на об'єкт дослідження невідомо росте, досвіду та інтуїції ОПР стає недостатньо, що стало поштовхом до активної розробки систем підтримки прийняття рішень (СППР), які вирішують неструктуровані і слабоструктуровані багатокритеріальні задачі.

СППР представляє собою автоматизований програмний інструмент для аналізу даних, моделювання та прийняття рішень. Застосування СППР надає можливість використовувати обчислювальні потужності комп'ютерів для забезпечення об'єктивного та математично обгрунтованого результату при прийнятті рішень в складних умовах.

Недоліком більшості існуючих СППР є те, що вони не враховують зміни об'єкта на проміжку часу, тому отримані результати будуть актуальними лише для конкретного моменту і можуть стати недійсними в будь-який інший. Саме тому ціллю даної роботи є модифікація методу морфологічного аналізу (ММА) для врахування змін об'єкта в часі та розробка СППР на базі створеного алгоритму. Як метод якісного аналізу розглядається модифікований двохетапний ММА, що є потужним інструментом в процесі передбачення.

Врахування змін об'єкта з часом в ММА можливе шляхом додавання параметру часу до морфологічної таблиці в якості окремого характеристичного параметра [1], або шляхом введення додаткових елементів, що є зовнішніми по відношенню до морфологічної таблиці, у вигляді подій (миттєвих змін) та тенденцій (неперервних змін) [2]. Перший підхід є найпростішим, проте веде за собою низку зайвих оцінок, що є причиною значного зменшення точності методу. Тому в роботі використовуємо другий підхід.

Більшість процесів у природі досить складні і носять нелінійний характер, тому опису тенденцій в лінійному вигляді, як наведено в [2], не завжди достатньо для адекватного представлення змін об'єкта з часом. В своїй роботі ми використовуємо принципово новий підхід до визначення залежності тенденцій від часу, що носить логарифмічний характер і на відміну від підходу, розглянутого в [2], дозволяє більш точно аналізувати вплив тенденцій на початкові ймовірності альтернатив параметрів морфологічної таблиці.

В якості прикладу, на якому представлено роботу створеної СППР, розглядається проблема захворюваності населення. На першому етапі морфологічного аналізу аналізуємо неконтрольовані чинники захворюваності, на другому етапі – оцінюємо стратегії запобігання (пом'якшення) хвороб населення. В якості подій, що впливають на досліджуваний об'єкт, розглядаємо епідемії, а в якості тенденцій – сезонні спалахи захворюваності.

Отже, СППР на основі модифікованого двохетапного ММА з урахуванням змін об'єкта в часі дозволяє ОПР представляти, моделювати стан і поведінку об'єкта протягом тривалого часу після проведення дослідження без додаткового залучення фахівців з технологічного передбачення і експертів. За допомогою досить простої стратегії супроводу об'єкта забезпечується довгострокова актуальність морфологічної моделі об'єкта і можливість подальшого передбачення його поведінки в майбутньому.

Література. 1. Ritchey, T. Futures Studies using Morphological Analysis / Adapted from an article for the UN University Millennium Project: Futures Research Methodology Series, 2005. – 14 р. 2. Савченко И.А. Эволюция объекта исследования с привлечением модифицированного метода морфологического анализа // Системні дослідження та інформаційні технології. – 2015. – №2. – С. 122–130.

Шудренко Є.О.¹, Болдак А.О.^{1,2}

¹КПІ ім. Ігоря Сікорського, ФІОТ, Київ, Україна; ²Світовий центр даних з геоінформатики та сталого розвитку КПІ ім. Ігоря Сікорського, Київ, Україна

Практична реалізація методу Делфі

У даній роботі розглядається програмна реалізація методу Делфі, що є одним з найпопулярніших засобів проведення експертних опитувань та структуризації групових дискусій, що використовується для досягнення заданого рівня консенсусу по певному питанню, чи прогнозування [2]. Сутність методу полягає у ітеративному опитуванні запрошених фахівців, при якому, після кожної ітерації, розподіл відповідей поширюється серед учасників [3], що надає останнім можливість прослідкувати загальні тенденції та обдумати варіанти, запропоновані іншими експертами. Кожен має право додавати коментарі щодо обраного варіанту та всіляко обґрунтовувати його. При цьому гарантується анонімність респондентів та можливість вільно висловлювати свою думку. Даний підхід дозволяє за певну кількість турів згрупувати судження більшості експертів навколо одного напрямку [2], котрий, з високою часткою вірогідності, може бути прийнятий в якості резолюції або прогнозу на майбутній розвиток [3], за умови належного підбору учасників опитування для даної предметної області.

В процесі еволюції обчислювальних технологій метод Делфі отримав удосконалення у вигляді модифікації Real-time Delphi [1], що представляє собою проведення опитувань у мережі, зазвичай за допомогою веб-сервісів у режимі реального часу. Такий підхід відрізняється від традиційного наявністю технологічних можливостей відображення графіків розподілу відповідей на кожному етапі опитування, що виключає необхідність витрат часу на поширення результатів чергової ітерації серед учасників і пришвидшує процес в цілому. За використання веб додатку у якості платформи для проведення методу Real-time Delphi експерти можуть приймати участь в опитуванні у зручний для них час та редагувати свої відповіді онлайн, що підвищує рівень доступності метода [1] і забезпечує його асинхронність. Окрім цього суттєво збільшується максимальна кількість учасників та забезпечується оптимальний рівень анонімності.

На поточний момент існує декілька спеціалізованих сервісів для проведення опитувань за технологією Real-time Delphi, що мають обмежені можливості для налаштування як зовнішнього вигляду, так і структури анкет [1]. Для розробки останніх можуть бути також використані численні веб-портали – Google Forms [5], Survey Monkey [7], Typeforms [6] та інші. Хоча вони і забезпечують високий рівень візуалізації та якості відображення інформації, широкий вибір видів питань та налаштування їх поведінки, але жоден з них не надає інструментів для проведення опитувань методом Делфі, а саме: можливостей відображення розподілу відповідей учасників у режимі реального часу, проведення декількох турів опитувань у рамках однієї анкети, додавання власних варіантів відповідей, видимих іншим експертам, тощо.

Таким чином постає питання про необхідність створення веб-додатку з повним набором інструментів для реалізації експертних опитувань за методом Делфі. Варіантом рішення в цьому випадку може слугувати розроблений на платформі DJ-Portal додаток DJ-Forms, що забезпечує користувача необхідними засобами для проведення відповідних досліджень.

За допомогою DJ-Forms можливо створити анкету опитування з великою кількістю питань та гнучкими налаштуваннями їх відображення. У розпорядженні користувача знаходиться понад десяти видів питань, деякі з яких відсутні в інших сервісах для створення анкет. Для кожного виду наявний простий у використанні інструментарій для редагування відображення та організації питань в анкеті, приклад якого показаний на рис. 1.

Головними аспектами додатку DJ-Forms для проведення опитувань за методом Делфі є:

- Наявність поруч з кожним питанням гістограм розподілу відповідей усіх учасників опитування на попередньому етапі, що реалізує принцип впливу колективної думки при прийнятті рішення та дозволяє бачити віддаленість власної думки від найпопулярнішого варіанту, що демонструється на рис. 2.

- Можливість запрошення експертів за допомогою електронної пошти, що дозволяє залучити необхідну кількість обраних фахівців.
- Гнучкі налаштування доступу для різних категорій користувачів забезпечують достатню ступінь захищеності від випадкових відвідувачів.
- Інтуїтивний інтерфейс для додавання нових варіантів відповідей як адміністратору, так і учасникам опитування, що є основою ітераційного характеру методу.
- Можливість проведення необхідної кількості турів опитувань на одній анкеті.

Рис. 1. Вікно налаштувань питання типу *Scales*

Рис. 2. Діаграми розподілу відповідей з індикацією вибору

Таким чином, беручи до уваги усі вищеописані функціональні можливості та простоту використання, а також відсутність у мережі аналогічних рішень, можна стверджувати, що додаток DJ-Forms є ідеальним рішенням для проведення опитувань методом Делфі у режимі реального часу, що задовольняє потреби різноманітних за масштабами та складністю досліджень.

Література. 1. Gordon T.J. Real-time Delphi [Електронний ресурс] / Theodor J. Gordon. – 2009. – Режим доступу до ресурсу: <http://107.22.164.43/millennium/RTD-method.pdf>. 2. Linstone H.A. The Delphi Method Techniques and Applications [Електронний ресурс] / H.A. Linstone, M. Turoff, O. Helmer // ISBN 0-201-04294-0. – 2002. – Режим доступу до ресурсу: <https://web.njit.edu/~turoff/pubs/delphibook/delphibook.pdf>. 3. Chia-Chien H. The Delphi Technique: Making Sense Of Consensus [Електронний ресурс] / Hsu Chia-Chien // ISSN 1531-7714. – 2007. – Режим доступу до ресурсу: <http://pareonline.net/pdf/v12n10.pdf>. 4. Helmer O. Analysis of the Future: The Delphi Method [Електронний ресурс] / Olaf Helmer // RAND Corporation. – 1967. – Режим доступу до ресурсу: <https://www.rand.org/pubs/papers/P3558.html>. 5. Google Forms [Електронний ресурс] // Google – Режим доступу до ресурсу: <https://www.google.com/intl/en-gb/forms/about/>. 6. Sawers P. Try Typeform. [Електронний ресурс] / Paul Sawers. – 2014. – Режим доступу до ресурсу: <https://thenextweb.com/apps/2014/02/12/typeform/>. 7. Survey Monkey: The World's Most Popular Free Online Survey Tool [Електронний ресурс]. – 2018. – Режим доступу до ресурсу: <https://www.surveymonkey.com/>.



This page was left blank intentionally

Plenary talks

System analysis of
complex systems of
various nature

Intelligent systems for
decision-making

High Performance
Computing and
Microsystems Engineering

Progressive information
technologies

2

Intelligent systems for decision-making



Section 2

Intelligent systems for decision-making

1. Intelligent decision-making systems (IDMS) in finance-economical area (micro- and macro economical systems, banks, stock exchanges, insurance companies, etc.).
2. Decision-making systems in social processes' management.
3. Decision-making systems in technological processes' management in industry.
4. Intellectual analysis of data and knowledge; problems of data and knowledge mining, knowledge bases for IDMS.
5. Mathematical modelling and forecasting of complex objects and processes.
6. Decision-making under the data uncertainty conditions (systems of fuzzy logical deduction).
7. Modern methods and algorithms for IDMS (genetic and evolutionary algorithms, neural networks, etc.).

Секция 2

Интеллектуальные системы принятия решений

1. Интеллектуальные системы принятия решений (ИСПР) в финансово-экономической сфере (микро- и макроэкономические системы, банки, биржи, страховые компании и т.п.).
2. Системы принятия решений в управлении социальными процессами.
3. Системы принятия решений в управлении технологическими процессами в промышленности.
4. Интеллектуальный анализ данных и знаний; проблемы добыwania данных и знаний (Data&Knowledge Mining), базы знаний для ИСПР.
5. Математическое моделирование и прогнозирование сложных объектов и процессов.
6. Принятие решений в условиях неопределенности данных (системы нечеткого логического вывода).
7. Современные методы и алгоритмы ИСПР (генетические и эволюционные алгоритмы, нейронные сети и т.п.).

Секція 2

Інтелектуальні системи прийняття рішень

1. Інтелектуальні системи прийняття рішень (ІСПР) в фінансово-економічній сфері (мікро- та макроекономічні системи, банки, біржі, страхові компанії і т.д.).
2. Системи прийняття рішень в управлінні соціальними процесами.
3. Системи прийняття рішень в управлінні технологічними процесами в промисловості.
4. Інтелектуальний аналіз даних і знань; проблеми добування даних і знань (Data&Knowledge Mining), бази знань для ІСПР.
5. Математичне моделювання та прогнозування складних об'єктів і процесів.
6. Прийняття рішень в умовах невизначеності даних (системи нечіткого логічного виведення).
7. Сучасні методи й алгоритми ІСПР (генетичні та еволюційні алгоритми, нейронні мережі і т.д.).

Chertov O., Rudnyk T.

Igor Sikorsky Kyiv Polytechnic Institute, Faculty of Applied Mathematics, Kyiv, Ukraine

Search of phony accounts on Facebook and Twitter

Nowadays there is a lot of misinformation and disinformation on social networks. Our goal is to find such user accounts that are aimed at destabilizing the situation in a certain region or country by dissemination of the instigating statements on Facebook and Twitter social networks. To detect such so called phony accounts, it is insufficient to look at some structural features or content of their blogs or messages. Phony accounts can be detected only in dynamics, by analyzing their behavior in a given community. Since the harm done by such accounts to the corresponding community is maximal, and their presence can threaten the very existence of the community, the task of finding such accounts is one of the key ones in the information war against the enemy.

Experiment. In Ukraine the most popular politician Facebook page belongs to Mikheil Saakashvili, an opposition leader. Now he has more than one million followers and there are a lot of discussions under each post. Discovering Saakashvili Facebook page in period from 2017-10-01 to 2017-12-31 we have collected 316 posts with 51,970 comments. Four phony accounts were detected; all of them had more than 25 comments during observed period, mostly positive or neutral (see Fig. 1).

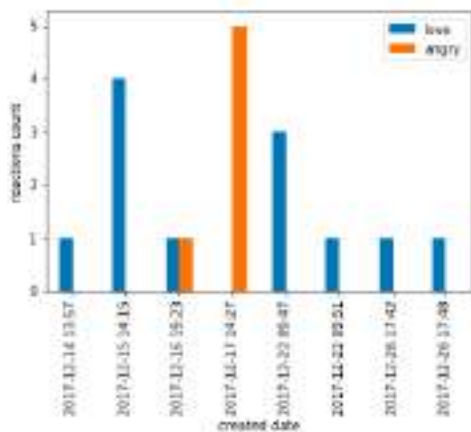


Figure 1. Phony account activity on Facebook

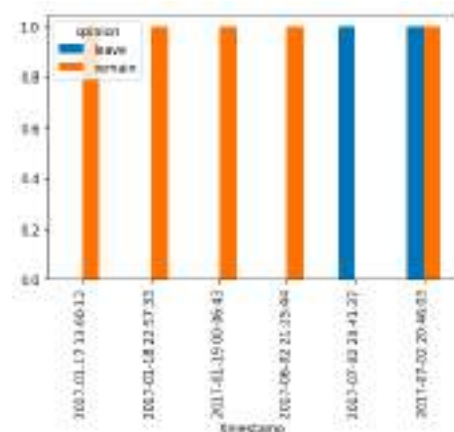


Figure 2. Phony account activity on Twitter

All of them started to write negative comments at 2017-12-17 – the date of the march “For Impeachment”. Such comments were about separatism, the incitement to riot, the lie of judges and the corruption of TV channels. These accounts can be related and followed by one leader.

Twitter is not popular in Ukraine unlike Europe and USA. For search of phony accounts 2,164 tweets about Brexit were collected. We assume hash tags LeaveEU and VoteLeave as opinion to leave European Union. Tags RemainEU and VoteRemain was marked as desire to remain in the EU. Two potentially phony accounts were found, because of their quick opinion change (see Fig. 2).

Conclusions. In this research, we proposed a novel approach for detecting phony accounts on Facebook and Twitter social networks. It was based on detecting groups of people who make mainly positive messages, but start write negative messages at the roughly same time. On Facebook other users who subscribed to the pages made the labeling of comments. On Twitter hash tags was used for labeling. The experimental results supported our approach by showing that such phony accounts were found on Facebook page of the Ukrainian politician Mikheil Saakashvili. Potentially phony accounts were found on Twitter where they try to change people opinion about Brexit. These findings can have significant consequences for fighting in information warfare.

Chertov O.R., Shanin V.O., Bobyr A.O.

Igor Sikorsky Kyiv Polytechnic Institute, Faculty of Applied Mathematics, Kyiv, Ukraine

Patterns of glycemia behavior for prediction of nocturnal hypoglycemia

Insular diabetes is a group of diseases with very high blood sugar levels over a prolonged period of time. These diseases require the correction of sugar level. The normal level of glucose in the blood is around 110–120 mg/dL (milligrams per deciliter). Hyperglycemia is registered when the blood glucose goes above 180 mg/dL. From the other hand, when the glucose level goes below 70 mg/dL hypoglycemia is faced. Hypoglycemia is often a result of diabetes treatment, when, for example, absence of a meal and increased physical activity together with diabetes treatment can cause the decline in blood glucose level. Nocturnal hypoglycemia (hypoglycemia during the night) is very dangerous. Due to the time of its appearance, it is difficult to detect and thus to treat.

Our aim is to predict nights when there is a risk of it. We have the data from DirecNet about the level of glucose of 57 patients during some period with the interval of 10 minutes. Six representative patients were chosen for evaluating the results, because these patients have more than five nocturnal hypoglycemia cases. These data have been already filtered, smoothed out, sampled and the most important peaks were identified. We reduce our task to the problem of finding patterns of glycemia behavior during days when nocturnal hypoglycemia event was noticed.

With the doctor's approval, the time for making predictions was boundaried from 7 pm until 11 pm. Also due to the natural processes, we were forced to leave out of attention the time after 3 am as the time when nocturnal hypoglycemia is possible. We found that the most important pattern is a parabolic behaviour of glycemia level (fig. 1). The parabolic approximation was built, but because the processes in a human body are slowed down at the night, it was a need of regularization. The idea of Tikhonov regularization was used, so the approximation has taken an individual character: individual regularization coefficients were selected.

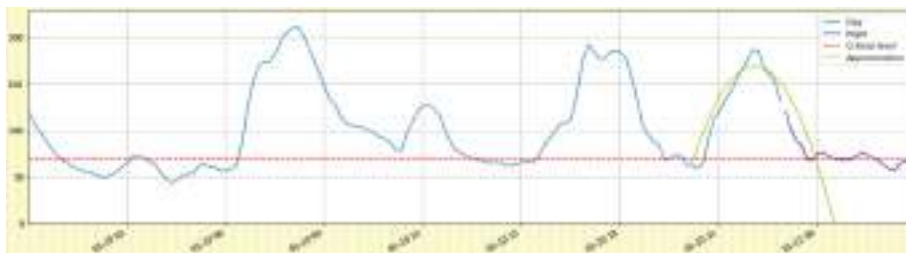


Figure 1. Parabolic behaviour of glycemia

The results in percentage are on the Table 1, as well as regularization Term for each of the presented patients.

Table 1. Results of prediction for some patients

ID	Term	MCC	Sensivity	Specifity	F1-score	F2-score
18	0.275	57.84%	90.91%	74.19%	68.87%	80.65%
25	0.150	58.97%	77.78%	84.91%	70.00%	74.47%
9	0.150	55.85%	63.64%	89.47%	70.00%	66.04%
35	0.250	42.43%	75.00%	74.42%	56.25%	66.18%
10	0.250	30.28%	61.90%	69.77%	55.32%	59.09%
11	0.800	46.16%	87.50%	65.22%	60.87%	74.47%

In perspective of the further research it is collecting the data that include the time and amount of meals and insulin injections, because this information has the great impact on the glycemia behavior.

Honcharevskyi D., Chertov O.

Igor Sikorsky Kyiv Polytechnic Institute, Faculty of Applied Mathematics, Kyiv, Ukraine

Violating group anonymity using machine learning techniques

Violating group anonymity task is gaining importance in modern world as source of uncovering purposely hidden information, which could harm proprietors of distributed data. As shown in paper [1] this task could be achieved by using evolutionary approaches, fuzzy models in particular. Therefore, aiming similar purposes one could achieve approximate results using machine learning techniques, such as decision trees (random forests as subordinate method) [2] and neural networks [3]. Both methods are being applied for classification problems. American Community Survey Public Use Microdata Sample concluded in state of Florida [4] is regarded as a training data. Sample consists of 196.829 recordings with 128 features. Occupancy of respondent is determined as violated feature, especially whether his or her occupation is military specific. Data aggregation, cleansing and partitioning are used to allocate essential parameters in problem's scope. It means that data aggregation serves purpose of transforming initial multifarious occupation feature into a binary relation establishing correspondence as true if a person is a military servant and false otherwise. In similar way, cleansing excludes data features of little importance in determining person's occupation. Partitioning divides data on training and testing data sets, where former one would be used for building random forests classifier and learning neural networks and latter one would be used with sole purpose of classifying unlabeled data using developed random forests and neural networks. Both built models are graphs. In first case, structure becomes an ensemble of decision trees, acyclic, connected graphs in particular. In second case, structure is being regarded as fully connected network, acyclic, fully connected graph in particular. For random forests decision making of classifying data is clean and explicit as Gini index, Entropy-based impurity measures are used for dividing training data on groups. That grouping goes on until some threshold is gained or subgroups consist only from one-class elements. In the end one could achieve partitioned domain, which is easily described and interpreted, as a nature of decision trees. For neural networks decision making of classifying data is implicit and hard to interpret, because tuning classifying accuracy could be achieved by neural network back propagation that minimizes loss of basis function being used. In general case, neural network could be built with arbitrary number of nodes on every hidden layer of network as well as arbitrary number of hidden layers itself. Every hidden layer's node is interpreted as some linear transformation of input space being feed into node. Thus, in the end one could achieve same partitioned domain or similar domain to the one built with random forest classifier. As both models serve same purpose of classifying unlabeled data, therefore one could assume that random forests as more explicit and comprehensive approach could explain neural network structure and attempt to correspond feature dividing decision in random forest with domain transformation being made with simple neuron. Therefore one should be able to attempt to reconstruct neural network structure by utilizing random forests structure. This approach applied to Florida State dataset could be generalized on other states of USA and even further trialed on some similar surveys of e.g. European countries. Comparison of results obtained on different datasets could improve overall accuracy of model, help better describe and explain partitioning decisions and uncover new, unknown dependencies and features that could reveal feature that are targeted for violation.

References. 1. Chertov, Oleg, Tavrov, Dan. Evolutionary approach to violating group anonymity using third-party data, SpringerPlus, 2016. 2. Murphy, Kevin P. Machine Learning: A Probabilistic Perspective. The MIT Press, 2012, pp. 544–552. 3. Bishop, Christopher M. Pattern Recognition and Machine Learning (Information Science and Statistics), Springer-Verlag New York, 2006, pp. 227–272. 4. American Community Survey Public Use Microdata Sample (PUMS) files. 2016 Florida Population Records. 2016 Florida Housing Unit Records. The Census Bureau, www.census.gov/programs-surveys/acs/technical-documentation/pums/documentation.2016.html.

Lazuka V.T., Rodionov A.M.

Igor Sikorsky Kyiv Polytechnic Institute, Institute of Physics and Technology, Kyiv, Ukraine

Combining Several Techniques of Adversarial Example Generation

Deep neural networks (DNNs) are popular machine learning models that achieve state-of-the-art results on challenging learning tasks in domains where there is adequate training data and compute power to train them. They have been shown to outperform humans in many tasks, for example, image classification and face recognition. Machine learning classifiers are known to be vulnerable to malicious inputs constructed to force misclassification. Those inputs are adversarial examples [1].

Adversarial examples. Adversarial examples are designed to cause a mistake. They are not defined to differ from human judgment. If adversarial examples were defined by deviation from human output, it would by definition be impossible to make adversarial examples for humans. On some tasks, like predicting whether input numbers are prime, there is a clear objectively correct answer, and we would like the model to get the correct answer, not the answer provided by humans. It is challenging to define what constitutes a mistake for visual object recognition, since after adding a perturbation to an image it likely no longer corresponds to a photograph of a real physical scene, and it is philosophically difficult to define the real object class for an image that is not a picture of a real object.

Adversarial examples are not defined to be imperceptible. If this were the case, it would be impossible by definition to make adversarial examples for humans, because changing the human's classification would constitute a change in what the human perceives. Moreover, in many domains, it is not possible to make imperceptible changes.

Generating adversarial examples. I consider three algorithms to generate adversarial examples.

Fast Gradient Sign Method (FGSM) [2]. This method linearizes the inner maximization problem in:

$$\mathbf{X}^{\text{adv}} = \mathbf{X} + \varepsilon \text{sign}(\nabla_x J(h(\mathbf{X}), y_{\text{true}})).$$

The Jacobian-based saliency map approach (JSMA) [3]. The method iteratively perturbs features of the input that have large adversarial saliency scores:

$$S(\mathbf{X}, t)[i] = \begin{cases} 0 & \text{if } \frac{\partial \mathbf{F}_t(\mathbf{X})}{\partial \mathbf{X}_i} < 0 \text{ or } \sum_{j \neq t} \frac{\partial \mathbf{F}_j(\mathbf{X})}{\partial \mathbf{X}_i} > 0, \\ \left(\frac{\partial \mathbf{F}_t(\mathbf{X})}{\partial \mathbf{X}_i} \right) \left| \sum_{j \neq t} \frac{\partial \mathbf{F}_j(\mathbf{X})}{\partial \mathbf{X}_i} \right| & \text{otherwise.} \end{cases}$$

Basic iterative method [4]. A straightforward extension of FGSM is to apply it multiple times with small step size:

$$\mathbf{X}_0^{\text{adv}} = \mathbf{X}, \mathbf{X}_{N+1}^{\text{adv}} = \text{Clip}_{x, \varepsilon} \{ \mathbf{X}_N^{\text{adv}} + \alpha \text{sign}(\nabla_x J(\mathbf{X}_N^{\text{adv}}, y_{\text{true}})) \}.$$

System. I evaluated attacks against VGG16 model built over Visual Geometry Group (VGG) [5] neural network. I use keras framework with TensorFlow backend to build neural network model.

The original VGG exhibited state-of-the-art results on the Labeled Faces in the Wild (LFW) benchmark, with 98.95% accuracy for face verification. The VGG architecture contains a large number of weights (the original DNN contains about 268.52 million parameters).

The model was trained on VGG-Face dataset. It contains 2622 identities. The dataset is designed to have no overlap with LFW.

The system is used for face verification (classification). There are 22 classes with 10 pictures for validation per each identity from LFW dataset.

Experiments. There are 7 experiments.

I use cleverhans v2.0 [6] to generate adversarial examples from validation data.

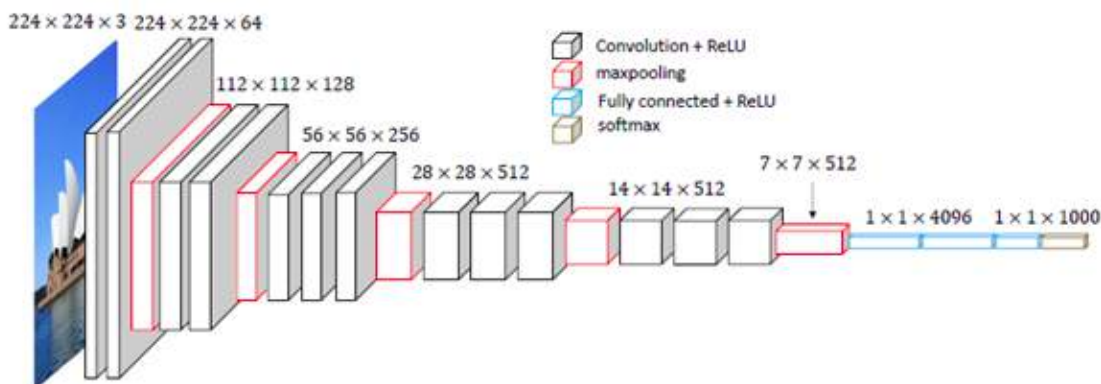


Figure 1. The architecture of VGG16 model

Table 1. Description of experiments

	1	2	3	4	5	6	7
FGSM	×			×	×		×
JSMA		×		×		×	×
Basic Iterative			×		×	×	×

Conclusion. Almost every machine learning classifier is vulnerable to adversarial inputs. The article shows it is more effectively for attacker to use combination of some 'simple' adversarial example generators than to use only one.

References. **1.** Christian Szegedy, Wojciech Zaremba, Ilya Sutskever, Joan Bruna, Dumitru Erhan, Ian Goodfellow, and Rob Fergus. Intriguing properties of neural networks. arXiv preprint arXiv:1312.6199, 2013. **2.** Ian J. Goodfellow, Jonathon Shlens, Christian Szegedy. Explaining and Harnessing Adversarial Examples. arXiv preprint arXiv:1412.6572, 2014. **3.** Nicolas Papernot, Patrick McDaniel, Somesh Jha, Matt Fredrikson, Z. Berkay Celik, Ananthram Swami. The Limitations of Deep Learning in Adversarial Settings. arXiv preprint arXiv:1511.07528, 2015. **4.** Alexey Kurakin, Ian Goodfellow, Samy Bengio. Adversarial examples in the physical world. arXiv preprint arXiv:1607.02533, 2016. **5.** Karen Simonyan, Andrew Zisserman. Very Deep Convolutional Networks For Large Image Recognition. arXiv preprint arXiv:1409.1556, 2014. **6.** Ian J. Goodfellow, Nicolas Papernot, Patrick D. McDaniel. cleverhans v2.0.0: an adversarial machine learning library. arXiv preprint arXiv:1610.00768, 2016.

Naderan M., Zaychenko Y.P.

Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

Diagnosing Lung Cancer Based on Deep Learning Algorithms: Review

Cancer is a leading cause of death worldwide. In the United States of America there are 225,000 people faced lung cancer every year, and the cost of the health care is about \$12 billion. The purpose of the paper is to review recent related studies and propose the best algorithm for diagnosis lung cancer in early stage.

Introduction. Most cancers that start in the lung, known as primary lung cancers, are carcinomas [1]. The two main types are small-cell lung carcinoma (SCLC) and non-small-cell lung carcinoma (NSCLC). The most common symptoms are coughing (including coughing up blood), weight loss, shortness of breath, and chest pains [2].

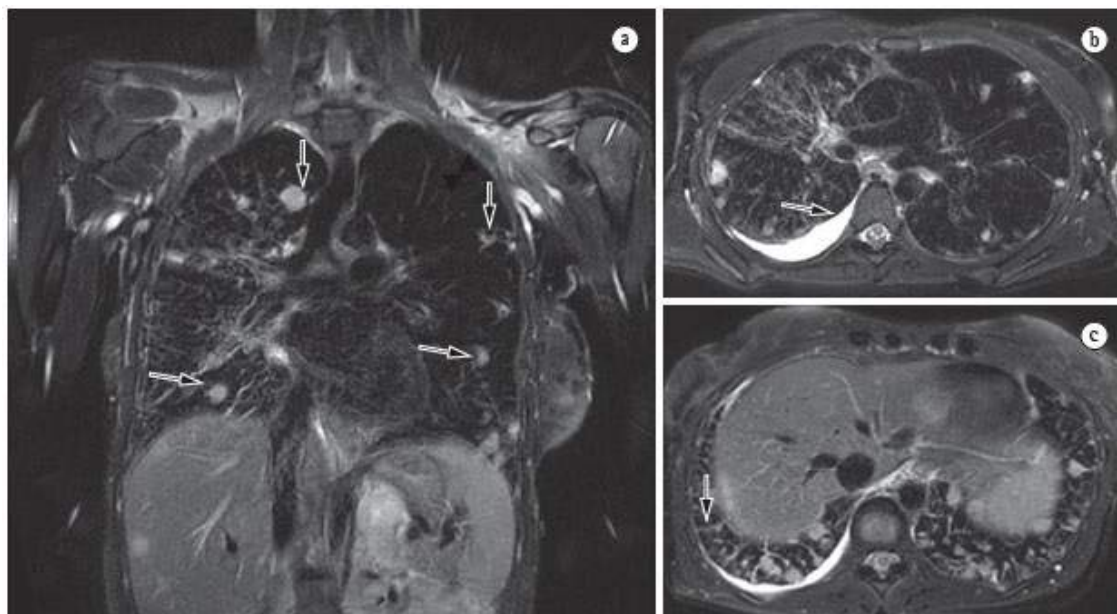


Figure 1. In a, control T2-weighted MRI with fat satyrating of a patient with a lung tumore, showing metastatic nodules with high signal intensity in the lung parenchyma (arrows). In b, axial T2-weighted MRI with fat saturation, showing pleural effusion (arrow). In c, axial T2-weighted MRI with fat saturation, showing septal lines suggestive of lymphangitic carcinomatosis (arrow).

Causes. Tobacco smoking for a long-term is the main cause (85%) of cases of lung cancer. However, the rest of cases of lung cancer is because of the genetic factors and exposure to radon gas, asbestos, second-hand smoke, or other forms of air pollution.

Methods. The focus of recent research [3] has been on using deep learning frameworks to automatically extract features from multi-modal preoperative medical images (i.e., T1 MRI, fMRI and DTI) of high-grade glioma patients. Specifically, they adopt 3D convolutional neural networks (CNNs), also propose a new network architecture for using multi-channel data and learning supervised features. Along with the pivotal clinical features, they train a support vector machine to predict if the patient has a long or short overall survival (OS) time. Their experimental results demonstrate that methods can achieve an accuracy as high as 89.9%. They also find that the learned features from fMRI and DTI play more important roles in accurately predicting the OS time, which provides valuable insights into functional neuro-oncological applications.

In [4], authors have compared CNN and endoscopists. The results obtained by authors in [4] have shown the sensitivity, specificity, accuracy, and diagnostic time were 88.9%, 87.4%, 87.7%, and 194 s, respectively, for the CNN. These values for the 23 endoscopists were 79.0%, 83.2%, 82.4%, and 230 ± 65 min respectively.

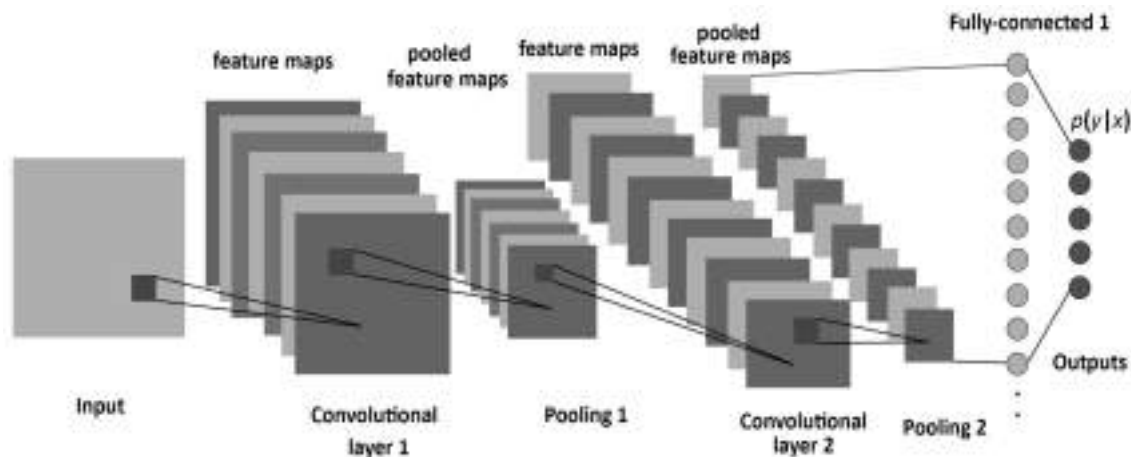


Figure 2. Typical Convolutional neural network architecture

Conclusion. From the research that has been carried out, it is possible to conclude that diagnosis ability of CNN is greater than endoscopists in general. Also, the diagnostic time of CNN is considerably shorter than manual diagnosis by endoscopists. Methods based on deep learning have attracted much attention due to their breakthrough performance for classification in many application domains, including image recognition and natural language processing.

References. 1. Bahmanyar S, Montgomery SM, Hillert J, Ekbom A, Olsson T. Cancer risk among patients with multiple sclerosis and their parents. *Neurology* (2009). 2. Wallace J. Brownlee FRACP, Todd A. Hardy, Franz Fazekas, David H. Miller. The diagnosis and differential diagnosis of multiple sclerosis: progress and challenges. 3. Dong Nie, Han Zhang, Ehsan Adeli, Luyan Liu, and Dinggang Shen. 3D Deep Learning for Multi-modal Imaging-Guided Survival Time Prediction of Brain Tumor Patients. *Med Image Comput Comput Assist Interv.* 2016 Oct; 9901: 212–220. 4. Satoki S., Shuhei N., Kazuharu A., Yoshitaka N., Motoi M. Application of Convolutional Neural Networks in the Diagnosis of *Helicobacter pylori* Infection Based on Endoscopic Images. *EBioMedicine*, Volume 25, November 2017, Pages 106–111.

Nikolaiev S.S.^{1,2}, Tymoshenko Y.O.¹

¹Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine; ²Dekom, Kyiv, Ukraine

RBF neural network probabilistic skin color detector for real-time video applications

Human skin color is a useful feature for several computer vision researches including but not limited to determining heart rate from video stream. Statistical models representing the human skin color property are used to build robust and reliable skin detectors but for real-time video-processing applications on some embedded devices computational complexity of the method is also important. In this paper RGB color space projection to luminance independent plane (LIP) is proposed for building fast probabilistic skin color detector. The probability distribution of skin colors from SFA dataset is explored. Obtained color distribution is modeled with RBF neural network that returns probability of RGB pixel being pixel from skin region on the image.

Introduction. Fast human skin detection in a video frames plays an important role in various real-time applications, e.g., face detection in videos [1], hand gesture recognition for human-computer interaction [2], heart rate variability (HRV) extraction using web- and video cameras [3]. Many approaches of skin detection have been extensively explored over the years, but detecting human skin in real-world applications is still nontrivial and challenging task due to many factors including variable ambient light, noisy cameras' sensors, backgrounds similar to skin, diversity of human races and many more.

Dataset. In this paper a human skin image database called SFA was used for models training. The SFA image database showed great potential to assist research which have skin color or texture as a relevant feature. SFA was constructed based on face images of FERET (876 images) and AR (242 images) databases, from which skin and non-skin samples and the ground truths of skin detection were retrieved [4].

Projecting RGB color space to luminance independent plane. Color spaces like RGB, CMY, LAB, LUV, YCbCr [5] and many others were extensively explored in literature for creating luminance independent representation of skin tint that would allow creating simple models for skin and face detection. But many of these color spaces contain non-linear computationally costly transformations, which are unnecessary in some applications.

So let's consider observing vector $\vec{x}$ with RGB values of some skin color. If more light illuminate the scene then each value in $\vec{x}$ will grow proportionally to luminance in direction $[1, 1, 1]^T$. So to compensate or even completely remove illumination dependence for vector $\vec{x}$, it can be projected to the plane P orthogonal to the vector $[1, 1, 1]^T$ so that two dimensional projection contains mostly color tint and almost no luminance information. P can be represented with:

$$P = \begin{bmatrix} 0 & 1 & -1 \\ -1 & 0.5 & 0.5 \end{bmatrix}, \quad (1)$$

This matrix can be used as projection operator for RGB column vector $\vec{x}$ into illumination independent space.

$$y(\vec{x}) = P \cdot \vec{x}, \quad (2)$$

$\vec{y}$ – denotes 2D coordinates on the color projection plane, given pixel RGB color vector $\vec{x} \in R^3$. Projecting 124,7 million skin color samples from dataset at LIP creates probability distribution of these samples.

Modeling probability distribution with RBF neural networks. For modeling probabilistic distribution D of colors on the projection plane P , artificial neural network with radial basis functions (Gaussian) as activation functions is used. The RBF network model is the following:

$$\phi(\vec{y}) = \sum_{i=1}^N a_i \cdot e^{-\|\vec{\beta}_i \cdot (\vec{y} - \vec{c}_i)\|^2},$$

where $\vec{\beta}_i$ can be thought as the vector that scales corresponding dimensions of $\vec{x}$ so that original data distribution could be covered with the Gaussian. To find the optimal model the following steps should be conducted:

1. Finding clusters and assigning vectors $\vec{c}_i$ as centers for each cluster.
2. Computing scales $\vec{\beta}_i$ for each neuron, so that the basis Gaussian functions cover not just some points in the centers of the clusters, but whole clusters.
3. Finding weights a_i with the gradient descent.

Rotation and stretching of the projection plane. Because the whole skin tints probability distribution represents one single cluster, to simplify and fasten the model lets translate the center of distribution $\vec{c}$ to 0-point and perform principal components analysis (PCA) rotating the axis of the plane in alignment with principal components vectors using rotation matrix G . Thus the model becomes:

$$\phi(\vec{y}) = a_0 \cdot e^{-\|W_0 \cdot (G \cdot P \cdot \vec{x} - \vec{c})\|^2} = a_0 \cdot e^{-\|R \cdot \vec{x} + B\|^2},$$

where $R = W_0 \cdot G \cdot P$ is the rotation-stretching matrix, $B = -W_0 \cdot \vec{c}$ – bias that shifts center (mean vector) of the distribution to point $[0, 0]$ on plane P . G and P are known from previous steps, $a_0 = 1$ since our network should return probabilities, and W_0 is stretch matrix estimated using Adam optimizer to fit the model to the probabilities distribution D .

Results. After 43 epochs of training, resulting vector W_0 was obtained:

$$W_0 = \begin{bmatrix} 11.54891205 & 36.16787338 \end{bmatrix}^T. \quad (3)$$

After substituting all matrices together the following numerical values are obtained:

$$R = \begin{bmatrix} 6.23055543 & 2.16381916 & -8.3943746 \\ 16.53261421 & -27.77861606 & 11.24600185 \end{bmatrix}$$

and

$$B = \begin{bmatrix} -1.32159356 & -0.22516318 \end{bmatrix}^T. \quad (4)$$

Conclusion. In this paper, human skin color space is explored using SFA dataset, efficient encoding of the skin colors distribution and its approximation with RBF neural network is introduced. As the result robust probabilistic color-based skin detector model is proposed for real-time video applications aiming to estimate the probability of each pixel being the pixel of skin or non-skin. Because models that use spatial information about pixel textures arrangements give better results than models relying on individual pixels, independent from its neighbors, Gaussian blur is applied to analyzed images to smooth deviations of single pixels colors so that color information is mixed between adjacent pixels. Using this approach before individual pixel classification, leads to analyzing skin in spatial context smoothing regional properties of skin and reducing misclassified outliers.

References. **1.** Hajraoui, A., Sabri, M., (2014). Face Detection Algorithm based on Skin Detection, Watershed Method and Gabor Filters, International Journal of Computer Applications (0975–8887) Volume 94–No 6, pp 33–39. **2.** Nalepa, J., Grzejszczak, T., Kawulok, M., (2014). Wrist localization in color images for hand gesture recognition, Man–Machine Interactions 3, Springer International Publishing, pp. 79–86. **3.** Kumar, M., Veeraraghavan, A., Sabharwal, A., (2015). DistancePPG: Robust non-contact vital signs monitoring using a camera, Electrical and Computer Engineering, RICE University, Houston, TX, USA, Vol. 6, No. 5, p. 24. **4.** Casati, J., Moraes, D., Rodrigues, E., (2013). SFA: A Human Skin Image Database based on FERET and AR Facial Images. In: IX Workshop de Visão Computacional, Rio de Janeiro. Anais do VIII Workshop de Visão Computacional. **5.** Jedynek, B., Zheng, H., Daoudi, M., (2003). Maximum entropy models for skin detection, in: Proceedings of IEEE Conference on Computer Vision and Pattern Recognition, pp. 180–193.

Onanko A.M.

Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

Prediction of Russian content on Ukrainian television

In the context of the active development of television and its use as a tool of propaganda and zombie of people, the task of forecasting the quantity of Russian content on Ukrainian television is one of the most important tasks in the field of television.

Historically, Ukraine and Russia were closely linked, not only by the border, in the course of many thousands of years of existence. “Undina”, “Not born beautiful”, “Carmelite” and “Brigade” – series that were watched back time by entire families in the evening prime time. The Ukrainians unwittingly studied the views of Moscow and the Russian way of life. After the Revolution of Dignity and the beginning of a military conflict in the Donbas, the annexation of the Crimea, Russian TV shows broadcast on Ukrainian television channels were used as an element of propaganda. The spectators in Ukraine watched and might have been proud of Russian “sea devils”, spies, FSB operatives and cops. In this regards, it is important to develop forecasting methods that will work under the conditions of uncertainty.

The key tasks of the work were:

1. To conduct research of tasks related to the definition of the level of Russian content on Ukrainian television, the allocation of factors that affect its number
2. To analyze input data and understand what indicators related to television and economic aspects of Russian product purchasing should be used as input for models in order to get the best results of the forecast (to achieve maximum efficiency), to lead them to a preliminary target, and also think of the architecture for monitoring Russian TV product
3. To conduct a comparative analysis and describe the main advantages and disadvantages of selected methods for forecasting the level of Russian television product on Ukrainian television.

Input data. The most significant indicators that were submitted to the input were the following parameters:

1. The number of views and displays of the Russian product on 6 major Ukrainian channels starting from 2014 and up to the beginning of 2018.
2. The main economic indicators for Ukraine and the change in the cost of purchasing a product in a neighboring country (monthly) from 2014 and up to the beginning of 2018.
3. Political indicators from 2014 and up to the beginning of 2018.

Methods. Two methods for forecasting were applied: Group method of data handling and Fuzzy neural network. The advantage of the inductive modeling method of GMDH is the possibility of constructing an objective model during the operation of the algorithm, as well as the ability to work on short samples.

Neural networks are now developing at a crazy pace, so it was interesting to combine these two methods to solve one problem and see which method showed itself better. Also they have a number of features, such as:

- the ability to work with incomplete and uncertain data;
- the possibility of taking expert knowledge into account in the form of fuzzy predicate rules of the output of type IF-THEN-ELSE.

Conclusions. The Ukrainian media market has been analyzed for the presence of Russian content in different years, the periods of economic crisis and stability, aggravation of political relations with the Russian Federation and the beginning of the war. The obtained results have been compared using various forecasting methods.

References. 1. Основы вычислительного интеллекта, М.З. Згуровский, Ю.П. Зайченко, Изд. «Наукова Думка». Киев, 2013.- 406 с. 2. Каллан Р. Основные концепции нейронных сетей / Роберт Каллан – М. : Вильямс, 2001. – 287 с.

Tkachenko D.A.

Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

Modeling spatial information with convolutional and recurrent neural networks for video classification

Video classification is an important problem which arises in robotics, security, content moderation, autonomous vehicle control and other fields. Datasets that contain recordings of various human actions are usually used as a benchmark for machine learning models that perform video classification.

Complete video classification systems usually use spatial, motion and audio information [1]. Spatial information refers to the stream of still video frames. Motion clues are derived from temporal information and can be encoded as an optical flow. Audio information can be represented as a spectrogram and derived features like Mel-Frequency Cepstral Coefficients (MFCC).

This paper focuses on modeling spatial information which is a crucial part of any video classification model.

Dataset. Solution is evaluated on UCF-101 action recognition benchmark [2] which contains 13320 clips (~27 hours @ 25 fps) with 320x240 resolution annotated into 101 classes. Experiments are performed on the first train/test split only.

Sequence modeling. Recurrent neural networks (RNN) are often used for processing sequential data. In the case of video classification, the input to RNN has the shape of (timestep, frame representation), thus for every timestep, there is some representation of a frame (raw pixels or embedding vector). “Vanilla” RNNs usually struggle with long sequences; to address this issue, variants like LSTM [3] and GRU [4] were developed.

Frame embedding. It’s not practical to feed raw pixels into RNN, so convolutional neural network (CNN) is used to extract features (embedding vector). This embedding vector is a dense representation of various features (visual patterns) that are present in a certain frame. For feature extraction, InceptionResNetV2 [5] model with weights pre-trained on ImageNet, is used. This model achieves 80.4% accuracy on the test set of the ImageNet classification (CLS) challenge and 95.3% top-5 accuracy. Fine-tuning upper layers (3 Inception blocks) on UCF-101 didn’t yield significant improvement, so pre-trained weights are used. To retrieve frame embedding, the upper dense layer (used for classification) is removed and the output of the top global average pooling layer, which produces a 1536-dimensional vector, is used.

Preprocessing. Before feeding frames to CNN, they are resized to 299x299x3 which is the default input size for InceptionResNetV2. Videos are also subsampled to 5 frames per second, and limited to the first 36 seconds (thus limiting the maximum length of sequence processed by RNN to 180).

Model architecture. CNN receives raw image pixels as input and produces 1536-dimensional embedding which is fed into RNN model (LSTM/GRU). The output of RNN model (video embedding) is processed by dense (fully-connected) layer which produces probabilities for every one of 101 classes. This architecture is visualized in the fig. 1. Dropout (with the rate of 0.5) is used on the input layer, between RNN model and top dense layer, and between LSTM/GRU layers.

Multiple RNN layers. It’s possible to make RNN model deeper to help it learn more complex representations at higher levels of abstraction. This can be done by using multiple LSTM/GRU layers. There are two approaches



Figure 1. Model architecture

to do so. First is *stacking* them on top of each other. In this approach, intermediate layer would return representation for each timestep, instead of a single embedding for an entire sequence. These sequences are input to a subsequent RNN layer. Another approach is using *bidirectional* layers. A bidirectional layer consists of two unidirectional layers, where one of them processes the input as is, and another one looks at the reversed sequence. Then representations from both layers are

concatenated. This allows to preserve information from both “past” and “future” at any point in time. Bidirectional layers can also be stacked.

Video embedding. An output of an RNN model provides compact video representation. Instead of feeding it into a fully-connected layer, it can be used as input to an ensemble model. For example, in a more complex video classification model which uses the fusion of spatial, audio and motion information, video representation can be combined (by some model, or using averaging) with other representations to produce a more accurate result.

Evaluation. Different RNN model architectures were trained with a varying number of LSTM/GRU units using stochastic gradient descent (SGD) with Nesterov momentum (also known as Nesterov Accelerated Gradient or NAG). Loss function used is the cross-entropy loss (top dense layer used softmax activation function). The best model (marked with (*)) was evaluated with linear support vector machine (SVM) (i.e., hinge loss, linear activation and L_2 regularization set to 0.01 on the last dense layer), results are marked with (**). Adaptive gradient methods (like Adam, RMSProp, and Adadelta) didn’t yield significantly faster convergence, and NAG usually produced better results, which is consistent with published observations [6]. Training has been run for some number of epochs until validation loss had stopped improving: if it doesn’t decrease for 10 epochs, the training is ceased. Also learning rate is reduced by a factor of 10 if validation loss hasn’t decreased for 5 epochs. Initial learning rate for SGD is 10^{-2} , and the lowest possible value (after reductions) is 10^{-6} . Batch size is 32.

Results. Presented in the table 1. “S.” denotes two *stacked* layers, and “Bi” prefix means *bidirectional* layer was used; specified number of units is per single LSTM/GRU layer in a model. Results show that RNN layers with 512 or 1024 layers generally perform best on this task. Stacked GRU layers provide better accuracy than stacked LSTM layers, likely because they are less prone to overfitting since they have fewer parameters. The best accuracy was achieved by training the best model (stacked bidirectional GRU) with SVM as last layer (instead of softmax and cross-entropy loss).

Conclusion. This paper proposed the architecture for modeling spatial information with CNN and RNN for video classification and experimentally compared the performance of different RNN models. The best model achieved the accuracy of 79.40% on the UCF-101 dataset.

References. 1. Z. Wu, Y.-G. Jiang, H. Ye, X. Wang, and X. Xue. “Multi-Stream Multi-Class Fusion of Deep Networks for Video Classification,” ACM Multimedia, 2016, pp. 791–800. 2. K. Soomro, A.R. Zamir, and M. Shah. “UCF101: A Dataset of 101 Human Actions Classes From Videos in The Wild,” CVCV-TR-12-01, 2012. 3. S. Hochreiter and J. Schmidhuber. “Long short-term memory,” Neural Computation. 9 (8), 1997, pp. 1735–1780. 4. J. Chung, Ç. Gülçehre, K. Cho, Y. Bengio “Empirical evaluation of gated recurrent neural networks on sequence modeling,” CoRR, 2014, abs/1412.3555. 5. C. Szegedy, S. Ioffe, V. Vanhoucke, and A. Alemi. “Inception-v4, Inception-ResNet and the Impact of Residual Connections on Learning,” AAAI, 2017, pp. 4278–4284. 6. A.C. Wilson, R. Roelofs, M. Stern, N. Srebro, and B. Recht. “The Marginal Value of Adaptive Gradient Methods in Machine Learning,” NIPS, 2017, pp. 4151–4161.

Table 1. Results

Model	Units	Top-1, %	Top-5, %
LSTM	256	76.24	93.49
LSTM	512	77.09	93.09
LSTM	1024	77.07	93.62
LSTM	2048	76.22	93.62
GRU	512	77.83	93.88
GRU	1024	77.81	93.17
BiLSTM	512	77.36	93.41
BiLSTM	1024	75.93	93.49
BiGRU	512	77.33	93.46
BiGRU	1024	77.52	93.41
S. LSTM	512	74.71	92.8
S. LSTM	1024	76.8	94.09
S. GRU	512	77.09	93.03
S. GRU	1024	77.22	93.25
S. BiGRU (*)	512	78.34	93.51
S. BiGRU	1024	77.15	93.17
S. BiLSTM	512	74.50	92.77
S. BiGRU (**)	512	79.40	90.84

Tkachenko D.A., Kharchenko K.V.

Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

Deep learning approach to audio analysis

Audio recognition is used for the variety of applications, e.g., speech recognition (voice control, speech to text, etc.), content-based music recommendation, video classification, etc.

Deep learning and particularly convolutional neural networks have recently provided a significant boost in classification accuracy and quality of learning audio representations.

Datasets. There is a large variety of datasets available for training machine learning models for audio recognition tasks which include labeled recordings of urban or environmental sounds, music, speech commands, etc. One of the largest datasets is *AudioSet* [1] which contains more than 2 million sound clips drawn from YouTube videos.

Feature engineering. Somewhat contrary to the idea of deep learning where a sufficiently deep model is expected to be able to automatically learn complex feature representations (potentially similar to hand-crafted features), in practice special features (specific to audio analysis) of a signal significantly improve model performance. A conventional approach [2] is to transform raw audio tracks (represented as signals in the time domain) into a *spectrogram* (by computing the *short-time Fourier transform* (STFT)), scale it and use as an input to a machine learning model.

Short-time Fourier transform. Fourier transform is the fundamental method to transform the time-domain signal into frequency domain. Music signals are highly non-stationary, so it is meaningless to compute a single Fourier transform over an entire audio recording (e.g., song). So to retrieve a spectrogram, short-time Fourier transform can be used. STFT is obtained by computing the Fourier transform for successive frames in a signal: $X(m, \omega) = \sum_n x(n)w(n-m)e^{-j\omega n}$.

STFT is a function of time m , and frequency, ω . As m increases, the window function w slides to the right. Fourier transform is computed for the resulting frame $x(n)w(n-m)$. An alternative approach to the STFT is the wavelet transform [3].

Spectrogram. The spectrogram shows the intensity of frequencies over time. A spectrogram is the squared magnitude of the STFT: $S(m, \omega) = |X(m, \omega)|^2$.

Scaling. The human perception of sound intensity is logarithmic in nature. Thus it makes sense to use log amplitude. In practice [2, 4], *log-mel spectrogram* is often used. The mel scale is a scale of pitches judged by listeners to be equal in distance one from another.

Mel-frequency cepstral coefficients. *Mel-frequency cepstral coefficients* (MFCC) are a small set of features (usually about 10–20) which concisely describe the overall shape of a spectral envelope. They can also be used as features input to a model [2].

Other features. Using spectrogram means that only the magnitude spectrum is considered. However, it has been shown [5] that processing *phase spectrum* yields a significant increase in classification accuracy. Other useful features include *energy* (how loud the signal is) and *filterbank* (band-pass filters).

Neural network architectures. Most state-of-the-art results on audio recognition tasks have been achieved using various convolutional neural networks [2–7]. Recurrent neural networks have also found their use [7] since audio recording can be viewed as a sequence so that sequence modeling techniques can be applied.

Unsupervised representation learning. A usual approach to audio feature learning is to use labeled tracks (supervised learning). However, one can also utilize some supplemental information (which accompanies audio track). In [6] acoustic representation is derived from unlabeled videos by leveraging natural synchronization between vision and sound.

Image classification networks. There have recently been lots of advancements in image recognition. Sophisticated models can achieve around 80% accuracy [8] on challenging datasets like ImageNet. These models have been successfully applied to audio classification [4]. One approach to harnessing

these models is to turn audio classification task into image classification. Spectrograms can be naturally visualized: one axis represents time, the other is frequency, and color intensity represents the third dimension – amplitude of a particular frequency at a specific time. It's also useful to use weights pre-trained on some image classification dataset (e.g., ImageNet) and fine-tune a number of top layers. Most of the layers need to be fine-tuned because spectrograms are very different from the images included in image classification datasets. However, transfer learning is still beneficial since it makes it possible to reuse the ability of network's lower layers to recognize basic geometric features, which accelerates convergence.

WaveNet. One of the impressive applications of deep learning in the audio analysis is *WaveNet* – a deep neural network for generating raw audio waveforms [9]. It has been developed for text-to-speech systems, but can also be used as a discriminative model wherever dilated convolutional neural network is applicable.

Feature extraction. Audio recognition models can be used as feature extractors to retrieve a semantically meaningful, high-level embedding vector for an audio recording from some input features. This embedding vector is much more compact than the raw audio input. For example, the *VGGish* model presented in [4] can extract 128-D embedding which can be fed as input to a downstream classification model. One of the applications of audio feature extraction is video classification: audio track embeddings can be used as input to fusion model, which combines them with other information streams like spatial and motion.

Interpretation through learning to generate. Recent work [10] has suggested the method where an agent is adversarially trained to reconstruct images by interacting with a simulator and a discriminator network. An agent is presented with a set of images, and it learns to produce (draw) pictures that resemble the real ones (similarity is measured by the discriminator). This way it can discover a sequence of actions (e.g., brush strokes) that re-create presented images. By doing so, an agent recovers structured representations from the raw input; this representation can be meaningful and succinct and can be useful for tasks like classification. The similar approach can be utilized for audio analysis: for example, an agent can interact with musical instrument simulator and learn to re-create some real recordings; a learned sequence of interactions (e.g., key presses on a piano) can be used as a concise representation of a musical composition.

Conclusion. In this paper, we investigated the application of deep learning models in various audio analysis tasks and described the commonly used feature engineering process. We proposed transforming audio tracks to images and utilizing image recognition networks and transfer learning, and also suggested learning representations using reinforced adversarial learning.

References. 1. J.F. Gemmeke et al. "Audio Set: An ontology and human-labeled dataset for audio events," IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), 2017, pp. 776–780. 2. A. van den Oord, S. Dieleman, and B. Schrauwen. "Deep content-based music recommendation," NIPS, 2013. 3. S. Qu, J. Li, W. Dai, and S. Das. "Understanding Audio Pattern Using Convolutional Neural Network From Raw Waveforms," CoRR, 2016, abs/1611.09524. 4. S. Hershey et al. "CNN architectures for large-scale audio classification," IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), 2017. 5. R.N. Tak, D.M. Agrawal, and H.A. Patil. "Novel Phase Encoded Mel Filterbank Energies for Environmental Sound Classification," PReMI, 2017. 6. Y. Aytar, C. Vondrick, and A. Torralba. "SoundNet: Learning Sound Representations from Unlabeled Video," NIPS, 2016. 7. Y. Xu, Q. Kong, W. Wang, and M.D. Plumbley. "Large-scale weakly supervised audio classification using gated convolutional neural network," CoRR, 2017, abs/1710.00343. 8. C. Szegedy, S. Ioffe, V. Vanhoucke, and A. Alemi. "Inception-v4, Inception-ResNet and the Impact of Residual Connections on Learning," AAAI, 2017, pp. 4278–4284. 9. A. van den Oord et al. "WaveNet: A Generative Model for Raw Audio," CoRR, 2016, abs/1609.03499. 10. Y. Ganin, T. Kulkarni, I. Babuschkin, S.M.A. Eslami, O. Vinyals. "Synthesizing Programs for Images using Reinforced Adversarial Learning." [Online]. Available: <https://deepmind.com/documents/183/SPIRAL.pdf> [Accessed: 28 Mar. 2018].

Zaychenko Y.¹, Galib H.^{1,2}

¹Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine; ²Azershig Co, Baku, Azerbaijan

Inductive Modeling Method GMDH in the Problems of Big Data Analysis

Last year's problems of Data Mining (DM) in data bases (DB) have become very crucial in IT-applications. Especially it refers to big DB, where mountains of raw data are accumulated and hidden laws in these data are to be detected and corresponding models to be constructed [1]. Previously several classes of methods were developed for finding unknown dependencies in data, in particularly statistical methods: ARMA, Logit and Probit models, ARCH and GARCH methods and neural networks. But they have drawbacks: statistical methods solve only problems of parametric identification and don't solve structural identification while neural networks allow to determine model structure but in an implicit form. The model structure is hidden in neural weights and its analytical form is unavailable.

Therefore the development of methods for detecting unknown dependencies in raw data and constructing structural models identification constitute important problem in BD mining. The main goal of this paper is development and investigation of novel method for constructing models in data accumulated in data warehouses. For this goal the method of inductive modeling- Group Method of Data Handling (GMDH) is suggested and investigated [2].

For finding unknown laws in data under uncertainty new version of GMDH – Fuzzy GMDH is suggested and explored [3,4]. Fuzzy GMDH enables to operate with incomplete or indefinite initial data and constructs fuzzy models whose coefficients are fuzzy.

The significant property of GMDH is that it may operate with high dimensional data (with many variables) and so-called "short samples" when the number of model coefficients m is greater than sample size N . This is achieved due to specificity of GMDG algorithm as at each step of it a set of so-called partial models are constructed consisting only of two variables instead of n initial input variables like other modeling methods. This enables to reduce substantially the dimension of model and decrease the time for its construction. This advantage rises with the increase of model complexity: the greater is model dimension (number of variables), the greater is cut in computational time for its construction as compared with conventional modeling methods.

As an important and wide-spread field of DM is forecasting of non-stationary time series in economy and financial sphere. This problem is traditionally solved with application of neural networks and classic statistical methods like ARIMA. But the main drawback of the conventional methods is that they doesn't allow to construct analytical forecasting models. In presented work inductive methods GMDH and Fuzzy GMDH are suggested and investigated for the problem of forecasting non-stationary financial processes at financial markets.

For estimating their efficiency they are compared with neural network Back propagation and conventional statistical methods. The experiments have shown advantages of the suggested inductive modeling methods over statistical forecasting methods. The experimental investigations results are presented and discussed in a report.

References. 1. Barsegyan A.A. Technologies of data analysis: Data mining, Visual Mining, Text Mining, OLAP. / A.A. Barsegyan, M.C. Kuprianov, V.V. Stepanenko, I.I. Holod. – 2-nd edition, revised and add.). – SPb.: BHV – Petersburg, 2008.- 384 p. 2. Ivakhnenko A.G., Mueller I.A. Self-organization of forecasting models.-Kiev: Publ. House "Technika". – 1985. 3. Zaychenko Yu.P. Fuzzy Group Method of Data Handling // System research and information technologies. – 2003. – №3. – pp. 25–45. 4. Zaychenko Yu.P. The Fuzzy Group Method of Data Handling and Its Application for Economical Processes forecasting // Scientific Inquiry, vol. 7, No 1, June, 2006. – pp. 83–98.

Акимов В.С.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Алгоритми багаторівневого навчання для класифікації захворювань шкіри

Захворювання шкіри сьогодні належать до розповсюджених медичних проблем. Кількість таких захворювань постійно зростає. Особливий інтерес представляє автоматизована класифікація захворювань шкіри (як доброякісних, так і злоякісних) на базі зображення ураженої ділянки тіла. Глибокі згорткові нейронні мережі показують потенціал для аналізу категорії дрібнозернистих зображень. Навчальна вибірка даних представлена міжнародною спільнотою з цифрової обробки зображень шкіри ISDIS і містить 13791 зображень уражених ділянок шкіри, зроблених за допомогою дермоскопу [1].

Розглянемо архітектуру згорткової мережі (рис. 1), що була використана для класифікації захворювань. Дана мережа представляє собою послідовність шарів наступних типів: згортковий, агрегувальний та повнозв'язний [2].

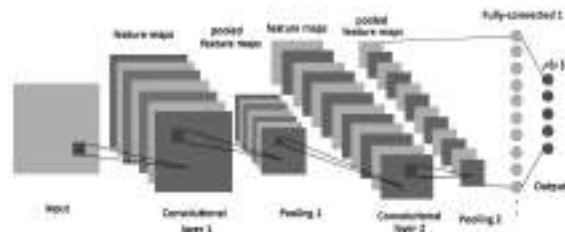


Рис. 1. Архітектура багат шарової згорткової нейронної мережі

Вхідний шар. Містить сирі значення інтенсивностей пікселів. Зображення мають ширину і висоту 150 px та три канали кольору R,G,B.

Шар згортки. Згортка – це лінійне перетворення вхідних даних (рис. 2). Нехай x^l – карта ознак в шарі під номером l , тоді результат двовимірної згортки з ядром розміру $2d + 1$ і матрицею ваги W розміром $(2d + 1) \times (2d + 1)$ на наступному шарі буде наступним:

$$y_{i,j}^l = \sum_{-d \leq a, b \leq d} W_{a,b} x_{i+a, j+b}^l, \quad (1)$$

де $y_{i,j}^l$ – результат згортки на рівні l , а $x_{i,j}^l$ – її вхід, тобто вихід всього попереднього шару. Інакше кажучи, щоб отримати компоненту (i, j) наступного рівня, необхідно використати лінійне перетворення до квадратного вікна попереднього рівня. Для внесення нелінійності в

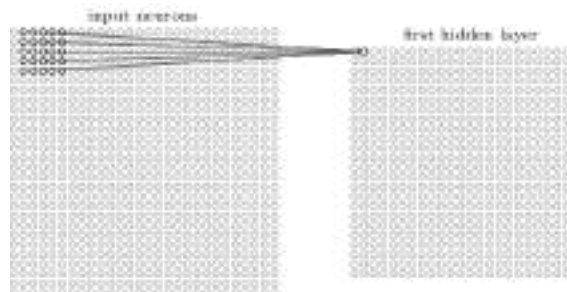


Рис. 2. Візуалізація операції згортки

модель використовується функція активації випрямляч (ReLU): $h(x) = \max(0, x)$.

Агрегувальний шар. Використовується для зменшення розмірності вхідної карти ознак за рахунок їх узагальнення і втрати незначної частини інформації про їх положення. В рамках

даної задачі рекомендовано використовувати метод вибору максимального елемента:

$$x_{i,j}^{l+1} = \max_{-d \leq a \leq d, -d \leq b \leq d} z_{i+a,j+b}^l. \quad (2)$$

Повнозв'язний шар. Після декількох проходжень згортки зображення і шарів підвибірki система перебудовується від конкретної сітки пікселів з високим розширенням до більш абстрактних карт ознак [3]. Ці дані об'єднуються і передаються на звичайну повнозв'язну нейронну мережу (рис. 3):

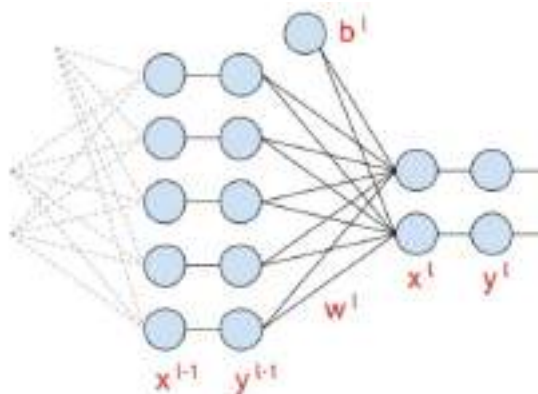


Рис. 3. Візуалізація повнозв'язного шару

$$x_i^l = \sum_{k=0}^m w_{ki}^l y_k^{l-1} + b_i^l, \forall i \in (0, \dots, n), \quad (3)$$

В якості функції активації для кінцевого шару використовується SoftMax, що перетворює вхідний вектор в вектор ймовірностей приналежності об'єкта до відповідного класу.

Оцінка точності отриманої моделі. Розглянемо поняття відносної ентропії (або відстані Кульбака – Лейблера), що представляє собою міру різниці між двома ймовірнісними розподілами P і Q [4]. Як правило, вважається, що P – це істинний розподіл, а Q – його наближення, тоді відстань Кульбака-Лейблера служить оцінкою якості наближення. У випадку, коли P і Q – дискретні випадкові величини на дискретній множині $X = x_1, \dots, x_N$, відстань Кульбака – Лейблера виглядає наступним чином:

$$KL(P||Q) = \sum_i p(x_i) \log \frac{p(x_i)}{q(x_i)}, \quad (4)$$

де $p(x_i)$ і $q(x_i)$ – власне ймовірності результату x_i . В якості функції помилки використовується перехресна ентропія:

$$H(p, q) = E_p[-\log q] = H(p) + D_{KL}(p||q), \quad (5)$$

Для дискретних p і q маємо:

$$H(p, q) = - \sum_{\forall x} p(x) \log q(x). \quad (6)$$

Література. 1. Lu, L., Zheng, Y., Carneiro, G., Yang, L. Deep Learning and Convolutional Neural Networks for Medical Image Computing. 978-3-319-42998-4 изд. Springer International Publishing, 2017. – 326 с. 2. Николенко С., Кадури А., Архангельская Е. Глубокое обучение. Погружение в мир нейронных сетей. 978-5-496-02536-2 изд. СПб. Питер: ООО Издательство “Питер”, 2018. – 480 с. 3. Nikhil Buduma Fundamentals of Deep Learning. Designing next-generation machine intelligence algorithms. 978-1-491-92561-4 O'Reilly Media, 2017. 277 с. 4. Francois Chollet Deep Learning with Python. 978-1617294433 изд. Manning Publications, December 22, 2017. – 384 с.

Бакурова А.В., Терещенко Е.В., Нагулов С.В.

Запорізький національний технічний університет, Запоріжжя, Україна

Логічна класифікація панельних даних

Панельні дані містять одномоментні значення зафіксованих параметрів певних об'єктів в динаміці. Структура панельних даних дозволяє формувати вибірки для дослідження за різними вимірами: за об'єктом, за параметром, за моментом часу. Зазвичай на основі панельних даних будують регресійні моделі для подальшого прогнозування або перевірки гіпотез [1]. В представленій роботі для пошуку прихованих закономірностей, що містять панельні дані, застосовувалися алгоритми логічної класифікації (АЛК) [2]. АЛК будують закономірності, що притаманні класам, у вигляді кон'юнкцій елементарних висловлювань, що робить інтерпретацію результатів досить прозорою. Класифікацію було представлено як розбиття множини об'єктів за певними значеннями одного з параметрів та формулювання у явному вигляді закономірностей на множині параметрів, що притаманні побудованим класам, в аналітичному виді та у вербальному виді на мові предметної галузі [3].

В представленому дослідженні використовувалася одна з ітерацій панельної вибірки Penn World Table версії 5.6, що представляє собою набір макроекономічних панелей розроблених і підтримуваних University of California, Davis і Groningen Growth Development Centre of the University of Groningen, для вимірювання реального ВВП в різних країнах з плином часу [4]. Для обробки було сформовано вибірку без пропусків та дублювання даних, яка містила значення 14-ти макроекономічних параметрів для 98-ми країн за період з 1960-го по 1989-ий роки. Класифікація країн за класами «низький», «середній», «високий» проводилася за параметром ω = «Standard of Living Index». На першому етапі було побудовано загальну панельну класифікацію країн за параметром ω . На другому етапі отримано сімейство класифікацій країн за параметром ω по щорічним перерізам. На третьому етапі проводилася класифікація параметрів за волатильністю за класами «слабко мінливі», «середньо мінливі» і «сильно мінливі». Волатильність визначалася як середньоквадратичне відхилення значень параметрів за період з 1960-го по 1989-ий роки.

В представленій роботі АЛК поєднували стратегії перебору кон'юнкцій «пошук в ширину» (breadth-first search BFS) і «пошук в глибину» (depth-first search DFS) з евристичним критерієм інформативності та критерієм бустінгу [2].

За результатами проведеного дослідження можна зробити висновок про високий потенціал застосування АЛК для видобування знань при панельному представленні інформації. Було побудовано класифікації країн загальна та щорічні. Також були виокремлені класи параметрів за волатильністю. Для кожного класу було побудовано закономірності, які отримали інтерпретацію на мові предметної галузі. Зроблено висновки щодо ефективності застосування різних поєднань стратегій DFS та BFS з евристичним критерієм та критерієм бустінгу в залежності від кількості найбільш інформативних кон'юнкцій.

Практичне застосування результатів дослідження дозволяє виявляти інформативні правила класифікації об'єктів або параметрів, обирати значущі параметри об'єктів у різні хронологічні періоди, передбачати переходи об'єктів між класами, враховуючи динаміку змін їх параметрів. Це буде корисним для попередньої обробки панельних даних при побудові економетричних моделей, що доводить доцільність подальших досліджень.

Література. 1. Лук'яненко І.Г. Сучасні економетричні методи у фінансах // І.Г. Лук'яненко, Ю.О. Городніченко – К.: Літера ЛТД, 2002. — 352 с. 2. Воронцов К.В. Лекции по логическим алгоритмам классификации [Електронний ресурс] / К.В. Воронцов – Режим доступу: <http://www.ccas.ru>. 3. Перепелица В.А. Задачи классификации и формирование знаний / В.А.Перепелица, И.В.Козин, Э.В.Терещенко – Saarbrücken, Germany: LAP LAMBERT Academic Publishing GmbH Co. KG, 2012. – 196 с. 4. Penn World Rable v.5.6 – Режим доступу: <http://dc1.chass.utoronto.ca/pwt56/>.

Буценко Ю.П., Лабжинський В.А.

КПІ ім. Ігоря Сікорського, Київ, Україна

Нейромережіві алгоритми розпізнавання та класифікації надзвичайних ситуацій на об'єктах критичної інфраструктури

Об'єкти критичної інфраструктури характеризуються такими особливостями [1]:

- великою кількістю контрольованих параметрів, причому різномірних як за фізичним походженням, так і за місцем отримання їх значень;
- складною структурою технологічних та інших зв'язків між елементами системи;
- неповнотою інформації про елементи системи та неоднозначністю їх властивостей;
- різноманітністю швидкостей протікання процесів у системі.

Для нормального функціонування системи, що містить вищезгадані елементи, природним є контроль над нею за допомогою нейронної мережі Хопфілда [2]. Такий режим роботи дозволяє розпізнавати стандартні варіанти функціонування системи та відповідним чином реагувати на них. В той же час специфічні особливості мережі Хопфілда обмежують кількість контрольованих нею елементів. Таким чином, вона не може бути використана для контролю стану всієї багатоелементної системи. Елементи, що входять до складу мережі Хопфілда, можна розглядати як опорні.

Розглянемо детальніше ситуації, пов'язані з невизначеностями. Зазвичай виникнення таких ситуацій пов'язане з нестандартними відхиленнями показників роботи деякого елемента системи. В такому випадку паралельно запускають процеси обстеження як фізичного оточення проблемного елемента (сукупності об'єктів, які зазнають впливу його фізичних полів – температури, випромінювання всіх діапазонів, тисків рідин та газів тощо), так і технологічного оточення, зокрема, керуючих пристроїв, контролю, енергозабезпечення тощо.

В обох випадках за умови виявлення елемента з відхиленнями показників від нормальних процедура повторюється. Це дозволяє:

- виявляти всі передумови виявлених відхилень від нормального режиму роботи;
- зробити це на якомога ранній стадії створення передумов до виникнення надзвичайної ситуації;
- повною мірою визначити сукупність елементів системи, на які впливають вищезазначені відхилення;
- мінімізувати вартість ліквідації наслідків таких відхилень.

Вказаний алгоритм ґрунтується на наявності як традиційної мережі датчиків, встановлених на опорних елементах системи, що підлягають опитуванню з певною частотою, так і сукупності мобільних та таких, що допускають перемикання, засобів контролю [3]. Існування останніх дозволяє використання для традиційної мережі, внаслідок обмеження кількості її елементів нейронної мережі Хопфілда. Відповідно, для дослідження сукупностей елементів, що виникають при розгляді фізичного та технологічного оточень проблемних елементів, виявляється доречним застосування нечітких нейронних мереж, які прискорюють прогнозування наслідків відхилень, що спостерігаються [4]. Такий підхід дозволяє поліпшити класифікацію надзвичайних ситуацій до їхнього виникнення за сукупністю передумов.

Література. 1. Зелена книга з питань захисту критичної інфраструктури в Україні: аналітична доповідь / Д.С. Бірюков, С.І. Кондратов, О.І. Насвіт, О.М. Суходоля. – К.: НІСД, 2015. – 35 с. 2. Згуровський М.З. Основы вычислительного интеллекта / М.З. Згуровський, Ю.П. Зайченко. – К.: Наукова думка, 2013. – 408 с. 3. Назаров А.В. Нейросетевые алгоритмы прогнозирования и оптимизации систем / А.В. Назаров, А.И. Лоскутов. – СПб.: Наука и техника, 2003. – 384 с. 4. Шматко А.В. Анализ эффективности нечетких нейронных сетей в задачах прогнозирования чрезвычайных ситуаций техногенного характера / А.В. Шматко, Е.А. Голубничая // Системи обробки інформації. – 2007. – № 1. – С. 147–151.

Галушко М.О.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Класифікація зображень за допомогою нейронних мереж на GPU

Вперше нейронні мережі привернули загальну увагу в 2012 році, коли Алекс Крижевський завдяки їм виграв конкурс ImageNet, знизивши рекорд помилок класифікації з 26% до 15%, що тоді стало проривом. На даний момент нейронні мережі показують дуже хороші результати саме для задач аналізу картинок і текстів. Іноді навіть моделі машинного навчання справляються із поставленими задачами краще, ніж люди [1].

Завдання класифікації зображень – це прийом початкового зображення і виведення його класу (кішка, собака і т.д.) або групи ймовірних класів, яка найкраще характеризує зображення. Для людей це один з перших навичок, який вони починають освоювати з народження.

У роботі вирішувалась задача категоризації предметів для дому. Так як для великих поставок категоризація товару може забирати багато часу у працівників, автоматизація цього процесу значно спростить їхню роботу. Також це може бути корисно, якщо на сайті не дуже очевидна категоризація (з категоріями і підкатегоріями), тоді користувачу, який захоче продати/купити товар, буде легше розмістити оголошення/купити товар по фото.

Для виконання задачі класифікації зображень за допомогою нейронних мереж використовувались згорткові мережі. Згорткові мережі являють собою варіацію архітектури багатошарового перцептрона, і включають в себе згорткові шари, шари підвибірки (субдискретизація), і повнозв'язні шари. Архітектура згорткової мережі використовує переваги двовимірної структури вхідних даних – зображень за допомогою методу локальної зв'язності, обмежуючи кількість зв'язків між нейронами прихованого згорткового шару і вхідними даними. Конкретно, кожен нейрон прихованого шару пов'язаний тільки з обмеженою локальною ділянкою зображення [2].



Рис. 1. Архітектура згорткової нейронної мережі

Ціллю роботи є створення такої моделі, що виконувала категоризацію предметів із точністю, яка нас влаштовує. У ході роботи важливо проаналізувати підходи до вибору архітектури моделей, методів та технологій, щоб вибрати такий, який найефективніше б підійшов для категоризації предметів на 128 класів по фото. Не менш істотним у даному контексті є глибокий аналіз даних та існуючих підходів до виконання такого типу завдань їх особливостей та можливостей використання у реальних системах.

Результатом роботи є протестована модель, що класифікує предмети по фото із прийнятною точністю (> 75% правильних передбачень) та здійснений аналіз її ефективності і надійності. Для виконання роботи використовувався фреймворк Keras, остаточної архітектури Xception з натренованими вагами на ImageNet і додано два повнозв'язних шари. Точність класифікації отриманої моделі – 82.5%.

Література. 1. Goodfellow I. Deep Learning / Goodfellow I., Bengio Y. and Courville A.; Cambridge MA: MIT Press [2017] – 777 pages. (дата звернення: 20.02.2018). 2. Ruder S. (2016) An overview of gradient descent optimisation algorithms. URL: <https://arxiv.org/abs/1609.04747> (дата звернення: 20.02.2018).

Гусак О.В.^{1,2} Семендяк Є.С.^{1,2}

¹КПІ ім. Ігоря Сікорського, ІТС, Київ, Україна; ²Факультет інформатики Дрезденського технічного університету, Дрезден, Німеччина

Адаптивний пошук гіперпараметрів за допомогою підходів машинного навчання

У даній роботі сформовано загальний підхід до скорочення порівняльного аналізу. Основна ідея полягає в тому, щоб почати з обмеженої кількості вимірювань і шляхом побудови моделі поліноміальної регресії та прогнозування еталонного рішення визначити необхідність проведення додаткових вимірювань на основі зростаючої множини проведених експериментів.

Вступ. Порівняльний аналіз являється необхідною частиною емпіричних досліджень, таких як вибір гіперпараметрів машинного навчання. За допомогою тестування і експериментального аналізу визначаються зв'язки множинних факторів з однією чи кількома досліджуваними змінними. Основним недоліком досліджень, орієнтованих на порівняння, є високі затрати, необхідні для їх завершення.

Мета. Типовим підходом є перетворення комплексного фактора на кілька псевдо факторів, кількість яких можна зменшити пізніше. Проте таке скорочення є грубим і неефективним. Виконуючи менше експериментів, отримується компромісне рішення якості результатів якого, необхідно максимізувати. В роботі досліджуються питання: 1. Чи можна ефективно зменшити кількість експериментів і при цьому максимально наблизитися до необхідного оптимуму. 2. Як динамічно оцінювати експеримент щоб прийняти рішення про доцільність його подальшого розрахунку і готовності регресійної моделі.

Опис алгоритму. Основна ідея підходу полягає в тому, щоб обробляти конфігурації як дані та використовувати методи моделювання даних для виявлення важливих конфігурацій. Ідея скорочення розміру пошуку – це ітеративне додавання конфігурацій, їх вимірювання та оцінка необхідності подальшого порівняльного аналізу шляхом прогнозування вимірювань для інших конфігурацій. Для передбачення оптимальної конфігурації на кожній ітерації використовується функція прийняття рішень, побудована на результатах попередніх вимірів. Для конфігурації моделі поліноміальної регресії, що являє собою функцію прийняття рішень, використовується підмножина вибраних параметрів, розділена на набори даних для навчання та тестування моделі. Модель містить повну схему з фактичними значеннями для вже згрупованих конфігурацій та прогнозованих значень. Після вимірювання кожної нової конфігурації наявна множинна, регресійна модель покращується. Виміряні дані розподіляємо між наборами для навчання та тестування. При розділенні даних, розмір тренувального набору потрібно мінімізувати, при цьому отримуючи необхідну точність цільової моделі. За точністю розуміється коефіцієнт визначення прогнозу об'єкту – алгоритму навчання.

Реалізація. Проводиться оцінка запропонованого комплексного підходу, застосовуючи його для пошуку оптимальних гіперпараметрів за мінімальний час. Описане рішення дає можливість скоротити кількість проведених експериментів на 70% при втратах в точності до 5% в порівнянні з цільовою ідеальною конфігурацією. Для динамічної оцінки експерименту про доцільність подальшого розрахунку, проводилося порівняння вже затраченого і передбаченого часу в порівнянні з попередніми експериментами. При аномальній поведінці експеримент скасовується. Для цільової моделі машинного навчання і побудови регресійної моделі застосовувалася бібліотека scikit-learn.

Висновки. Апробація підходу показала ефективність використання регресійної моделі для зменшення порівняльного аналізу в побудові моделі машинного навчання. Мінімізовано обсяг тестування гіперпараметрів машинного навчання, одночасно максимізуючи якість результатів.

Література. 1. Sebastian G. Benchmark Reduction via Adaptive Instance Selection / P. Dmytro, G. Sebastian // International Workshop on Green and Sustainable Software (GREENS'18). – 2018.

Дмитренко О.О.¹, Ланде Д.В.²

¹КПІ ім. Ігоря Сікорського, ФТІ, Київ, Україна; ²Інститут проблем реєстрації інформації НАН України, Київ, Україна

Алгоритм пошуку оптимального сценарію впливу вершин у когнітивних картах заснований на принципі Парето

Вступ. Одним із напрямків моделювання складних мереж є створення когнітивних карт, їх опис та аналіз [1]. Під час дослідження когнітивних карт виникає задача обґрунтування вибору оптимального сценарію впливу одного вузла на інший серед доступних сценаріїв впливу. Тож важливим кроком є введення критеріїв оптимальності впливу [2].

Критерії оптимальності впливу. В роботі представлені два критерії оптимальності впливу C_1 та C_2 : сила впливу та швидкість реалізації сценарію.

Вплив вершини u_i на u_j ($i, j = 1, 2, \dots, n$, де n – кількість вершин у когнітивній карті) називається найсильнішим, якщо вплив на кінцеву вершину u_j початкової вершини u_i на k -му простому шляху є найбільшим за абсолютною величиною серед всіх впливів на інших простих шляхах, що сполучають вершини u_i та u_j .

Вплив вершини u_i на u_j вважається найшвидшим в реалізації, якщо він здійснюється по найкоротшому шляху. Швидкість реалізації k -го сценарію $(u_i, u_j)_k$ визначається відповідно до кількості ребер $(m - 1)$, що сполучають вершини u_i та u_j у k -му простому шляху (де m – кількість вершин, що входять до k -го шляху).

Для отримання єдиного розв'язку поставленої в даній роботі задачі застосовується багатокритеріальна оцінка отриманих сценаріїв за принципом Парето [3]. В загальному випадку, коли множина Парето містить більше одного елемента, для визначення оптимального сценарію в цій роботі запропоновано такий алгоритм:

1. Спочатку визначається НСК – найменше спільне кратне (англ. LCM – least common multiple) оцінок критерію C_2 для всіх елементів множини Парето. Якщо вважати, що C_2 – це час, за який реалізується відповідний сценарій, то НСК всіх значень – це найменший час, за який реалізується цілочисельна кількість кожного із сценаріїв, що входять до множини Парето. Тобто за один і той же час $LCM(c_2^{(1)}, \dots, c_2^{(d)})$ кількість реалізацій різних сценаріїв, що входять до множини Парето, буде відповідно різною $(a^{(1)}, a^{(2)}, \dots, a^{(d)})$:

$$a^{(k)} = \frac{LCM(c_2^{(1)}, \dots, c_2^{(d)})}{c_2^{(k)}},$$

- де $c_2^{(k)}$ – оцінка критерію C_2 для k -го сценарію; d – кількість елементів множини Парето.
2. Далі для кожного із сценаріїв, що входять до множини Парето, значення їх оцінок за критерієм C_1 $\{c_1^{(1)}, c_1^{(2)}, \dots, c_1^{(d)}\}$ домножаються на відповідне значення $(a^{(1)}, a^{(2)}, \dots, a^{(d)})$. Тобто визначається, яким буде загальний вплив вершини u_i на u_j за час $LCM(c_2^{(1)}, \dots, c_2^{(d)})$ за k -м сценарієм.
3. Для того, щоб визначити оптимальний сценарій, необхідно знайти найбільше значення добутку $c_1^{(k)} a^{(k)}$ визначеного на етапі 2):

$$\max_k (c_1^{(k)} a^{(k)}),$$

де $k \in \{1, \dots, d\}$ – номер оптимального сценарію, якому відповідає максимальний добуток $c_1^{(k)} a^{(k)}$.

В результаті обґрунтування вибору оптимального за введеними критеріями C_1 та C_2 сценарію впливу для кожної пари вершин (u_i, u_j) зваженого оґрафа ($i, j = 1, 2, \dots, n$, де n – кількість вершин зваженого оґрафа), можна побудувати матрицю впливу Z , що складатиметься з елементів z_{ij} , та матрицю T , що складатиметься з елементів t_{ij} .

Загальний вплив z_{ij} – це вплив вершини u_i на вершину u_j за оптимальним сценарієм (тобто значення оцінки критерію C_1 за оптимальним сценарієм). Загальний час t_{ij} – це час, що

необхідний для реалізації оптимального сценарію впливу вершини u_i на вершину u_j (значення оцінки критерію C_2 за оптимальним сценарієм). Якщо u_j недосяжна із вершини u_i , то $z_{ij} = 0$ та $t_{ij} = 0$.

Для того, щоб визначити, яким буде вплив кожної із вершин на інші при $t \rightarrow \infty$, необхідно виконати наступне. Спочатку потрібно визначити z_{ij}^1 – вплив кожної із вершин при $t = 1$. Враховуючи час необхідний для реалізації кожного із сценаріїв для матриці впливу Z шляхом ділення кожного її ненульового елемента z_{ij} ($z_{ij} \neq 0$) на відповідний елемент t_{ij} матриці T , буде отримано матрицю Z_1 , елементами якої є

$$z_{ij}^1 = \begin{cases} \frac{z_{ij}^t}{t_{ij}}, & z_{ij} \neq 0, \\ 0, & z_{ij} = 0. \end{cases}$$

Далі вплив, який здійснює кожна вершина за час t , обраховується як $z_{ij}^{t+1} = z_{ij}^t + z_{ij}^1$. На кожному кроці $t = 2, 3, 4, \dots$ здійснюється процес нормування:

$$z_{ij}^t = \frac{z_{ij}^t}{\sum_{k=1}^n \sum_{l=1}^n z_{kl}^t}.$$

Приклад. В роботі [1] розглядається імпульсно-стійкий зважений оргграф (рис. 1).

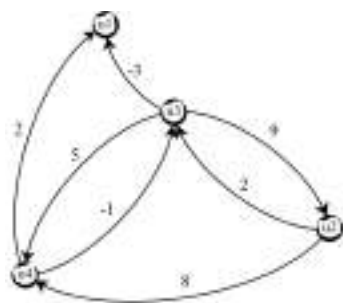


Рис. 1. Зважений оргграф

Табл. 1. Приклад таблиці Парето (Сценарій 2 є оптимальним за введеними критеріями C_1 та C_2)

Сценарій $(u_3, u_4)_k$	1	2
C_1	6,92	5
C_2	2	1

Визначивши оптимальний за введеними критеріями C_1 та C_2 сценарій впливу для кожної пари вершин (u_i, u_j) ($i, j = 1, 2, 3, 4$) когнітивної карти представленої на рисунку 1, можна побудувати матрицю впливу Z , матрицю T та матрицю Z_t :

$$Z = \begin{pmatrix} 0 & 0 & 0 & 0 \\ 1,66 & 0 & 2 & 8 \\ -3 & 9 & 0 & 5 \\ 2 & -1,79 & -1 & 0 \end{pmatrix}, T = \begin{pmatrix} 0 & 0 & 0 & 0 \\ 2 & 0 & 1 & 1 \\ 1 & 1 & 0 & 1 \\ 1 & 2 & 1 & 0 \end{pmatrix}, Z_t = \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0,038 & 0 & 0,091 & 0,365 \\ -0,137 & 0,41 & 0 & 0,228 \\ 0,091 & -0,04 & -0,046 & 0 \end{pmatrix}.$$

Висновок. Отже, в даній роботі запропоновано алгоритм пошуку оптимального сценарію впливу вершин у когнітивних картах. Використовуючи принцип Парето, відповідно до введених критеріїв (сила впливу та швидкість реалізації сценарію впливу), на конкретному прикладі було обгрунтовано вибір оптимального сценарію впливу.

Література. 1. Roberts F. S. Discrete Mathematical Models with Applications to Social, Biological, and Envi-ronmental Problems / Fred Roberts. – New Jersey: Rutgers University, Prentice-Hall Inc., 1976. 2. Снарский А.А. Ранжирование понятий, извлекаемых из потоков сетевых новостей / Снарский А.А., Ланде Д.В., Зоринец Д.И. // Информационные технологии и безопасность. Материалы XVI Международной научно-практической конференции ИТБ-2016. – К.: ИПРИ НАН Украины, 2017. – С. 130–131. 3. He Z. Fuzzy-based Pareto optimality for many-objective evolutionary algorithms / Z. He, G.G. Yen, J. Zhang. // Transactions on Evolutionary Computation. – 2014. – №18. – С. 269–285.

Доманецька І.М.¹, Хроленко Я.О.²

¹Київський національний університет ім. Тараса Шевченка, Київ, Україна; ²КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Нейромережеві технології у задачах автоматизації контролю знань для on-line сервісів відкритої освіти

У всьому світі набирають популярності портали, націлені на дистанційне навчання – від короткострокових тренінгів до вищої освіти. Динамічний розвиток ринку on-line освіти посилює існуючу на ньому конкуренцію. Але в порівнянні з традиційним підходом такі незаперечні переваги on-line навчання, як доступність і зручність, помітно знижуються через недоліки, основні з яких полягають в наступному:

- складність об'єктивного контролю знань (індивідуальну перевірку завдань замінюють різні форми тестування);
- складність об'єктивного оцінювання якості навчання.

Один із шляхів усунення цих недоліків базується на останніх розробках в області штучного інтелекту, і зокрема, нейромережевих технологіях.

Дана робота присвячена питанням інтелектуалізації методів автоматичного контролю знань. Завдання контролю знань розбивається на наступні підзадачі [1]:

1. визначення ступеня вірності відповіді на питання;
2. визначення рівня засвоєння одиниці перевірки за ступенями вірності відповідей на питання;
3. визначення рівня засвоєння вищестоящої по ієрархії одиниці перевірки за рівнями засвоєння нижчих одиниць перевірки (дисципліна, розділ, тема, питання).

Найцікавішою з точки зору авторів є задача визначення ступеню вірності відповіді на відкриті тестові питання. Питання відкритого типу характерні тим, що для отримання відповіді на них студент повинен ввести символічний рядок, що представляє власну відповідь. Звичайно такі відповіді вводяться на природній для людини мові, максимально наближеній до розмовної. Як правило, тягар перевірки завдань відкритого типу лягає на плечі викладача. Більшість сучасних систем тестування в основному ґрунтуються на завданнях закритого типу, автоматична перевірка завдань відкритого типу зводиться до перевірки повного збігу з одним із можливих еталонних варіантів відповіді.

Адекватна автоматична перевірка природномовних відповідей є важким завданням. Шаблони відповідей у вигляді стандартних виразів не в змозі прийняти до уваги різноманіття властиве рідній мові. Крім того, потрібно автоматичне виявлення випадкових помилок (друкарських помилок) і помилок правопису.

Тому актуальним є завдання створення/розвиток інтелектуальних механізмів, які б дозволили вирішити задачу якісної автоматизованої обробки природномовних текстових відповідей, і можуть бути покладені в основу інтелектуальних систем оцінювання знань.

На сьогоднішній день існує декілька підходів до оцінки правильності відповідей на відкриті тестові питання на базі нейромережевих технологій опрацювання природномовних текстів. Заслужують уваги роботи Т.В. Батури, Olschimke M., Moon B., Arnaudo P. [2], що досліджують використання згорткових мереж, рекурсивних (мережа Елмана), мереж, що працюють із символічними поданнями слів або змішаними поданнями, динамічних мереж, Міцеля А.А., Погуди А.А., що пропонують для проведення семантичного аналізу відповіді відкритої форми використовувати самоорганізаційні карти (SOM – self-organizing map), роботи Д. Моріфуджи, Т. Іннуї, Д.Є. Шкуліна, А.А. Солодовнікова, що працюють у напрямі використання нейронних мереж для вирішення задач морфологічного і синтаксичного аналізу тексту і задачі аналізу словозмін.

Література. 1. Титов А.М. Нейросетевые алгоритмы компьютерного контроля знаний: разработка и исследование, <http://www.dissercat.com/content/ixzz58tSdKc1a>. 2. International Journal of Artificial Intelligence in Education (IJAIED), <http://aied.inf.ed.ac.uk>.

Ерохин Г.В.

КПИ им. Игоря Сикорского, ИПСА, Киев, Украина

Прогнозирование курса криптовалют на основе анализа тональности новостей и записей в социальных сетях

Криптовалюты – это новая концепция в глобальной экономике. Они существуют менее десяти лет, и они уже привлекли большое внимание. Мы можем рассматривать криптовалюту как цифровой носитель обмена, основанный на принципах криптографии, который позволяет осуществлять безопасные, децентрализованные и распределенные экономические транзакции. Наиболее известной криптовалютой является *Bitcoin*, он же и появился первым в 2009 году. [1]

В последнее время (начиная с 2013 года) появилась еще одна область использования криптовалют. В 2009 году, когда *Bitcoin* только появился, никто о нем не знал и не пользовался. Поэтому его стоимость была равна \$0. Даже за следующий год стоимость *Bitcoin* не превышала \$0.39. Однако, в течение 2013 стоимость 1 *Bitcoin* превысила \$1000, а в течение 2017 его стоимость иногда приближалась к \$20000. Очевидно, что для многих людей это стало несколько похожим на рынок акций. Однако, в конце 2017 года стоимость снизилась до почти \$13000, что означает, что волатильность *Bitcoin* является очень высокой. Именно поэтому задача прогнозирования курса криптовалют является очень актуальной.

На то, что у курса настолько высокая волатильность, есть несколько причин:

- Первый фактор заключается в том, что криптовалюта имеют меньшие размеры рынка по сравнению с установленными формами валюты. Это означает, что даже небольшие движения криптовалюты могут иметь видимое влияние на ее цену. Более того, распределение богатства в сфере криптовалют является довольно необычным, люди с большими долями в криптовалюте имеют непропорционально большую власть над её ценами.
- Второй фактор заключается в том, что общественное восприятие криптовалют является достаточно дихотомическим. Например, подорожание во втором квартале 2017 может быть связано с рядом факторов (например, увеличение интереса во всех странах Азии, увеличение принятия предприятиями), которые увеличили позитивные настроения в этом пространстве. С другой стороны, негативные изменения в законности криптовалют, а также поломки в системе могут очень легко уничтожить общественный интерес к ним.

Есть мнение, что большинство людей, которые покупают или продают криптовалюты ориентируются в первую очередь не на числовое значение курса, а на реакцию и популярность объекта вложения в социальных сетях. Можно сказать, что большинство этих людей имеют доступ к интернету, являются активными пользователями социальных медиа, читают форумы, блоги и сайты. Поэтому их решения должны основываться на том, что происходит в интернете. Так как курс криптовалют очевидно коррелирует с интересом вкладчиков, а их интерес напрямую отображается в новостях и социальных сетях, имеется возможность прогнозировать курс криптовалют с помощью анализа тональности вышеупомянутых записей.

Предлагается следующий алгоритм прогнозирования:

- Обучение классификатора тональности на основе определенной маркированной выборки новостей.
- Сбор информации о курсе криптовалют и новостей о них за определенный период из открытых API.
- При идентификации корреляции между изменением цен на *Bitcoin* и изменением тональности в новостях, учитываются два временных аспекта: длина частоты и смещение. Краткосрочные прогнозы будут проводиться на определенных интервалах времени. Соответственно, при длине интервала, l и смещении равном n , если событие (под событием подразумевается изменение тональности) произошло во время X , то прогнозируется курс криптовалюты во время $X + n * l$.

Литература. 1. Vejačka, Martin. (2014). Basic Aspects of Cryptocurrencies. Journal of Economy, Business and Financing. 2. P. 75–83.

Жданова О.Г., Попенко В.Д.

КПІ ім. Ігоря Сікорського, ФІОТ, Київ, Україна

Економічна модель прийняття рішень щодо народження дітей у сім'ї

Розглянуті моделі прийняття рішень щодо кількості дітей в сім'ях з двох і трьох поколінь з урахуванням продуктивності праці та жорсткого бюджетного обмеження.

Вступ. Питання зміни кількості населення є визначальним для перспектив людства. Вчені тривалий час будують математичні моделі демографічних процесів. Основи економічного аналізу репродуктивної поведінки заклали Г.Беккер [1,2], Р.Акофф [3]. Для багатьох українських родин бюджетне обмеження є жорстким і реально впливає на прийняття важливих рішень. Зокрема, на репродуктивну поведінку впливає те, скільки дітей зможе прогодувати родина і скільки дітей потрібно для підтримки в старості. Розглянемо математичні моделі планування сім'ї з урахуванням її доходів і витрат, починаючи з найпростішої.

Модель 1. Сім'я з двох поколінь, батьки утримують неповнолітніх дітей, діти утримують батьків у старості. Маємо сім'ю з двох поколінь, пенсійне забезпечення відсутнє. Різниця між зайнятістю чоловіків і жінок відсутня.

Нехай $p \geq 1$ – кількість членів сім'ї, включаючи себе, яку може забезпечити один працюючий (аналог продуктивності праці); d – кількість дітей у сім'ї; c – витрати на утримання дитини як частка витрат на утримання дорослого (далі – відносні витрати на дитину).

Умова підтримки батьків дітьми у старості: кількість людей, яких можна забезпечити на заробітки дітей, не менше кількості членів сім'ї:

$$dp \geq d + 2. \quad (1)$$

Умова утримання батьками дітей до працездатного віку:

$$2p \geq dc + 2. \quad (2)$$

Разом (1) та (2) дають таку систему нерівностей:

$$\begin{cases} d \leq 2\frac{p-1}{c}, \\ d \geq \frac{2}{p-1}. \end{cases} \quad (3)$$

Причинно-наслідковий зв'язок параметрів моделі є наступним: продуктивність праці визначає допустиму кількість дітей, а остання визначає частку витрат на дитину від витрат на дорослого, тобто $p \rightarrow d \rightarrow c$. Тому побудуємо залежність c від p . З системи нерівностей (3) випливає справедливості

$$\frac{2}{p-1} \leq \frac{2(p-1)}{c}, \quad (4)$$

звідки отримуємо:

$$c \leq (p-1)^2. \quad (5)$$

На рис. 1 показано графік залежності максимально можливого значення відносних витрат c на дитину при заданому p і відповідна множина допустимих значень «відносні витрати на дитину – показник продуктивності праці» у моделі 1.

Модель 1 дозволяє зробити такі висновки:

- точка A ($p = 1$) відповідає стану крайніх злиднів і неможливості прогодувати дітей, у цьому випадку сім'я руйнується;
- точка B відповідає стану бідності, коли на дітей не витрачається майже нічого, і тим самим бідність відновлюється в наступному поколінні; така сім'я потребує чотирьох дітей;
- точка C відповідає сталому стану, коли в наступному поколінні відтворюється стан попереднього – зокрема, двоє батьків народжують двох дітей; ця точка є межею між розвитком і деградацією;
- точка D відповідає стану забезпеченості при відносно високій продуктивності праці, адже один працюючий може крім себе утримувати двох дорослих. Така сім'я здатна утримувати

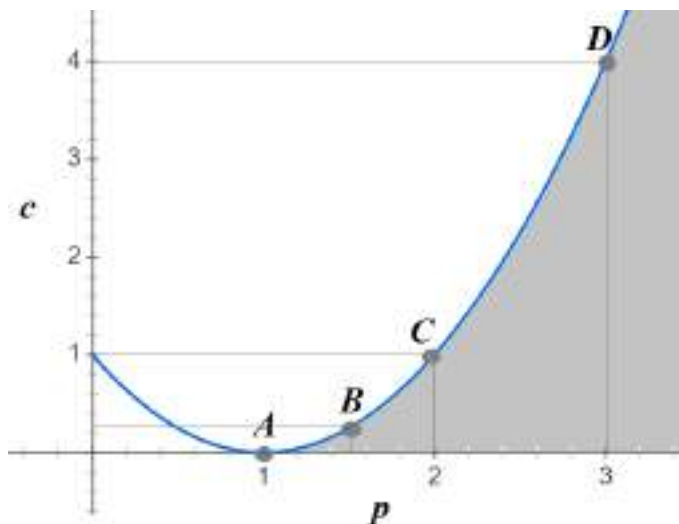


Рис. 1. Множина допустимих значень «відносні витрати на дитину – показник продуктивності праці» у моделі 1; точка $A : p = 1$; точка $B : p = 1.5; d = 4; c = 0.25$; точка $C : p = 2; d = 2; c = 1$; точка $D : p = 3; d = 1; c = 4$.

лише одну дитину з максимально можливою долею витрат відносно утримання дорослого $c = 4$. Але в порівнянні з точкою B у батьків є вибір: вони можуть завести двох дітей, встановивши для них $c = 2$.

Модель 2. Сім'я з трьох поколінь. В цій моделі передбачається, що працюючі батьки утримують дітей, а крім того, беруть участь в утриманні своїх батьків. Отже, заробіток батьків, який забезпечує утримання $2p$ осіб, іде на дітей (dc), самих себе (2), а також на своїх батьків, тобто дідусів і бабусь своїх дітей ($\frac{4}{2d} = \frac{2}{d}$). При отриманні останньої формули ми виходили з того, що дідусів і бабусь дві пари і кожна пара має d дітей, які всі беруть участь у їх утриманні. Отже,

$$dc + 2 + \frac{2}{d} \leq 2p, \quad (6)$$

звідки

$$cd^2 - 2(p-1)d + 2 \leq 0. \quad (7)$$

Дискримінант рівняння

$$cd^2 - 2(p-1)d + 2 = 0 \quad (8)$$

має значення $D = 4(p-1)^2 - 8c = 4((p-1)^2 - 2c)$. Його дійсні корені такі:

$$d_{1,2} = \frac{(p-1) \pm \sqrt{(p-1)^2 - 2c}}{c}. \quad (9)$$

Якщо $D \geq 0$, то парабола (8) перетинається з горизонтальною віссю і рівняння має дійсні корені, а якщо в замкненому інтервалі між коренями є хоча б одне натуральне число (адже рішення сім'ї щодо кількості дітей має бути натуральним числом), то умова (6) виживання сім'ї з дітьми виконується. На рис. 2 для $p = 3$ та $c = 1$ показані допустимі значення кількості дітей, при яких ця умова виконується.

Природньо очікувати, що модель 2 дає подружжю менший вибір, ніж модель 1. Зведемо результати моделей 1 і 2 при деяких однакових комбінаціях вхідних даних в таблицю 1. Як бачимо, необхідність мати багато дітей спостерігається в «найбіднішій» комбінації параметрів (перші рядки при відповідних c), хоча вона не є можливою в моделі 3. «Найзаможніша» комбінація параметрів дає батькам певний вибір (від 1 до 3 дітей), але цей вибір звужується,

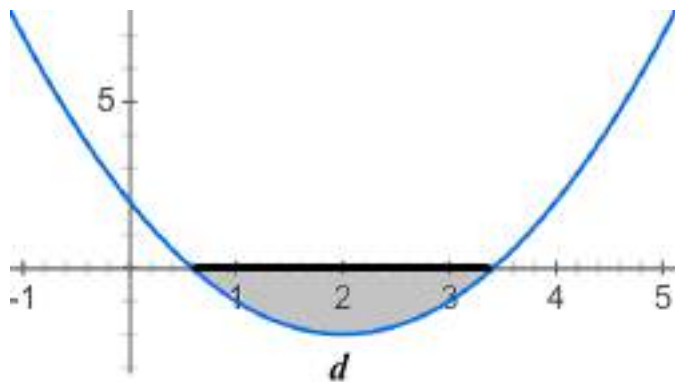


Рис. 2. Допустимий інтервал значень кількості дітей при $p = 3$ та $c = 1$, за яких, згідно моделі 2, виконується умова виживання сім'ї

якщо батьки намагаються вкласти в розвиток дитини якнайбільше. Тоді мала кількість дітей є економічною необхідністю навіть у відносно заможних сім'ях.

Табл. 1. Результати моделей 1 і 2 при різних комбінаціях вхідних даних

c	p	Модель 1, умова утримання дітей	Модель 1, умова утримання батьків	Модель 2
0.25	1.5	$d \leq 4$	$d \geq 4$	$D < 0$
0.25	2	$d \leq 8$	$d \geq 2$	$1.17 \leq d \leq 6.83$
0.25	3	$d \leq 16$	$d \geq 1$	$0.52 \leq d \leq 15.5$
0.5	1.5	$d \leq 2$	$d \geq 4$	$D < 0$
0.5	2	$d \leq 4$	$d \geq 2$	$d = 2$
0.5	3	$d \leq 8$	$d \geq 1$	$0.54 \leq d \leq 7.46$
0.75	1.5	$d \leq 1.33$	$d \geq 4$	$D < 0$
0.75	2	$d \leq 2.67$	$d \geq 2$	$D < 0$
0.75	3	$d \leq 5.33$	$d \geq 1$	$0.56 \leq d \leq 4.77$
1	1.5	$d \leq 1$	$d \geq 4$	$D < 0$
1	2	$d \leq 2$	$d \geq 2$	$D < 0$
1	3	$d \leq 4$	$d \geq 1$	$0.59 \leq d \leq 3.41$

Висновок. Незважаючи на спрощений опис реальних процесів, розглянуті моделі уявляють важливі залежності:

- визначальним параметром для всіх інших характеристик моделей є продуктивність праці, її достатньо високі значення забезпечують певний вибір батькам і покращання добробуту сім'ї (втім, це очевидно);
- за наявності бажання забезпечувати максимальний розвиток дітей їх невелика кількість є економічно необхідною, так само, як необхідна їх значна кількість в разі невеликої продуктивності праці.

Література. 1. Becker, Gary S. Family Economics and Macro Behavior // American Economic Review. – 1988. №78 (1): р. 1–13. 2. Г.С. Беккер. Человеческое поведение. Экономический поход, Москва, ГУ ВШЭ, 2003. 3. Акофф Р. Акофф о менеджменте [Текст] / Рассел Л. Акофф; СПб.: Питер, 2002.– 448с.: ил. – ISBN 5-318-00286-2.

Забелин С.И.

КПИ им. Игоря Сикорского, ИПСА, Киев, Украина

Интеллектуальный анализ вулканических газов с помощью интеллектуальных методов анализа данных

Вступление. Длительное время изучение вулканов вызывает всемирный научный интерес [1]. На острове Гавайи расположен один из самых активных и известных в мире вулканов – Килауэа. Несмотря на подробные отчеты об его активности с 1823 года, было проведено небольшое число исследований в области анализа вулканических газов с использованием современных методов интеллектуального анализа данных. В данной работе показаны результаты такого анализа, с помощью полиномиальных нейронных сетей на основе воздушных эмиссий различных газов с 1970 года и до 2013 года.

Постановка задачи. Были выбраны следующих параметры, что были поданы на вход нейронной сети:

1. С – температурная аномалия (С/месяц).
2. R – осадки (мм/месяц).
3. CO₂ – двуокись углерода (ppm).
4. SO₂ – диоксид серы (ppm).
5. MgO – оксид магния (процентное содержание, относительно общей массы вулканических выбросов).
6. CaO – оксид кальция (процентное содержание, относительно общей массы вулканических выбросов).
7. TiO₂ – оксид титана (процентное содержание, относительно общей массы вулканических выбросов).

Все выше представленные параметры имеют периодичность наблюдений в 1 месяц и охватывают в общей сложности период в 43 года активности вулкана Килауэа. Выходными переменными были выбраны следующие параметры, которые являются смещенными на один месяц вперед относительно входных параметров:

1. CO₂ – двуокись углерода (ppm).
2. SO₂ – диоксид серы (ppm).

Результаты исследования. Было произведено обучение нейронной сети Больцмана на обучающей выборке и проведена последующая кросс-валидация. Была достигнута высокая степень точности прогнозирования (91% точности). Были обнаружены скрытые зависимости выбросов диоксида серы от содержания оксида магния в окрестности вулкана.

Выводы. Поскольку полевые работы в районах активного вулканизма могут быть дорогостоящими, трудоемкими, логистически сложными и опасными для тех, кто участвует, представленный подход к анализу и прогнозированию выбрасываемых газов имеет потенциал для анализа и прогнозирования вулканических процессов, связанных с новыми извержениями.

Литература. 1. Holcomb, R.T. (1981). Kilauea Volcano, Hawaii: Chronology and morphology of the surficial lava flows (USGS Numbered Series). Reston, VA: U.S. Geological Survey. 2. USGS: Volcano Hazards Program. (n.d.). Retrieved March 24, 2018, from <https://volcanoes.usgs.gov/vhp/gas.html>. 3. Hinton, G.E. (2006). Reducing the Dimensionality of Data with Neural Networks. Science, 313(5786), 504–507. doi:10.1126/science.1127647. 4. Connor, C.B., Stoiber, R.E., Malinconico, L.L. (1988). Variation in sulfur dioxide emissions related to earth tides, Halemaumau Crater, Kilauea Volcano, Hawaii. Journal of Geophysical Research: Solid Earth, 93(B12), 14867–14871. doi:10.1029/jb093ib12p14867. 5. Sheth, H. (2003). The active lava flows of Kilauea volcano, Hawaii. Resonance, 8(6), 24–33. doi:10.1007/bf02837866.

Зайченко Е.Ю., Ови Нафас Агаи аз Гамии

КПИ им. Игоря Сикорского, ИПСА, Киев, Украина

Анализ и оптимизация компьютерных сетей нового поколения в условиях неопределенности

Важной задачей проектирования сетей нового поколения является задача анализа и оптимизации функциональных характеристик компьютерных сетей при ограничениях на показатели качества обслуживания. Постановка и метод решения данной задачи рассмотрены и исследованы в работе [1]. В этих работах предполагалось, что исходные данные проектирования, в частности стоимости каналов связи и интенсивности потоков передаваемых сообщений (пакетов/с) известны достоверно. Вместе с тем на практике эти величины известны проектировщикам на стадии проектирования лишь приблизительно. Поэтому задача оптимизации характеристик сетей и проектирования должна решаться в условиях неполноты исходной информации и неопределенности.

Одним из конструктивных подходов решения оптимизационных задач в условиях неопределенности является использование аппарата нечетких множеств и нечеткой логики, разработанного в работах Лотфи Заде. Дальнейшее развитие этот подход получил в работе [2], в которой были сформулированы задачи нечеткого математического программирования и разработан метод их решения при известных функциях принадлежности нечетких параметров. Именно такой подход предлагается использовать в рассматриваемой задаче. Рассмотрим ее постановку.

Постановка задачи. Пусть имеется сеть MPLS, которая описывается орграфом $G = \{X, E\}$, где $X = \{x_j\}$ – множество узлов сети (УС), $E = \{(r, s)\}$ – множество каналов связи (КС); μ_{rs} – пропускные способности КС. Допустим, что в сети передается K классов потоков ($K = \overline{1, 8}$) (CoS) в соответствии с матрицами требований $H(k) = \|h_{ij}(k)\|$ $i = \overline{1, N}$, $j = \overline{1, N}$ (Мбит/с). Для каждого класса k введен показатель качества (QoS) в виде величины средней задержки $T_{cp,k}$, выражение для которой приведено в [1].

Пусть компоненты матрицы требований $h_{ij}(k)$ являются нечеткими числами с функциями принадлежности $\mu(h_{ij})$, а стоимости каналов связи c_{rs} – нечеткими числами с ФП $\gamma(c_{rs})$. Требуется найти такие пропускные способности КС и распределения потоков всех классов $f_{rs}^{(k)}$, при которых общая стоимость сети будет минимальной, а средние задержки пакетов всех классов не будут превышать заданные значения $T_{cp,k}$, $K = \overline{1, 8}$.

Авторами разработан алгоритм решения данной нечеткой задачи оптимизации, который базируется на нахождении недоминируемых альтернатив уровня α и сведении данной задачи к решению соответствующих задач четкой оптимизации “пессимиста” и оптимиста для данного уровня [2] и композиции полученных решений. Проведены экспериментальные исследования предложенного алгоритма оптимизации в нечетких условиях. Сам алгоритм и результаты его исследований приводятся в докладе.

Литература. 1. Зайченко Е.Ю., Зайченко Ю.П. Сети с технологией MPLS: моделирование, анализ и оптимизация. Киев, НТУУ “КПИ”, 2008. – 240 с. 2. Згуровский М.З., Зайченко Ю.П. Модели и методы принятия решений в нечетких условиях. Киев, “Наукова думка”, 2011. – 278 с.

Караюз І.В., Бідюк П.І.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Інтегрована модель для прогнозування макроекономічних процесів

Сучасні процеси в економіці та фінансах, особливо процеси економіки перехідного періоду, характеризуються нестаціонарністю, нелінійністю та наявністю різних режимів функціонування. Так, більшість процесів містять тренди, деякі характеризуються наявністю гетероскедастичності (дисперсія змінюється у часі) або нелінійності того чи іншого типу (нелінійності стосовно змінних або параметрів моделі). Для дослідження розвитку таких процесів необхідно будувати інтегровані (комбіновані) моделі, які надають можливість враховувати негативні впливи, що призводять до погіршення адекватності моделей і остаточних результатів – оцінок прогнозів розвитку досліджуваних процесів.

Для побудови адекватних математичних моделей з метою прогнозування фінансово-економічних процесів на макрорівні, необхідно враховувати такі фактори:

- фактори попиту: реальна заробітна плата; процентні ставки; обмінний курс; ціни на активи; споживчі настрої;
- фактори пропозиції:
 - у довгостроковій перспективі: рівень інфраструктури; розвиток технологій; людський капітал; ринок праці;
 - у короткостроковій перспективі: товарні ціни; політична нестабільність;
- фактори, що не вимірюються грошима, але мають вплив на попит і пропозицію: невиробничі види діяльності; якість навколишнього середовища; виснаження ресурсів; якість життя; бідність та економічна нерівність; погодні умови.

Для врахування перелічених невизначеностей запропонована інтегрована модель, функціональна схема якої подана на рис. 1. Серед вибраних факторів можуть бути такі, що не вимірюються, наприклад, об'єм капіталу, що безконтрольно пересилається за кордон. Для оцінювання значень таких змінних застосовується експертне оцінювання або спеціальні алгоритми, у тому числі алгоритми оптимальної фільтрації [1–3]. Попередня обробка статистичних даних забезпечує розв'язання таких задач: нормування вимірів у заданому діапазоні значень, заповнення пропусків, обробку екстремальних значень, зменшення впливу похибок вимірів та випадкових зовнішніх збурень (шляхом оптимальної фільтрації даних), виявлення нелінійності та ідентифікація її типу. Після належної підготовки даних будуються регресійні моделі (оцінювання їх структури і параметрів), які перетворюються у простір станів і використовуються надалі для оцінювання прогнозів розвитку вибраних процесів. Фільтр Калмана (ФК) використано для розв'язування таких задач: оптимальне оцінювання станів досліджуваних процесів; оцінювання коваріацій похибок вимірів і випадкових збурень стану (адаптивна фільтрація), а також для оцінювання невимірюваних компонент вектора стану моделі.

Для забезпечення високої якості остаточних результатів на кожному етапі обробки даних використовується відповідна множина критеріїв якості: критерії



Рис. 1. Функціональна схема інтегрованої моделі

якості даних, критерії адекватності моделей-кандидатів і критерії якості оцінок прогнозів. У випадку, коли не вдається отримати оцінки прогнозів бажаної якості, для додаткового вибору факторів впливу на основі статистичних даних будується і використовується байєсівська мережа (БМ). Це ймовірно-статистична модель у формі спрямованого графа, вершинами якого є вибрані змінні, а дуги відображають фактичні причинно-наслідкові зв'язки між змінними [4]. Перевагами БМ є такі: можливість використання в одній моделі статистичних даних, експертних оцінок та категорійних даних; допускається висока розмірність моделі стосовно кількості змінних; структура моделі оптимізується з використанням альтернативних критеріїв якості; існує множина альтернативних методів оцінювання ймовірного висновку, який можна формувати у довільному напрямі на графі.

Використання інтегрованої моделі надає можливість урахування факторів, що не вимірюються, виключення впливу випадкових факторів, покращення адекватності моделей та якості оцінок прогнозів. Реалізація запропонованого підходу до побудови математичних моделей на основі статистичних даних, оцінювання коротко- та середньострокових прогнозів виконується за допомогою системи підтримки прийняття рішень (СППР), спроектованої з використанням принципів системного аналізу [5].

Приклад використання СППР для прогнозування фінансових процесів ціноутворення за фактичними даними подано нижче. Результати короткострокового прогнозування волатильності прибутку за акціями Microsoft подані у таблиці. Для аналізу точності оцінок прогнозів вибрані такі статистичні критерії: середня абсолютна похибка (САП) і середня абсолютна похибка у процентах (САПП). У табл. 1 подані статистичні характеристики точності короткострокового прогнозування волатильності на навчальній та перевіірчій вибірках для досліджених моделей авторегресії з умовною гетероскедастичністю (АРУГ), узагальненої АРУГ (УАРУГ), експоненційної УАРУГ (ЕУАРУГ) і моделі стохастичної волатильності (МСВ). Значення середньої абсолютної похибки (САП) і середньої абсолютної похибки у відсотках (САПП), отримані на навчальній вибірці для моделей УАРУГ НВ, ЕУАРУГ НВ, МСВ НВ, показують менші значення похибок прогнозування ніж оцінки прогнозів моделей на перевіірчій вибірці (УАРУГ ПВ, ЕУАРУГ ПВ та МСВ ПВ), що і очікувалось. Зазначимо, що модель АРУГ відрізняється значною неточністю оцінок прогнозів як на навчальній вибірці, так і поза її межами, що пояснюється суттєвим спрощенням її структури. Застосування фільтра Калмана для розв'язування згаданих вище задач дало можливість підвищити якість оцінок прогнозів волатильності (без моделі АРУГ) від 7,1% до 50%. Подібні результати моделювання і прогнозування волатильності отримані для цін акцій компаній Google і Cisco. Обчислювальні експерименти, виконані з метою моделювання і прогнозування макроекономічних процесів, також свідчать про корисність практичного застосування запропонованої інтегрованої моделі.

Література. 1. Haykin S. Kalman filtering and neural networks. – New York: John Wiley & Sons, Inc., 2001. – 284 p. 2. Evensen G. Data assimilation: the ensemble Kalman filter. – Berlin: Springer, 2009. – 314 p. 3. Gibbs B.P. Advanced Kalman filtering. – Hoboken (New Jersey): John Wiley & Sons Inc., 2011. – 627 p. 4. Згуровський М.З., Бідюк П.І., Терентьев О.М., Просьянкіна-Жарова Т.І. Байєсівські мережі у системах підтримки прийняття рішень. – Вінниця: Едельвейс, 2014. – 300 с. 5. Бідюк П.І., Гожий О.П., Коршевніюк Л.О. Комп'ютерні систем підтримки прийняття рішень. – Миколаїв: Вид-во ЧДУ ім. Петра Могили, 2012. – 380 p.

Табл. 1. Прогнозування волатильності прибутку акцій Microsoft

Модель	САП без ФК	САПП без ФК	САПП з ФК
АРУГ НВ	0,000359	9454,4	9188,7
УАРУГ НВ	0,0000791	45,960	36,270
ЕУАРУГ НВ	0,4353	4,9400	3,7530
МСВ НВ	0,64	7,5000	4,9023
АРУГ ПВ	0,00041	5123,5	2494,4
УАРУГ ПВ	0,00013	51,993	28,396
ЕУАРУГ ПВ	0,5053	5,93	4,0710
МСВ ПВ	0,84	10,90	6,9590

Клімова О.В.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Методи прогнозування рекламної активності на телебаченні України

Реклама – це інформація, представлена в різноманітних формах, яка адресована особам для привернення уваги до об'єкта рекламування та підтримки інтересу до нього.

В роботі досліджується проблема прогнозування рекламної активності на телебаченні України. Основна задача полягає в прогнозуванні рейтингів рекламних посилок, що, в свою чергу, дозволить медіапланерам створити схему ефективного розміщення рекламних роликів.

Медіапланування – складання плану рекламної кампанії при оптимальному виборі необхідних каналів розміщення реклами в ЗМІ на основі даних маркетингових і медіадосліджень. Основними показниками якості медіапланування є ефективне охоплення і вартість схеми розміщення. Ефективність впливу рекламного звернення на цільову аудиторію буде сильно залежати від того, яку частину цієї аудиторії досяг сигнал рекламного повідомлення і скільки рекламних контактів було у представників цільової аудиторії.

Отже, найбільш вагомими показниками при плануванні рекламної кампанії є:

1. Affinity Index — індекс відповідності, визначається як співвідношення між рейтингом щодо цільової аудиторії і рейтингом щодо населення в цілому. Показує, наскільки краще цільова аудиторія контактувала з подією, ніж базова. Очевидно, що чим більше даний індекс, тим менше коштів буде витрачено на закупівлю рейтингів.
2. Gross rating point — сумарний рейтинг, набраний в результаті рекламної кампанії. Розраховується як сума рейтингів всіх рекламних повідомлень.
3. Target rating point — цільовий рейтинг, визначається як сумарна кількість рейтингів цільової аудиторії, яка контактувала з рекламним повідомленням. Головна відмінність від GRP полягає в тому, що в розрахунках використовується не вся аудиторія, а лише цільова.

Основна проблема даної роботи полягає в побудові прогнозуючих моделей та знаходженні невідомої функціональної залежності між величиною, що прогнозується, та заданим набором показників, отриманих експериментальним шляхом. При цьому аналітичний вигляд моделі невідомий.

Для вирішення даної задачі обрано метод індуктивного моделювання МГУА, розроблений А.Г. Івахненко та його учнями. Метод надає можливість побудови об'єктивної моделі в процесі роботи алгоритму.

Для даної задачі характерними є такі риси:

1. вигляд функціональної залежності невідомий, а визначений лише клас моделей;
2. коротка вибірка даних;
3. часові ряди нестаціонарні.

Варто зауважити, що МГУА використовує ідеї механізму еволюції:

1. схрещування або гібридизацію батьківських пар (аргументів) і генерацію нащадків;
2. селекцію та відбір найкращих.

Висновки. В роботі розглянута задача прогнозування рекламної активності, визначені та описані переваги і недоліки досліджених моделей. Проведені експериментальні дослідження та визначена ефективність даного методу в контексті його застосування в даній задачі.

Література. 1. Основы вычислительного интеллекта, М.З. Згуровский, Ю.П. Зайченко, Изд. «Наукова Думка». Киев, 2013.- 406 с. 2. Котлер Ф., Вонг В., Сондерс Д., Армстронг Г. Основы маркетинга, = Principles of Marketing: European Edition 4th. — 4-е европейское издание. — М.: «Вильямс», 2007. — 1200 с. 3. Балабанов А.В. Занимательное медиапланирование. — РИП-холдинг, 2001. — 104 с.

Кондратенко Н.Р., Снігур О.О.

Вінницький національний технічний університет, Вінниця, Україна

Використання інтервальних нечітких множин типу 2 в задачах ідентифікації складних об'єктів

В багатьох задачах ідентифікації складних об'єктів процес ідентифікації тісно пов'язаний з наявним рівнем знань про об'єкт та ефективністю засобів, що дозволяють зробити його опис у повному обсязі. Звідси виникає відома задача кількісного оцінювання ступеня ідентичності (адекватності) моделі та об'єкта. Розв'язання задачі визначення ідентичності моделі об'єкту-оригіналу є доцільним і тоді, коли об'єкт ідентифікації є таким, що слабо формалізується.

Залежно від ступеня нечіткості нечітких множин, що враховується при побудові нечіткої моделі, розрізняють нечіткі моделі типу 1, загальні моделі типу 2 та інтервальні типу 2.

Математичну модель для розв'язання поставленої задачі подамо у вигляді інтервальної нечіткої моделі ідентифікації. Структуру нечіткої моделі з інтервальними функціями належності для багатовимірних об'єктів типу багато входів – багато виходів зображено на рис. 1.

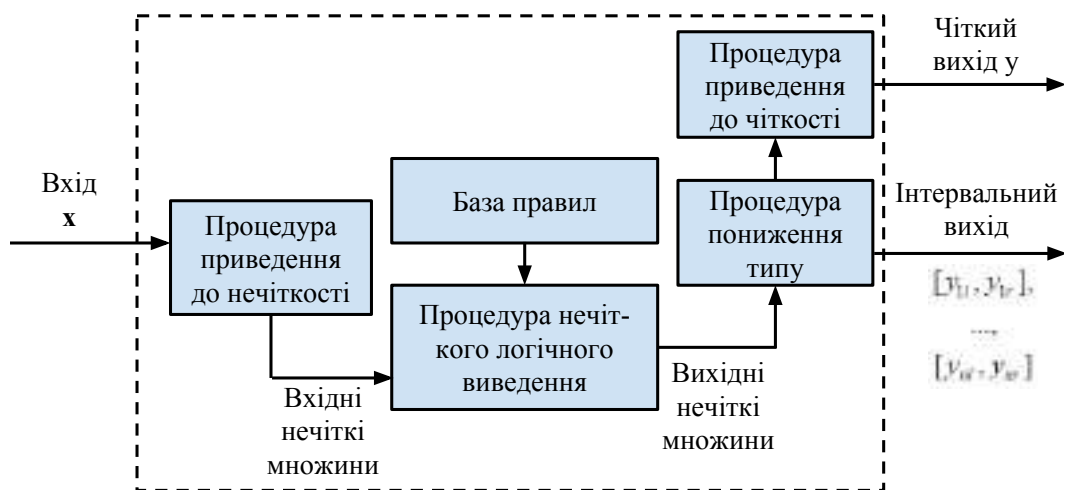


Рис. 1. Структура інтервальної нечіткої моделі типу-2

Модель відображає чіткі входи $x = (x_l, K, x_p)$ в інтервальні виходи $Y = ([y_1l; y_1r], \dots, [y_nl; y_nr])$ та чіткі виходи $y = (y_1, \dots, y_n)$. Для опису нечітких термів лінгвістичних змінних будемо використовувати інтервальні нечіткі множини типу 2. Тоді математична модель являє собою інтервальну нечітку модель типу 2, що включає базу правил (нечітку базу знань), процедуру приведення до нечіткості, процедуру нечіткого логічного виведення, процедуру пониження типу та процедуру приведення до чіткості (рис. 1).

Запропонуємо метод прямого генерування інтервальних нечітких моделей з експериментальних даних.

Наведемо його основні етапи.

Етап 1. Реалізація цього етапу за наявності вихідних даних можлива в такий послідовності кроків.

Нехай є експериментальна вибірка X :

$$X = X_1, X_2, \dots, X_n,$$

$$\text{де } X_i = (x_{1i}, x_{2i}, \dots, x_{ki}, y_i), i = 1..n;$$

n – кількість експериментальних прикладів,

k – кількість вхідних змінних,

y – вихідна величина.

Послідовність кроків така.

Проводимо генерування звичайної нечіткої моделі на основі експериментальної вибірки. Для цього використовуємо підхід до побудови нечітких моделей, коли нечітка модель будується на основі експериментальних даних, що визначають центри нечітких множин антецедентів та консеквентів правил. Отримуємо модель, для якої значення всіх вхідних параметрів відомі. Зазначимо, що така модель є надлишковою. В якості функції належності для вхідних змінних обираємо гаусову функцію належності:

$$\mu(x) = e^{-\left(\frac{x - m}{[\min(c), \max(c)]}\right)^2},$$

де $[\min(c), \max(c)]$ – діапазон зміни параметра c (належить інтервалу $\sigma \in [\sigma_l, \sigma_u]$) гаусової функції належності, зображено на рис. 2.

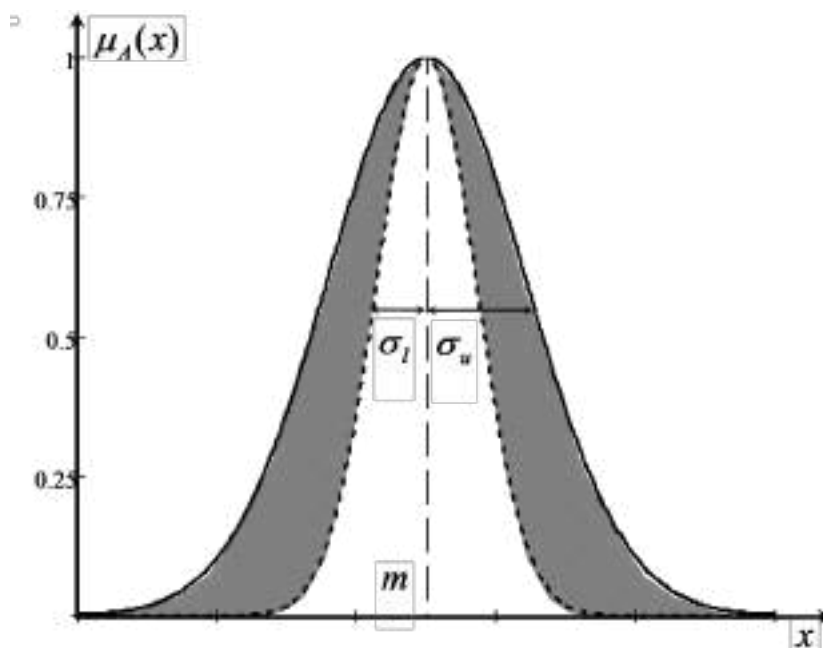


Рис. 2. Гаусова первинна функція належності зі сталим центром і невизначеним відхиленням

Для виконання переходу від звичайної до інтервальної функції належності при одночасному збереженні адекватності рішень, що приймаються системою, запропонуємо алгоритм. Алгоритм містить такі кроки: по-перше, побудова звичайної нечіткої моделі по зменшеній вибірці, оскільки початкова модель є надлишковою. На цьому кроці вилучаються правила, які не змінюють відгук моделі по всій вибірці. Далі алгоритм шукає верхню та нижню межу діапазону для параметру c гаусової функції належності, і одночасно при кожній зміні діапазону параметра c розраховує відгук моделі на всій вибірці.

Етап 2. На другому етапі реалізуємо корекцію вхідного вектора та експериментальної вибірки за визначенням розробника.

Останній крок для побудованої моделі – це тестування моделі на верифікованих даних та перевірка якості її функціонування.

В доповіді наведено комп'ютерне моделювання для задач ідентифікації багатовимірних об'єктів.

Кузнєцова Н.В., Фомін О.В.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Прогнозування ризику втрати користувачів онлайн-платформи

Швидка зміна технологій, використання мобільних телефонів, велика кількість соціальних мереж, розвиток доповненої та віртуальної реальностей суттєво впливають на життя людей і змінюють уподобання клієнтів. Якщо декілька років тому онлайн-ігри займали весь простір вільного часу користувачів-гравців, то тепер дедалі актуальнішим постає задача втримання існуючих клієнтів, оскільки залучення нових гравців до системи онлайн-ігор стає лише короткостроковою подією, – наслідком реклами або переадресації з іншої гри. У компанії, що здійснює підтримку користувачів онлайн-ігор, виникає необхідність оцінювання ризику втрати клієнта та прогнозування моменту, коли це може відбутися, та, за ідеального випадку, коли він може повернутися до гри.

Вхідні дані. Компанія А, що здійснює підтримку двох онлайн-ігор, Казино та Покер, збирає статистику щодо даних клієнтів, стратегії гри, та їх фінансової поведінки (суми ставок, виграші та програші, кількість ставок протягом доби) тощо.

Постановка задачі.

1. Провести кластеризацію гравців за типами агресивності гри;
2. Спрогнозувати, як буде клієнт грати в наступну гру;
3. Спрогнозувати, коли даний гравець повернеться знову в гру, тобто через скільки днів з моменту останньої гри він знову буде клієнтом компанії.

Методи вирішення задачі. Для описаної постановки задачі були відібрані наступні методи розв'язання: xgboost, ієрархічна кластеризація та моделі виживання. Кластерні дерева використовувались для типізації клієнтів, а xgboost, логістична регресія – для прогнозування факту повернення клієнта з певною ймовірністю. Моделі виживання у вигляді пропорційних ризиків оцінювали ймовірність повернення клієнта (функція виживання) та час, через який це може відбутися.

Розглянемо результати використання моделей ієрархічної кластеризації та xgboost.

Модель ієрархічної кластеризації. В рамках двох типів гри Покер (Холдем і Омаха) проводилася кластеризація гравців на основі простору статистик їх гри. Розмірність цього поля складала 135. Приклади статистик:

- VPIP – добровільний вклад своїх грошей в банк;
- PFR – підвищення ставки на «префлопі»;
- AF – фактор агресивності $AF = (\text{кількість ставок} + \text{кількість підвищень}) / \text{кількість урівнювань}$;

Кількість гравців: Холдем – 10 тис.; Омаха – 3 тис.

Результуючі дендрограми. В результаті кластеризації вдалося виділити по 6 основних груп в кожному із типів Покеру (рис. 1, 2), що дозволили ще краще проаналізувати поведінку гравців гри Покер. Варто відзначити, що аналіз менших кластерів дозволяє відслідковувати схожість в поведінці гравців, що свідчить про потужність апарату ієрархічної кластеризації.

Модель xgboost. Були агреговані та відібрані наступні вхідні характеристики: унікальний номер гравця, дата гри, сума ставок, сума виграшів, сума доходу, кількість зіграних годин, сума та кількість депозитів за день, сума та кількість виведень за день, кількість днів з першої гри, кумулятивний середній дохід з гравця, кумулятивний середній депозит та кумулятивний середній бонус, що він отримав. Для побудови моделей також бралися значення цих характеристик із лагом в 1 день.

При побудові цільового поля для моделі постала проблема цензурування даних, оскільки неможливо напевно сказати, чи покинув гравець платформу назавжди, або якщо і можна, то лише для невеликого їх числа. Ця проблема вирішувалася введенням поняття «значимого»

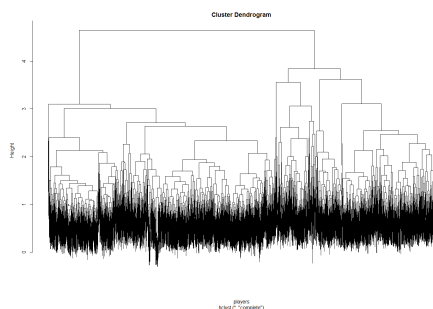


Рис. 1. Холдем

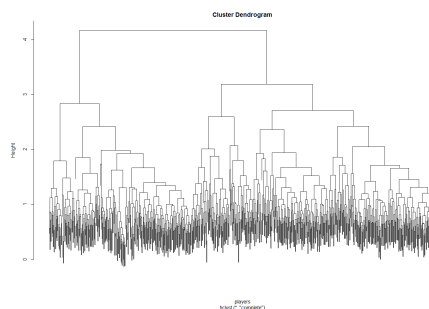


Рис. 2. Омаха

пропуску. І таким чином модель оцінює ймовірність настання саме значимого пропуску. Варто зазначити, що поєднання її з моделлю розподілу значимих пропусків дає можливість оцінювати ймовірність відтоку.

Цільове поле – це ознака того, що розрив між 2 послідовними іграми був значимим. При цьому значимим являється той пропуск, який входить до групи, що покриває 97,5% усіх пропущених гравцем днів. Значимий пропуск не може складати менше 4 днів, а якщо гравець за всю свою історію не пропускав більше 14 днів, то жоден з його пропусків не вважається значимим.

Гرادієнтний бустинг дерев рішень. Звичайні дерева рішень схильні до перенавчання. І досить часто не дозволяють досягти бажаних результатів у випадку, коли залежність між регресорами та цільовим полем є нелінійною. Для покращення результатів використовуються ансамблі дерев, такі як random forest, tree bagging та tree boosting. Одним із прикладів такого ансамблю є градієнтний бустинг:

$$F_m(x) = F_{m-1}(x) + \gamma_m h_m(x), \quad \gamma_m = \arg \min \sum_{i=1}^n L(y_i, F_{m-1}(x_i) + \gamma h_m(x_i)),$$

в якому кожне наступне дерево будується на основі похибки попереднього.

Для розв'язання проблеми перенавчання використовується регуляризація:

$$F_m(x) = F_{m-1}(x) + \nu \cdot \gamma_m h_m(x), \quad 0 < \nu \leq 1,$$

де ν – параметри, що регулюють швидкість навчання та перенавчання.

Саме такі моделі використовувалися для оцінок ймовірностей повернення гравців. Для кожної з ігор (Покер та Казино) була побудована своя модель. Графіки важливості, як середній інформаційний приріст за всіма деревами ансамблю, показників вказують на те, що виділяються основні регресори і для кожної гри вони свої.

В якості функції втрати застосовувався показник AUC.

Табл. 1. Значення AUC для кінцевих моделей

Тип гри	Тренувальна вибірка	Тестова вибірка
Покер	86,6%	86,5%
Казино	82%	77%

Висновки. Для онлайн-платформ ключовою задачею є збільшення доходу від клієнтів, що в даний момент користуються платформою, а також підтримання зацікавленості. Вартість залучення нового гравця до онлайн-платформи є значно вищою, ніж утримання існуючого. Тому ключовим моментом є дослідження і прогнозування моменту часу, коли гравець лише замислюється над призупиненням користування послугами, заохочення його новими бонусами, або переадресацію його на інший продукт компанії.

Луданов Д.К.¹, Зайченко Ю.П.²

¹КПІ ім. Ігоря Сікорського, ФМФ, Київ, Україна; ²КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Нейронні мережі та їх застосування

Штучні нейронні мережі – це математична програмна модель, побудована за принципом функціонування біологічних нейронних мереж – мереж нервових клітин живого організму [1]. Наразі використання штучних нейронних мереж надзвичайно поширене. Нижче наведено типи застосування нейронних мереж з прикладами досліджень у відповідних галузях.

Розпізнавання образів та класифікація. В якості образів можуть виступати різні за своєю природою об'єкти: символи тексту, зображення, зразки звуків і т.д. [1]. Здатність нейронних мереж навчатися є однією з основних можливостей, яка робить їх більш гнучкими і незалежними в порівнянні зі статистичними підходами в розв'язуванні задач класифікації. Зчитування тексту з цифрових зображень та класифікація об'єктів – це завдання, де широко застосовуються нейронні мережі [2].

Прийняття рішень та управління. Це завдання близьке до задачі класифікації. Класифікації підлягають ситуації, характеристики яких надходять на вхід нейронної мережі. На виході мережі повинна з'явитися ознака рішення, яке вона прийняла. Зокрема, на базі нейронних мереж було розроблено класифікатор, що надав змогу детектувати полум'я у відеопотоці [3].

Кластеризація. Під кластеризацією розуміється розбиття множини вхідних сигналів на класи, при тому, що ні кількість, ні ознаки класів заздалегідь не відомі. Після навчання така мережа здатна визначати, до якого класу належить вхідний сигнал [1].

Прогнозування. Здібності нейронної мережі до прогнозування безпосередньо впливають з її здатності до узагальнення та виділення прихованих залежностей між вхідними і вихідними даними. Після навчання мережа здатна передбачити майбутнє значення якоїсь послідовності на основі декількох попередніх значень або якихось існуючих зараз чинників [1].

Стиснення даних і асоціативна пам'ять. Здатність нейронних мереж до виявлення взаємозв'язків між різними параметрами дає можливість висловити дані великої розмірності більш компактно, якщо дані тісно взаємопов'язані між собою. Зворотній процес – відновлення вихідного набору даних з частини інформації – називається асоціативною пам'яттю.

Глибинне навчання. Це – набір методів машинного навчання для розв'язання складних завдань моделювання високорівневих абстракцій з великою кількістю вхідних даних [4]. При глибинному навчанні до великих штучних нейронних мереж передаються алгоритми навчання і об'єми даних, що постійно збільшуються. І чим більші об'єми даних, що обробляються, тим ефективніше процес. Навчання називають «глибинним», оскільки протягом часу нейронна мережа охоплює все більшу кількість рівнів, і чим «глибше» проникає мережа, тим вище її продуктивність. Незважаючи на те, що наразі більша частина глибинного навчання здійснюється під контролем людини, мета вчених – створити нейронні мережі, здатні формуватися та «навчатися» самостійно і незалежно [5].

Література. 1. Штучні нейронні мережі: що це таке? [Електронний ресурс]. – Режим доступу: <https://futurum.today/shtuchni-neironni-merezhi-shcho-tse-take/>. 2. Бандровський Г. Використання штучних нейронних мереж для розпізнавання тексту на графічних зображеннях. Тези доповідей VI Міжнародної наукової конференції «Інформація, комунікація, суспільство – 2017». – 2017. – С. 196–197. 3. Максимів О.П. Каскадний метод детектування полум'я у відеопотоці з використанням глибоких згорткових нейронних мереж. Науковий вісник НЛТУ України, 2017. – Том 27. – № 9. – С. 115–120. 4. Data trading. Whitepaper [Електронний ресурс]. – Режим доступу: https://data-trading.com/pdf/datatrading_whitepaper_ru_2017-12-27.pdf. 5. Что такое глубокое обучение? [Електронний ресурс]. – Режим доступу: <https://www.hpe.com/ru/ru/what-is/deep-learning.html>.

Павлов Д.Г., Чертов О.Р.

КПІ ім. Ігоря Сікорського, Київ, Україна

Використання нечіткої логіки для визначення мережевого шахрайства на прикладі шаблону інформаційної атаки

На відміну від булевої логіки, теорія нечітких обчислень дозволяє оперувати поняттями, більш подібними до людської мови, і, відповідно, будувати логічні висновки, наближені за своєю структурою та різноманітністю до таких, якими оперує людина. Переваги теорії нечіткої логіки були доведені багатьма дослідниками та результатами численних практичних застосувань [3]. Метою цієї роботи є демонстрація можливості та переваг застосування теорії нечіткої логіки для визначення мережевого шахрайства. Відомо, що мережа Інтернет стала невід'ємною частиною багатьох сфер діяльності людини, і рекламування не є виключенням. Переважна більшість рекламних Інтернет-оголошень оплачується за кліки, тобто, з рекламодавця стягується певна плата кожного разу, коли користувач натискає на рекламне оголошення. Така бізнес-модель породжує специфічний різновид шахрайства: *склікування* або *click fraud* [4], коли зацікавлена особа генерує штучні кліки із метою вичерпування бюджету рекламодавця. Постачальники послуг Інтернет-реклами постійно вдосконалюють свої системи захисту від склікування (наприклад, система захисту компанії Google [1]). Однак, як стверджується в звіті компанії PPC Protect [2], рівень шахрайства в галузі Інтернет-реклами залишається високим.

В роботі [5] авторами було запропоновано альтернативний підхід для визначення мережевого шахрайства типу склікування. На відміну від існуючих методів, запропонований підхід проводить аналіз інтернет-трафіку на стороні рекламодавця, тобто використовує лише доступну йому інформацію, та базується на визначенні типових шаблонів поведінки шахраїв. Одним з таких шаблонів є *шаблон інформаційної атаки* [6], зображений на рис. 1. Характерною особливістю цього шаблону є наявність чотирьох фаз: (I) фоновий шум, (II) проба (probe), (III) затишшя та (IV) атака (attack). У випадку мережевого шахрайства типу склікування ці чотири фази відповідають: (I) природньому рівню кліків (звичайна активність користувачів), (II) першій спробі проведення нападу, (III) вичікуванню результатів та (IV) склікуванню оголошень.

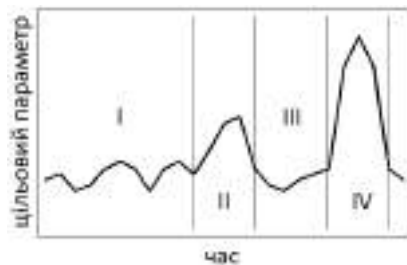


Рис. 1. Шаблон інформаційної атаки

При проведенні нападу за шаблоном інформаційної атаки, сплески активності, що відповідають другій та четвертій фазам, будуть певним чином відображені в розподілах параметрів рекламної кампанії [7]. Для спрощення, в цій роботі розглядаються лише два найважливіші параметри: часові розподіли кількості кліків (clicks) та кількості конверсій (conv). Останній параметр, *конверсія*, визначається як певна дія користувача на сайті, що приносить рекламодавцеві прибуток, наприклад, купівля товару або реєстрація [5]. Очевидно, що кількість кліків зростатиме впродовж фаз проби та атаки, а кількість конверсій буде зменшуватись. Для формалізації тверджень "*кількість кліків збільшилась*" або "*кількість конверсій зменшилась*" в цій роботі пропонується використовувати засоби нечіткої логіки.

Для побудови системи нечіткого виведення необхідно визначити множину вхідних і вихідних лінгвістичних змінних, а також сформувати систему нечітких продукцій, які будуть зв'язувати ці значення між собою. В нашому випадку вхідними змінними є відхилення розподілів параметрів рекламного трафіку (кількості кліків та конверсій) від очікуваних. Як відхилення розподілів від очікуваних пропонується брати відносне відхилення, оскільки абсолютний розкид залежить як від досліджуваного параметру, так і від якості рекламної кампанії.

Після визначення вхідних змінних необхідно визначити їх значення та для кожного з них задати функції приналежності. Для відносного відхилення кількості кліків та конверсій будемо розглядати такі значення: "менше очікуваного" (Low), "трохи менше очікуваного" (sLow),

“відповідає очікуваному” (Exp), “трохи більше очікуваного” (sHigh), “більше очікуваного” (High).

Функції приналежності кожного нечіткого параметру будемо визначати за допомогою гаусіани загального вигляду $g(x) = \exp\left(-\frac{(x - m)^2}{2\sigma^2}\right)$. Таким чином, для визначення вигляду деякої функції приналежності потрібно задати два параметри: математичне сподівання m та середнє квадратичне відхилення σ . Будемо вважати відхилення великим (“більше очікуваного” або “менше очікуваного”), якщо воно вдвічі перевищує модельне значення. Виходячи з цих міркувань, отримаємо $m_{Low} = -1$ та $m_{High} = 1$. Очевидно, що параметр m для функції приналежності нечіткої змінної “відповідає очікуваному” буде дорівнювати 0. Математичні сподівання інших функцій приналежності визначимо з умови їх рівномірного розташування на відрізку $[-1; 1]$, а середні квадратичні відхилення з двох умов: 1) значення параметру σ є однаковим для всіх функцій приналежності; 2) в точці перетину двох сусідніх функцій приналежності їх значення дорівнюють 0,5. Враховуючи всі поставлені вище умови, матимемо функції приналежності з параметрами $m_{Low} = -1$, $m_{sLow} = -0,5$, $m_{Exp} = 0$, $m_{sHigh} = 0,5$, $m_{High} = 1$, $\sigma = \left(4\sqrt{2\ln 2}\right)^{-1}$, зображені на рис. 2.

Зисновком нечітких продукцій буде присвоєння вихідній змінній “ймовірність атаки” (pAtt) одного з наступних значень: “дуже низька” (vLow), “низька” (Low), “нижче середньої” (sLow), “середня” (Avg), “вище середньої” (sHigh), “висока” (High), “дуже висока” (vHigh). Для формування системи нечіткого виведення необхідно визначити правила, що зв’язують значення вхідних та вихідних лінгвістичних змінних. В табл. 1 наведено систему нечітких продукцій для випадків, коли ймовірність атаки вважається різного ступеню високою (sHigh, High та vHigh). Наприклад, останній рядок цієї таблиці інтерпретується як таке правило: “якщо кількість кліків в періоді проби й атаки була *більше очікуваного* значення, а кількість конверсій в ці періоди була *менше очікуваного* значення, то ймовірність атаки *вище середньої*”. Таким чином, застосування нечіткої логіки дозволяє побудувати математичну модель задач із неочевидною математичною формалізацією і надати системі управління можливість оперувати “людськими” поняттями.

Література. 1. Google’s Protection against Invalid Clicks, <https://www.google.com/intl/en/ads/adtrafficquality/invalid-click-protection.html>. 2. List of Fraud Statistics 2018, <https://ppcprotect.com/ad-fraud-statistics/>. 3. Н. Singh, M.M. Gupta, T. Meitzler, Z.G. Hou, K.K. Garg, A.M.G. Solo and L.A. Zadeh. “Real-life applications of fuzzy logic” // Advances in Fuzzy Systems, vol. 2013, 2013, 3 p. 4. K. Asdemir and M.A. Yahya. “Legal and strategic perspectives on click measurement” // SEMPO Institute Opinions and Editorials, 2006, 11 p. 5. О.Р. Чертов, Д.Г. Павлов, В.В. Мальчиков, and М.В. Александрова. “Виявлення аномальної поведінки користувача системи контекстної реклами” // Штучний інтелект, no. 4, pp. 476–483, 2010. 6. В.М. Фурашев, Д.В. Ланде. “Практичні засади прогнозування можливих загроз та ризиків шляхом аналізу взаємозв’язку подій з інформаційним простором” // Открытые информационные и компьютерные интегрированные технологии: сб. науч. трудов, no. 42, pp. 194–203, 2009. 7. Д.Г. Павлов. “Параметри трафіку для побудови поведінкових шаблонів склікування оголошень у мережі Інтернет” // Вісник Східноукраїнського національного університету ім. В. Даля, no. 8 (179), pp. 299–307, 2012.

Табл. 1. Нечіткі продукції

probe		attack		pAtt
clicks	conv	clicks	conv	
sHigh	sLow	sHigh	Low	High
sHigh	sLow	sHigh	sLow	High
sHigh	sLow	High	sLow	vHigh
sHigh	sLow	High	Low	vHigh
sHigh	Low	sHigh	sLow	sHigh
sHigh	Low	sHigh	Low	High
sHigh	Low	High	sLow	vHigh
sHigh	Low	High	Low	vHigh
High	sLow	High	Exp	sHigh
High	sLow	High	sLow	sHigh
High	sLow	High	Low	High
High	Low	High	sLow	sHigh
High	Low	High	Low	sHigh

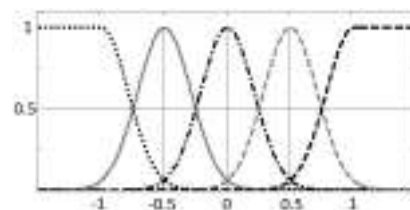


Рис. 2. Функції приналежності

Поворознюк А.И., Филатова А.Е., Шехна Х.

Национальный технический университет “Харьковский политехнический институт”, Харьков, Украина

Согласованная морфологическая фильтрация полутоновых изображений на основе фрактального анализа

Постановка проблемы и анализ литературы. Создание биомедицинских систем поддержки принятия решений (СППР), как части аппаратно-программных диагностических комплексов, может повысить качество проведения инструментальных обследований пациентов и снизить риски принятия неправильных решений. Процесс обработки информации в биомедицинских СППР состоит из последовательности соответствующих этапов, одним из которых является этап морфологического анализа биомедицинских изображений (БМИ) с локально сосредоточенными признаками (ЛСП). БМИ с ЛСП – это изображения, имеющие такую структуру, при которой диагностические признаки сосредоточены на небольших фрагментах их области определения. Задачей морфологического анализа является выделение на фоне помех информативных фрагментов (структурных элементов) БМИ, в результате которого формируются диагностические признаки в виде параметров найденных структурных элементов. Морфологический анализ БМИ с ЛСП является одним из ответственных этапов, так как ошибки на этом этапе приводят к принятию ошибочных диагностических решений или к отказу от принятия решения вообще. Этот этап требует использования специализированных методов морфологического анализа данных, учитывающие особенности БМИ с ЛСП и методы их преобразования [1]. Наиболее полно теория морфологического анализа разработана для обработки изображений и сцен различной природы [2, 3]. Однако в медицинской практике для морфологического анализа БМИ с ЛСП чаще используются различные эвристические методы обработки, которые на интуитивном уровне оперируют понятием формы сигнала или изображения в целом или отдельных их участков.

Таким образом, классические методы цифровой обработки изображений не учитывают особенности БМИ с ЛСП. Поэтому актуальность темы состоит в необходимости решения проблемы развития теоретических основ и средств поддержки принятия решений при проектировании биомедицинских систем на основе морфологического анализа БМИ с ЛСП с целью повышения качества инструментального обследования пациента.

Цель работы. Разработка метода согласованной морфологической (СМ) фильтрации биомедицинских полутоновых изображений с учетом фрактальных размерностей.

Морфологический анализ БМИ на основе анализа фрактальных размерностей. На основе обобщенной постановки задачи морфологической фильтрации, рассмотренной в [2, 3], авторами предложен СМ-фильтр следующего вида:

$$\varphi_{MM}^w(p, a, M) = o + K_{MM}^w(p, a, M)(a - o), \quad (1)$$

где $a, p \in \Omega$ – образ (элемент изображения) и эталон соответственно; Ω – множество (пространство) образов (изображений), характерных для некоторой морфологической системы; M – модель формы; $K_{MM}^w(p, a, M)$ – локальный морфологический коэффициент согласования (морфологический коэффициент согласования, заданный в окне w).

В работе предлагается метод морфологического анализа маммограмм на основе анализа фрактальных размерностей, в котором выделение патологических структур на маммограмме выполняется с помощью СМ-фильтра вида (1) без явного задания модели формы патологических структур. Исследования на реальных маммограммах показали, что значения фрактальной размерности, рассчитанные для небольших участков изображения, можно использовать в качестве параметров преобразования исходного маммографического изображения в задаче морфологического анализа БМИ. В выражении (1) в качестве простейшего образа o предлагается принять изображение с нулевым значением яркости пикселей:

$$\phi_{MM}[x, y] = K_{MM}^{w[x, y]}(g, f^{w[x, y]})g[x, y]; \quad (2)$$

где $g[x, y]$ – исходное полутоновое изображение; $f^{w[x, y]}$ – нормированное значение фрактальной размерности в окне $w[x, y]$; $K_{MM}^{w[x, y]}(g, f^{w[x, y]})$ – локальный морфологический коэффициент согласования, учитывающий фрактальные свойства фрагмента изображения в окне $w[x, y]$.

Коэффициент $K_{MM}^{w[x, y]}(g, f^{w[x, y]})$ рассчитывается по следующему выражению:

$$K_{MM}^{w[x, y]}(g, f^{w[x, y]}) = \frac{\max(g[x, y], mx - f^{w[x, y]}) - \min(g[x, y], mx - f^{w[x, y]})}{g[x, y]}. \quad (3)$$

В работе была выполнена СМ-фильтрация маммограмм на основе анализа значений фрактальной размерности. Для выполнения морфологического анализа с помощью разработанного СМ-фильтра (3) был выполнен расчет фрактальной размерности в заданном окне с помощью метода «Differential Box Counting». Исследования реальных цифровых маммограмм показали,

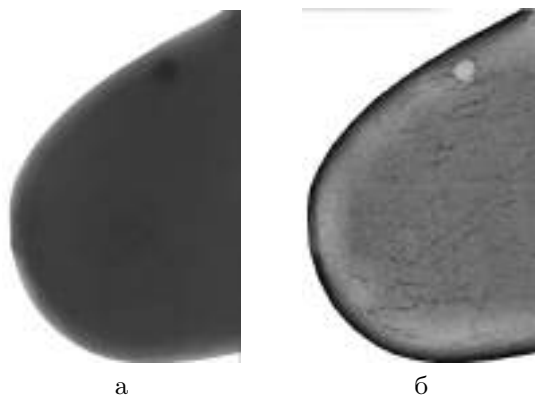


Рис. 1. Цифровая маммограмма: а – до фильтрации; б – после СМ-фильтрации

что при обработке всего изображения различия в значении фрактальной размерности маммограмм с наличием и отсутствием патологических структур не обнаружено. Это связано с тем, что патологические участки имеют достаточно небольшие размеры по сравнению с размерами изображения молочной железы, поэтому на изменение значений фрактальной размерности практически не влияют. При этом анализ значений фрактальной размерности, рассчитанных для участков маммограмм со здоровыми и пораженными тканями, показал, что для участков маммографических изображений, которые содержат уплотнение тканей, значения фрактальной размерности меньше, чем для участков маммографических изображений без видимых патологий. При этом независимо от размеров и положения патологических структур после СМ-фильтрации (3) выделяются светлые однородные по яркости участки, а в случае морфологического анализа нормальных маммограмм светлые участки не выделяются. Таким образом, значения фрактальной размерности, рассчитанные для небольших участков изображения, можно использовать в качестве параметров преобразования исходной маммограммы в задаче морфологического анализа БМИ с ЛСП.

Выводы. В работе предложен СМ-фильтр на основе фрактального анализа, который позволяет решить задачу выделения патологических структур на маммограмме.

Литература. 1. Guo, Q. A novel approach to mass abnormality detection in mammographic images / Q. Guo, V. Ruiz, J. Shao, F. Guo // In Proceedings of the IASTED International Conference on Biomedical Engineering. – Innsbruck. – 2005. – P. 180–185. 2. Рубис, А.Ю. Морфологическая фильтрация изображений на основе взаимного контрастирования / А.Ю. Рубис, М.А. Лебедев, Ю.В. Визильтер, О.В. Выглов // Компьютерная оптика. – 2016. – Т. 40, № 1. – С. 73–79. 3. Визильтер, Ю.В. Сравнение изображений по форме с использованием диффузной морфологии и диффузной корреляции / Ю.В. Визильтер, В.С. Горбачевич, А.Ю. Рубис, О.В. Выглов // Компьютерная оптика. – 2015. – Т. 39, № 2. – С. 265–274.

Пономаренко Р.М.^{1,2}

¹Институт кибернетики им. В.М.Глушкова НАН Украины, Киев, Украина; ²Национальный университет биоресурсов и природопользования Украины, Киев, Украина

Нечеткий логический вывод на основе многоуровневого параллелизма

Рассматривается организация иерархических систем нечеткого логического вывода с применением параллелизма нескольких уровней. Приводится теоретическое обоснование, а также экспериментальные оценки ускорения иерархических нечетких систем, построенных на основе многоуровневого параллелизма.

Использование нечеткой логики позволяет решать множество промышленных и научных задач в условиях неопределенности. На сегодняшний день, при решении задач большой сложности с использованием математического аппарата нечеткой логики, строятся иерархические (многоуровневые) нечеткие системы, которые призваны бороться с числом правил, проблема которого есть нерешена в пределах элементарной нечеткой системы [1, 2].

Целью данной работы является исследование и обоснование иерархических систем нечеткого логического вывода на основе многоуровневого параллелизма [3, 4].

Нечеткие системы типа Такаги-Сугено. Нечеткие системы типа Такаги-Сугено состоят из блоков нечетких предикатных правил *если-то (if-else)*:

$$R_j : \text{if } x_1 = A_1 \text{ AND } x_2 = A_2 \text{ AND } \dots \text{ AND } x_n = A_n \text{ then } y = g_j(x_1, x_2, \dots, x_n), \\ j = 1, 2, \dots, N,$$

где $g_j(x_1, x_2, \dots, x_n)$ – четкая функция от входных параметров. Пусть $w = \min_{x_i} \mu(x_i)$, $i = 1, 2, \dots, n$ Тогда общее заключение системы Такаги-Сугено можем получить по формуле (1):

$$y = \frac{\sum_{j=1}^N g_j w_j}{\sum_{j=1}^N w_j}. \quad (1)$$

Параллелизм уровня 1. Пусть $Y = y^q$, $q = 1, 2, \dots, Q$ – множество заключений элементарных нечетких систем в пределах одного яруса. Понятие яруса иерархических операций раскрывается в [5]. Тогда в пределах каждого яруса будем иметь Q параллельно запущенных элементарных нечетких систем.

$$y^q = \sum_{j=1}^{N^q} g_j^q w_j^q / \sum_{j=1}^{N^q} w_j^q. \quad (2)$$

Параллелизм уровня 2. Блоки правил в нечетком логическом выводе могут быть активированы в произвольном порядке. Следовательно, для каждой элементарной нечеткой системы Такаги-Сугено определим набор параллельно обрабатываемых операций, каждая из которых сводится к вычислению функции $g_j(x_1, x_2, \dots, x_n)$ и коэффициента w_j по каждому правилу R_j . После чего следует операция редукции с последующим нахождением заключения элементарной нечеткой системы по формуле (1).

Теорема (Пономаренко, 2017 [4]). Суммарное ускорение вычислений при использовании параллелизма уровня n определяется как произведение ускорений каждого из уровней (3).

$$S_p^{sum} = S_{p_1}^1 * S_{p_2}^2 * \dots * S_{p_n}^n = \prod_{i=1}^n S_{p_i}^i, \quad (3)$$

где p_i является максимально доступным количеством процессоров на каждом уровне i . За основу возьмем закон Амдала [5]:

$$S_p = \frac{1}{a + \frac{1-a}{p}} = \frac{p}{a * p + 1 - a}. \quad (4)$$

Введем новый коэффициент при a , который является величиной, обратной к (4). Тогда докажем теорему методом математической индукции для $n = 2$.

$$\begin{aligned}
 S_p^{sum} &= \frac{1}{a_1 * \left(\frac{1}{a_2 + \frac{1-a_2}{p_2}} \right)^{-1} + \frac{(1-a_1) * \left(\frac{1}{a_2 + \frac{1-a_2}{p_2}} \right)^{-1}}{p_1}} = \\
 &= \frac{1}{a_1 * (S_{p_2}^2)^{-1} + \frac{(1-a_1) * (S_{p_2}^2)^{-1}}{p_1}} = \frac{1}{\frac{a_1 * (S_{p_2}^2)^{-1} p_1 + (1-a_1) * (S_{p_2}^2)^{-1}}{p_1}} = \\
 &= \frac{p_1}{(a_1 * p_1 + (1-a_1)) * (S_{p_2}^2)^{-1}} = \frac{S_{p_2}^2 * p_1}{a_1 p_1 + 1 - a_1} = S_{p_1}^1 * S_{p_2}^2.
 \end{aligned} \tag{5}$$

Теорема доказана.

На основе графических ускорителей Nvidia CUDA было реализовано программную систему для исследования теоретической оценки ускорения иерархических систем нечеткого логического вывода при наличии сложных зависимостей между блоками нечетких правил, используя параллелизм уровня 2 (рис. 1), в сравнении с параллелизмом уровня 1 (технология MPI).

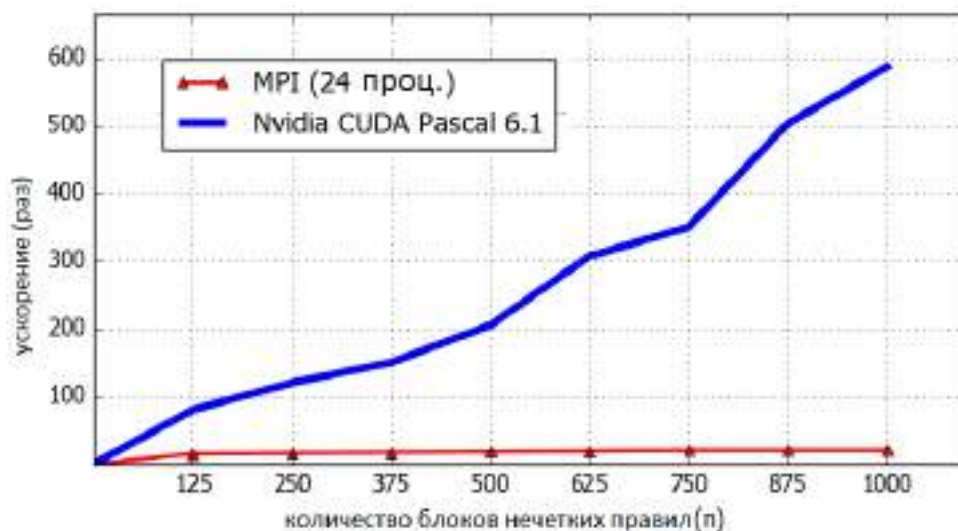


Рис. 1. Экспериментальная оценка ускорения иерархических нечетких систем параллелизма уровня 2

Литература. 1. Ершов С.В. Модель интеллектуальных агентов, основанная на нечеткой логике высшего типа // Компьютерная математика. – 2012. – №1. – С. 10–16. 2. Леоненков А.В. Нечеткое моделирование в среде MATLAB и fuzzyTECH / А.В. Леоненков. – СПб.: БХВ, 2005. – 736 с. 3. Ершов С.В., Пономаренко Р.М. Метод побудови паралельних систем нечіткого логічного виведення на основі графічних прискорювачів // Проблеми програмування. – 2017. – №4. – С. 3–15. 4. Пономаренко Р.М. Моделі паралельних ієрархічних систем для нечіткого логічного виведення // Комп'ютерна математика. – 2017. – №2. – С. 37–45. 5. Воеводин В.В. Параллельные вычисления / В.В. Воеводин, Вл.В. Воеводин. – СПб.: БХВ-Петербург, 2002. – 608 с.

Путренко В.В.¹, Пашинська Н.М.², Назаренко С.Ю.¹

¹Світовий центр даних з геоінформатики та сталого розвитку КПІ ім. Ігоря Сікорського, Київ, Україна; ²Київський національний університет ім. Тараса Шевченка, Київ, Україна

Використання 3D Space-Time Cube для інтелектуального аналізу даних

Починаючи з 2013 року, національна безпека України набуває все більшого значення. Кожного дня громадяни України зустрічаються із найрізноманітнішими загрозами природного, техногенного, соціального та військового характеру. Небезпечні процеси, надзвичайні явища, катастрофи, віртуальний та реальний тероризм та інше – це саме ті речі, на які органам державної влади потрібно реагувати практично кожного дня. Окрім цього, швидке реагування на надзвичайну ситуацію значно збільшує шанси на зменшення кількості жертв або іншого негативного наслідку, починаючи від конкретної людини до масштабів держави. Так, наприклад, міжнародна статистика свідчить, що кількість врятованих після землетрусу прямо залежить від початку рятувальних робіт. Якщо рятувальники прибудуть в зону землетрусу в перші три години, вони можуть врятувати до 90 відсотків людей, які залишилися живими, через шість годин – 50 відсотків. У подальшому шанси на порятунок зменшуються, а через 10 днів ведення рятувальних робіт практично втрачає сенс. Тільки за рахунок швидкого реагування можна зменшити кількість жертв на 20–30 відсотків [1].

Важливим питанням, що виникає при аналізі даних, є те, як події розподіляються в просторі та часі. Відповідне представлення даних необхідне для відповіді на це запитання. Для вирішення поставленої проблеми автори пропонують використовувати геоінформаційні системи (ГІС), а саме використання ГІС для використання методу аналізу багатомірних даних за допомогою побудови просторово-часового кубу даних (3D Space-Time Cube).

Просторово-часовий куб – це техніка 3D-візуалізації, призначена для одночасного представлення просторових та часових характеристик руху. Згідно з цим, точки траєкторій відображаються у тривимірному просторі, де вертикальна вісь зазвичай виражає час [2].

На початку 70-х років Хейгерстранд [3] розробив графічний погляд на час як додатковий просторовий вимір. Він запропонував тривимірну діаграму, так званий просторово-часовий куб, щоб показати життєві історії людей і те, як люди взаємодіють у просторі та часі. Площина куба являє собою двовимірний географічний простір, а висота куба – тимчасовий вимір. У той час, коли була представлена концепція, варіанти створення такої графіки обмежувались ручними методами. Це було очевидним обмеженням підходу, оскільки створення кожної нової діаграми виявилось трудомістким тренуванням. Сьогодні ж, коли сучасні комп'ютерні технології мають набагато кращі можливості для візуалізації даних, графічний погляд Хейгерстранда отримує друге життя.

Використання просторово-часового кубу потребує як просторові, так і часові дані, за для аналізу певних подій. Прикладами таких подій є землетруси, дорожньо-транспортні пригоди, випадки захворювань або спостереження рідкісних тварин, для яких діє трендове рівняння Манна-Кендалла [4].

Непараметричний тест Манна-Кендалла зазвичай використовується для виявлення монотонних тенденцій в серії даних. Нульова гіпотеза H_0 полягає в тому, що дані надходять із виборки з незалежними реалізаціями і ідентично розподіляються. Альтернативна гіпотеза H_A полягає в тому, що дані слідує за монотонічною тенденцією. Статистика тесту Манна-Кендалла розраховується за:

$$S = \sum_{k=1}^{n-1} \sum_{j=k_1}^n \operatorname{sgn}(X_j - X_k), \quad (1)$$

$$\operatorname{sgn}(x) = \begin{cases} 1, & x > 0; \\ 0, & x = 0; \\ -1, & x < 0. \end{cases} \quad (2)$$

Середнє значення $S \in E[S] = 0$ та дисперсія σ^2

$$\sigma^2 = \{n(n-1)(2n+5) - \sum_{j=1}^p t_j(t_j-1)(2t_j+5)\}/18, \quad (3)$$

де p – кількість пов'язаних груп у наборі даних, а t_j – кількість точок даних в j прив'язаній групі. Статистичне S є приблизно нормальним розподілом за умови, що використовується наступне Z-перетворення:

$$Z = \begin{cases} \frac{S-1}{\sigma}, S > 0; \\ 0, S = 0; \\ \frac{S+1}{\sigma}, S < 0. \end{cases} \quad (4)$$

Статистичне S є близьким до Кендал τ та дано як: $\tau = \frac{S}{D}$, де

$$D = \left[\frac{1}{2}n(n-1) - \frac{1}{2} \sum_{j=1}^p t_j(t_j-1) \right]^{\frac{1}{2}} \left[\frac{1}{2}n(n-1) \right]^{\frac{1}{2}}. \quad (5)$$

Хейгерstrand пропонував застосовувати просторово-часовий куб до даних про рух об'єктів, тобто про зміни просторових місць. У даній роботі автори пропонують застосувати концепцію Хейгерstrand до іншого типу даних, а само до аналізу подій.

За для прикладу авторами використовуються дані виходу в інтернет (дзвінки, повідомлення, пошук інформації...) користувачами мережі Vodafone, за період з 1 червня 2017 року по 31 серпня 2017 року (рис. 1). За своєю природою дані про події у мережі мобільного зв'язку відносяться до класу Big Data, який має просторову прив'язку до найближчої станції передачі сигналу. В роботі розглянуто набір даних, що містить близько 2 млн. записів, що ускладнює його аналіз за допомогою традиційних інструментів обробки даних.

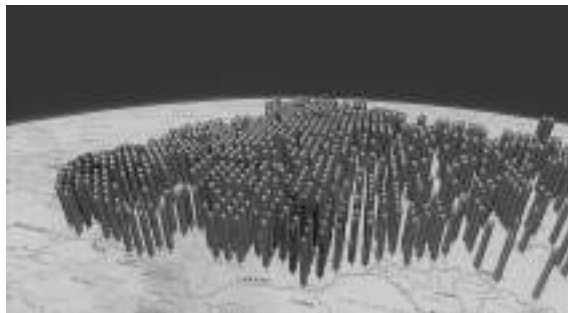


Рис. 1. Розподіл активності виходу в мережу інтернет серед користувачів Vodafone

Надалі за допомогою просторово-часового кубу можна аналізувати статистичні викиди активності користувачів мобільного зв'язку з прив'язкою до певної території, що дає змогу ідентифікувати певну екстраординарну ситуацію. Приклад такого використання просторово-часового аналізу даних надає змогу швидко виявляти та реагувати на небезпечні інциденти соціального або природного характеру.

Література. 1. Запорожець О.І., Заплатинський В.М., Халмурадов Б.Д. – Безпека життєдіяльності. 2. Jacek Rak, John Bay, Igor Kotenko, Leonard Popyack, Victor Skormin, Krzysztof Szczypiorski Computer Network Security: 7th International Conference on Mathematical Methods, Models, and Architectures for Computer Network Security, MMM-ACNS 2017, Warsaw, Poland, August 28–30, 2017, Proceedings. 3. T. Hägerstrand. What about people in regional science? Papers, Regional Science Association, 24, 7–21. 1970. 4. Peter Gatalsky, Natalia Andrienko and Gennady Andrienko Interactive Analysis of Event Data Using Space-Time Cube // Proceedings of the Eighth International Conference on Information Visualisation (IV'04).

Роговий А.В., Бідюк П.І.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Побудова скорингових моделей для оцінки платоспроможності клієнтів телекомунікаційного оператора

Актуальність дослідження. Для побудови надійного кредитного портфеля кожному банку необхідно постійно вдосконалювати свої скорингові моделі. Для українських фінансових установ основними ресурсами для побудови таких моделей є:

1. Демографічні показники: вік, стать, національність, місце проживання, тривалість проживання в актуальному місці, освіта, професія, тривалість працевлаштування, наявність власності, сімейний стан, наявність дітей та інше. Недоліком такого джерела інформації є відсутність достовірності вказаних клієнтом даних.
2. Фінансова історія клієнта: як часто в минулому клієнт повертав кредит, яка сума надходжень була на його рахунки. Але таке джерело можна використовувати лише у випадку повторного приходу клієнта у банк [1].

Але яким чином можна покращити такі моделі? Як діяти у випадку швидкого кредитування, коли людина залишає лише ім'я, паспортні дані та номер телефону? Для розв'язання таких задач необхідно шукати додаткові джерела інформації. Наприклад, дані телеком-оператора.

Побудова моделі. Для побудови моделі використано дані компанії "Київстар". На початку побудови моделі потрібно було відібрати характеристики, які могли найбільше впливати на платоспроможність клієнта. В якості характеристик було обрано такі: затрати на інтернаціональні дзвінки, кількість відвіданих протягом останнього часу країн, час з моменту підключення (часто шахраї купляють нові sim-карти), поповнення рахунку, використання інтернет трафіку, модель телефону.

Для збільшення стійкості моделей за часом над всіма неперервними характеристиками був проведений бінінг.

Основна ідея: По кожній з неперервних характеристик ми будуємо своє дерево рішень для прогнозу цільового значення. Кожна гілка представляє собою нову категоріальну характеристику. Бажано будувати не глибокі дерева (до 5–6 гілок) для відсутності перенавчання.

Переваги підходу:

1. Модель при застосуванні бінінгу стає менш чутлива до змін характеристик.
З часом в деяких характеристиках може змінюватись їх математичне очікування. Наприклад, показник середніх витрат на рахунок з часом може збільшуватись за рахунок інфляції. За рахунок перетворення неперервної характеристики на категоріальну ми більше застраховані від подібних змін.
2. Автоматизація обробки пропусків (для null значень заводиться новий бін).
Це особливо відчутно при великому обсязі характеристик: в нас немає необхідності придумувати особливий спосіб заповнення null значень, оскільки ми можемо завести для таких значень окремий бін.
3. Боротьба з викидами.
Якщо в якійсь характеристиці з'явиться не притаманне для неї велике або мале значення, то модель не буде перенавчатись на цьому окремому прикладі, оскільки наша характеристика потрапить до окремого біна, який буде символізувати велике або мале значення.
4. Чутливість до нелінійних залежностей.
Припустимо, що в нас є показник «вік». Значення таргету має сильний зв'язок з особами,

вік яких складає від 22 до 30. Якщо в якості основної моделі ми використовуємо, наприклад, логістичну регресію, то дана характеристика має свою певну вагу, що значить, що при збільшенні цього показника ми створюємо більший вплив на таргетне значення, а при зменшенні – навпаки, зменшуємо, але таргетне значення найбільше залежить для показників, які знаходяться посередині. Якщо ж ми застосовуємо бінінг, то в нас з'явиться нова характеристика, яка буде символізувати, що вік особи складає від 22 до 30, і при наявності такої характеристики в нас буде посилюватися вплив на таргетне значення.

Основними характеристиками побудованих дерев для бінінгу були:

1. максимальна глибина дерев – 5;
2. критерій оптимізації – ентропія;
3. збалансованість класів;
4. мінімальний процент екземплярів на кожній гілці – 20%.

Після цього до всіх змінних для подальшої роботи було застосовано OneHotEncoding.

В якості основної моделі було використано логістичну регресію [2]

$$h_{\Theta}() = g(\Theta^T X) = g\left(\sum_{i=0}^n \theta_i x_i\right) \text{ where } x_0 = 1 \text{ and } g(x) = 1/(1 + \exp(-x)).$$

На наступному етапі була побудована модель логістичної регресії зі штрафною функцією l_2 (оскільки основний інтерес був в тому, щоб перевірити, яких якісних показників можна досягнути при використанні всього об'єму інформації) та з $C = 0,001$.

В результаті побудови моделі після розбиття вибірки на 5 стратифікованих фолдів було досягнуто середнє значення AUC = 0.69 на навчальній вибірці та AUC = 0.68 на тренувальній. Це доволі непогані результати, які вирішують поставлену перед нами задачу, але в моделі приймає участь велика кількість змінних (близько 800). Оскільки в українських реаліях математичне очікування цих змінних з часом може змінюватися. Наприклад, середня сума затрат на оплату телеком послуг має тенденцію до постійного росту, через нестабільність в українській економіці. Тому було вирішено зменшити кількість характеристик з метою покращення стійкості моделі до часових змін.

За допомогою XGBoost було відібрано 20 характеристик, які мали найбільші ваги. Обидва дали однаковий набір ключових характеристик. Після чого на цих 20 характеристиках була знову побудована модель логістичної регресії з тими самими параметрами регуляризації. В результаті чого точність моделі стала в порівнянні з минулим тестом більшою. Середнє значення AUC на 5 стратифікованих навчальних датасетах стало дорівнювати 0,69.

Висновок. В результаті проведених досліджень було підтверджена гіпотеза про те, що по даним телеком-оператора можна побудувати доволі непогану модель оцінки платоспроможності клієнтів.

Література. 1. Ковальов М.С. Методика построения банковских скоринговых моделей для оценки кредитоспособности физических лиц [Електронний ресурс] / Ковальов М.С. // Інтернет-журнал «Науковедение» Выпуск 2, март – апрель 2014. – Режим доступа: <http://www.bsu.by/Cache/pdf/49623.pdf>. 2. Сравнение методов интеллектуального анализа данных при оценивании кредитоспособности физических лиц / А.Н. Терентьев, П.И. Бидюк, А.В. Миронова [та ін.] // Проблемы управления и информатики. – К.: ИКИ НАНУ-НКАУ, 2009. – № 5. – С. 141–149.

Савченко Д.А.

КПИ им. Игоря Сикорского, ИПСА, Киев, Украина

СППР для прогнозирования уровня продаж торговой фирмы

Система поддержки принятия решений предназначена для поддержки многокритериальных решений в сложной информационной среде. При этом под многокритериальностью понимается тот факт, что результаты принимаемых решений оцениваются не по одному, а по совокупности многих показателей (критериев), рассматриваемых одновременно. При создании СППР используют экспертные оценки, нейронные сети, методы оптимизации, регрессионный анализ, байесовские модели и методы, и множество других современных методов и подходов. Система поддержки принятия решений обычно решает две основные задачи:

1. Оптимизация – выбор наилучшего решения из множества возможных;
2. Ранжирование – упорядочение возможных решений по предпочтительности.

В обеих задачах принципиальным моментом является выбор совокупности критериев, на основе которых в дальнейшем будут оцениваться и сопоставляться возможные (альтернативные) решения. Система СППР помогает пользователю сделать такой выбор.

В данной работе СППР предназначена для прогнозирования уровня продаж на предприятии. При выборе типов математических и статистических моделей для описания процессов были выбраны регрессионные модели. Примером для тестирования является магазин одежды и обуви.

В общем виде модель имеет вид:

$$y(k) + a_1(k-1) + a_2y(k-2) + \dots + a_ny(k-n) = b_1u(k-1) + b_2u(k-2) + \dots + b_mu(k-m) + v(k) \quad (1)$$

где $y(k)$ – зависимая переменная, которая представлена в виде линейной комбинации вектора наблюдений с компонентами $u(k)$ и шумовой составляющей $v(k)$; $u(k)$ – независимая переменная, которую можно измерять; $v(k)$ – случайное возмущение (шум), которое в большинстве случаев рассматривается как белый шум.

Коэффициенты рассчитываются по методу наименьших квадратов. Среди множества вариантов моделей выбирается лучшая по показателям адекватности, таким как коэффициент детерминации, информационный критерий Акайке, статистика Дарбина-Уотсона [1]. Критерий Акайке учитывает сумму квадратов ошибок, количество измерений и количество оцениваемых параметров модели. Для лучшей модели критерий имеет меньшее значение, так как он зависит от суммы квадратов погрешностей (СКП). Однако, кроме СКП, он учитывает длину выборки и количество оцениваемых параметров, что делает его более информативным.

В докладе приводится описание работы с предложенной СППР и результатов ее применения на основе данных с предприятия. Был получен краткосрочный прогноз уровня продаж на 3 месяца, что позволяет запланировать расходы и варианты развития предприятия. Результаты прогноза сравнивались с реальными и была оценена их точность с помощью средней погрешности в процентах и среднеквадратической погрешности [1]. Также были выявлены основные причины неточности системы, такие как недостаточно большая выборка данных, и определено направление дальнейшего развития системы.

Литература. 1. Бідюк П.І., Коршевніук Л.О. Проектування ІСППР. – Київ: НТУУ „КПІ”, 2010. – 340 с.

Тимощук О.Л., Дорундяк К.М.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Порівняльний аналіз методів прогнозування банкрутства підприємств України

Зростаючий рівень банкрутства українських підприємств дозволяє стверджувати про актуальність прогнозування та розвиток методів оцінки ймовірності банкрутства. Наявність надійних методів оцінювання ймовірності банкрутства дозволяє попереджати кризовий стан підприємств заздалегідь, на етапах зародження банкрутства.

Метою дослідження є здійснення порівняльного аналізу різних методик прогнозування банкрутства, виявлення їх переваг і недоліків.

Однією з найбільш відомих дискримінантних моделей є універсальна дискримінантна модель українського дослідника О. Терещенка [1]. Для порівняльного аналізу було обрано також модель Матвійчука [2] та підхід, який базується на використанні штучних нейронних мереж [3].

Побудовано штучну нейронну мережу типу перцептрон, який обчислює вихідні змінні через послідовне нелінійне перетворення у нейронах вхідного шару, зважених ваговими коефіцієнтами міжнейронних зв'язків. В якості незалежних вхідних даних було обрано такі фінансові показники: коефіцієнт трансформації (оборотності активів), коефіцієнт рентабельності активів, коефіцієнт швидкої ліквідності, коефіцієнт фінансової незалежності, частка оборотних виробничих фондів в обігових коштах. Залежною змінною став клас кризового стану, визначений на основі фінансового становища компанії: 1 – підприємству нічого не загрожує, в нього гарна фінансова стійкість; 2 – наявні певні фінансові труднощі, існує невелика ймовірність настання банкрутства; 3 – підприємства має високий рівень ризику стати банкрутом.

Для побудови та навчання нейронної мережі була обрана мова програмування Python із використанням бібліотеки scikit-learn та функції MLPClassifier. Навчальна вибірка складалася з 43 спостережень підприємств різних галузей та різного фінансового стану (зокрема, Публічне акціонерне товариство “Світоч”, Приватне акціонерне товариство “Росинка”, тощо). Вихідні класи кризи були визначені за допомогою методики Міністерства Фінансів України. Загалом, 18 спостережень належало до 1 класу, 10 – до другого і 15 – до третього. Точність моделі нейронної мережі на навчальній вибірці склала 97,7%, що свідчить про гарний результат.

Для перевірки точності та якості моделі було обрано 2 підприємства – Публічне акціонерне товариство “Оболонь” та Публічне акціонерне товариство “АВК”, що досі функціонують і дані яких є достатньо інформативними. Для аналізу були відібрані дані з фінансової звітності за 2011–2016 роки із бази даних smida.gov.ua.

Варто зазначити, що модель нейронної мережі для обох розглянутих підприємств дала точніші результати, ніж моделі Матвійчука та Терещенка. Модель Матвійчука характеризується досить високою точністю, але відносить підприємства лише до двох категорій – банкрут і не банкрут. Модель Терещенка правильно розпізнає збиткові підприємства, але при цьому іноді відносить до банкрутних ті підприємства, які ще перебувають у задовільному стані.

Висновки. Проведено порівняльний аналіз дискримінантних моделей Матвійчука та Терещенка з розробленою моделлю штучної нейронної мережі. Підхід, оснований на використанні штучної нейронної мережі, у порівнянні з дискримінантними моделями показав вищу точність оцінок ймовірності банкрутства та здатність виявлення прихованих залежностей між основними фінансовими показниками.

Література. 1. Терещенко О.О. Фінансова санація та банкрутство підприємств: [навч. посіб.] / О.О. Терещенко – К.: КНЕУ, 2009. – 412 с. 2. Матвійчук А.В. Нечіткі, нейромережеві та дискримінантні моделі діагностування можливості банкрутства підприємств / А.В. Матвійчук // Нейро-нечіткі технології моделювання в економіці. – 2013. – С. 71–118. 3. Хайкин С. Нейронные сети: полный курс / С.Хайкин — М.: «Вильямс», 2006. — 1104 с.

Фіногенов О.Д., Грачова О.А.

КПІ ім. Ігоря Сікорського, ФІОТ, Київ, Україна

Аналіз даних сайтів погоди на основі методу аналізу ієрархій

У сучасному світі зі збільшенням негативних факторів впливу все важливіше місце займає дослідження здоров'я людини. Першим кроком до усунення негативних ефектів є відповідні дії людини на погодні умови, прогнозування яких є вкрай трудомісткою та в багатьох випадках нечіткою задачею. На даний момент, все більше людей отримують дані погоди завдяки мережі інтернет, але на багатьох сайтах, що інформують про погодні умови, є відмінності як у показниках, так і в наведених даних. Вибір одного з декількох джерел інформації про погоду можна представити у вигляді багатокритеріальної задачі, одним з методів вирішення якої є метод аналізу ієрархій (МАІ), що містить аналіз не лише кількісних, але і якісних показників. Вибір людиною погодних сайтів також обумовлений великою інерцією – звичне налаштування або інтуїтивно зрозуміле подання інформації може бути більш впливовим фактором, а ніж достовірність. У задачі, що є предметом досліджень, вихідними множинами елементів було обрано дані про температуру (Т), вологість (Н), швидкість і спрямованості вітру (W), атмосферний тиск (P), імовірність опадів (Pr), хмарність (C) та оціночну температури (FL) за чотири періоди доби: ранок, день, вечір та ніч. Оскільки дані на сайтах можуть оновлюватися у режимі реального часу, показники знімаються двічі на добу – вранці і ввечері, що дозволяє в подальшому аналізувати достовірність. Для спрощення аналізу будемо вважати, що геометрично населений пункт являє собою точку, оскільки через великі розміри міста показники у різних районах можуть відрізнятися.



Рис. 1. Приклади стислих прогнозів погоди [1–3]

Інформація, що подається на сайтах, порівнюється із даними європейської метеостанції [4] у вигляді ступеню абсолютного відхилення за допомогою парних порівнянь. Угрупування даних протягом доби доцільно проводити тетрадами: по чотири показники кожного періоду доби об'єднуються в групу – це формує другий рівень ієрархії. Оскільки ознаки впливають на досягнення кінцевої мети нерівномірно, необхідно визначити інтенсивність впливу ознак на кінцеву оцінку альтернативи.

Висновки. Оскільки деякі з альтернатив не надають повний набір показників для аналізу, ієрархія є неповнозв'язною. Для формування повнозв'язної ієрархії загального виду доцільно ввести оцінку відхилення даних у шкалі оцінок $1 \div 7$ або менше, а для даних, які відсутні для альтернативи, використовувати оцінки $1/8$, $1/9$, щоб мінімізувати вплив на загальну оцінку [5, 6].

Література. 1. Meta Погода: <http://pogoda.meta.ua/> 2. Gismeteo: <https://www.gismeteo.ua/> 3. Український гідрометеорологічний центр: <http://meteo.gov.ua/> 4. World Weather Online: <https://www.worldweatheronline.com/> 5. Саати Т. Аналитическое планирование / Т. Саати, К. Кернс. – М.: Радио и связь, 1991. – 320 с. 6. Саати Т. Принятие решений. Метод анализа иерархий / Т. Саати. – М.: Радио и связь, 1993. – 210 с.

Чапалюк Б.В., Зайченко Ю.П.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Автоматична медична діагностика на базі зображень комп'ютерної томографії

В рамках дослідження, розглядалася задача діагностування ракової пухлини в легенях пацієнтів. Знімки легень пацієнтів, що були отримані за допомогою методів комп'ютерної томографії та мітки із позначенням присутності ракової пухлини, були взяті із публічно доступного датасету, що був опублікований в рамках змагання Data Science Bowl 2017 [1]. В даному наборі даних представлені знімки для більше ніж 1000 пацієнтів, які знаходяться в зоні високого ризику наявності ракової пухлини в легенях. Дані одного пацієнта складаються із серії рентгенівських знімків, де кожен знімок представляє собою певну частину легень. Чим вище роздільна здатність пристрою, тим більша кількість рентгенівських знімків може бути зроблена для одного пацієнта. Тобто різні пацієнти, можуть мати різну кількість знімків в залежності від типу машини, яка виконувала сканування. Склавши серію знімків пацієнта та після певної обробки даних, можливо отримати 3D зображення легень пацієнта, яке можна використовувати для задачі виявлення раку.

Однією із особливостей наявних даних пацієнтів є сильна варіація кількості зроблених знімків в рамках серії комп'ютерної томографії. Це призводить до того, що результуюче 3D зображення легень пацієнта може сильно відрізнятися один від одного по свої розмірності. Інша проблема полягає в тому, що вимоги до обчислювальної потужності та пам'яті для обробки таких даних росте кубічно разом із роздільною здатністю зображення. Тобто чим більшу роздільну здатність матиме машина, тим більш ймовірно, що ми не зможемо навіть завантажити необхідну модель в пам'ять відеокарти для виконання необхідних розрахунків. Із наявними даними у вибірці, неможливо подати ціле 3D зображення, так як воно просто не поміститься в пам'яті відеокарти. Існує декілька підходів, які обходять дану проблему:

1. Обробка кожного зображення серії знімків пацієнта окремо. В цьому випадку зазвичай використовується такі нейронні мережі як DenseNet [2] або ResNet [3], що були попередньо навчені на ImageNet, а потім навчанні та оптимізовані під знімки легень пацієнтів. Такий підхід зазвичай не показує найвищу точність, так як втрачається трьохмірна природа даних.
2. Обробка серії зображень:
 - Із серії знімків пацієнта створюється 3D зображення легень пацієнта, яке потім подається на вхід 3D згорткової мережі [4]. Цей підхід враховує трьохмірну природу даних і показує одні із найкращих показників точності роботи системи, але є дорогою з точки зору використання пам'яті та обчислювальної потужності.
 - Послідовно проходять по зображенням, передаючи їх на вхід рекурентної нейронної мережі, таких як LSTM або GRU [5].

Література. 1. Kaggle, "Data Science Bowl 2017" [Онлайн]. Доступ: <https://www.kaggle.com/c/data-science-bowl-2017>. [Дата звернення: 02.04.2018]. 2. G. Huang, Z. Liu, L.v.d. Maaten and K.Q. Weinberger, "Densely Connected Convolutional Networks" 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Honolulu, HI, 2017, pp. 2261–2269. 3. K. He, X. Zhang, S. Ren and J. Sun, "Deep Residual Learning for Image Recognition" 2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Las Vegas, NV, 2016, pp. 770–778. 4. F. Liao, M. Liang, Z. Li, X. Hu, S. Song, "Evaluate the Malignancy of Pulmonary Nodules Using the 3D Deep Leaky Noisy-or Network" [Онлайн]. Доступ: <https://arxiv.org/abs/1711.08324>. [Дата звернення: 02.04.2018]. 5. M. Grewal, M.M. Srivastava, P. Kumar, S. Varadarajan, "RADNET: Radiologist Level Accuracy using Deep Learning for Hemorrhage detection in CT Scans" [Онлайн]. Доступ: <https://arxiv.org/abs/1710.04934>. [Дата звернення: 02.04.2018].

Чечельницкая А.Б.

КПИ им. Игоря Сикорского, ИПСА, Киев, Украина

Прогнозирование цен на авиабилеты бюджетной авиакомпании Ryanair с помощью нейронных сетей и авторегрессии со скользящим средним

Сейчас все больше и шире нейронные сети применяются в задачах прогнозирования в авиационной отрасли. В работе были применены метод искусственных нейронных сетей (ИНС) с алгоритмом обратного распространения ошибки и метод авторегрессии со скользящим средним для определения модели прогнозирования продажи авиабилетов.

Модели авторегрессии со скользящим средним (ARIMA) являются популярным и гибким классом моделей прогнозирования, в котором все базируется на значениях исторических данных. Этот тип представляет собой базовый метод прогнозирования, который можно использовать в качестве основы для более сложных моделей. Первым шагом в применении методологии ARIMA является проверка стационарности. «Стационарность» подразумевает, что с течением времени серия остается постоянным. Если существует тенденция, как в большинстве экономических или бизнес-приложений, то данные НЕ являются стационарными. Данные также должны демонстрировать постоянную дисперсию колебаний со временем. Без выполнения этих условий стационарности многие значения, связанные с процессом, не могут быть вычислены.

Принцип построения искусственной нейронной сети основан на работе биологических (человеческих или других животных) нейронных сетей. ИНС, как и человек, обладает умениями решения и оценки, которые более выгодны по сравнению с формальными логическими рассуждениями. ИНС может изменять внутреннюю структуру в соответствии с внешней информацией как тип адаптивной системы. Входная информация ИНС — это сумма всех входных данных, умноженных на соответствующие им веса. Выходные данные вычисляются, подставив входное значение в функцию активации. Когда есть выходные данные, они передаются дальше и так повторяется для всех слоев, пока не доходит до выходного нейрона. Когда сумма весов превышает пороговое значение, нейрон запускает значение с плавающей запятой, которое затем передается через функцию активации для отображения значения результата. Когда функция затрат высока, алгоритм обратного распространения ошибки используется для изменения весов, чтобы уменьшить стоимость. В следующем цикле подается следующий вход и одновременно обновляются веса. Вышеупомянутый процесс рекурсивно выполняется, пока не будет низкой стоимости. Множество таких нейронов, соединенных вместе, известно как нейронная сеть. Основная операция обработки входных данных для прогнозирования вывода позволяет упростить проектирование ИНС. Бесспорным преимуществом является: возможность изучать и моделировать нелинейные и сложные отношения, гибкость, широкий арсенал алгоритмов обучения параметров сети, в отличие от многих других методов прогнозирования, ИНС не накладывает никаких ограничений на входные переменные.

Работа посвящена прогнозированию цен на авиабилеты, а данные, которые я использую в работе, координаты аэропорта вылета, координаты аэропорта назначения, дата и время отправки, цена билета взяты из официального сайта Ryanair методом парсинга.

Выводы.

1. В работе рассмотрена задача прогнозирования цен на авиабилеты с помощью искусственных нейронных сетей и авторегрессии со скользящим средним.
2. Рассматриваемая система позволяет эффективно решить задачу прогнозирования цен на авиабилеты, поскольку учитывает частоту полетов и спрос, что важно в этой отрасли.

Литература. 1. S.D. Barrett How do the demands for airport services differ between full-service carriers and low-cost carriers Journal of Air Transport Management, 10 (1) (2004), pp. 33–39. 2. S.J. Forbes The effect of air traffic delays on airline prices International journal of industrial organization, 26 (5) (2008), pp. 1218–1232. 3. D. Pitfield Predicting air-transport demand Environment and Planning A., 25 (4) (1993), pp. 459–466.

Чечула М.В., Погорілий С.Д.

Київський національний університет ім. Тараса Шевченка, Київ, Україна

Розробка методів роботи з BigData на мобільних пристроях з використанням технології GPGPU

Вступ. Сьогодні переважна більшість інформації створюється в необробленому медіаформаті за допомогою мобільних пристроїв, апогей розповсюдження яких досі недосягнутий. Це висуває глобальну проблему контролю за інформацією та її систематизацією і поширенням. Звідси постає потреба у програмному засобі обробки зібраної інформації на мобільних пристроях перед її безпосередньою публікацією в мережі Internet в режимі реального часу та узгодження такого засобу з вже існуючими для централізованих комп'ютерних систем зберігання і обробки даних. Перевагами такого методу є відсутність необхідності створення нових централізованих систем для обробки даних та залучення потужностей мобільних пристроїв [1, 2].

Постановка задачі. Поставлено задачею є створення програмного комплексу для розпізнавання та систематизації відео та аудіо інформації в реальному часі як на платформі Laptop під управлінням ОС Windows або Linux так і на мобільній платформі з ОС Android з використанням технології GPGPU з передбаченою можливістю інтеграції в високопродуктивні централізовані системи зберігання та обробки даних.

Реалізація. Як основний інструмент розв'язання цієї задачі використано бібліотеку TensorFlow, це дозволяє вирішити низку питань, пов'язаних з масштабованістю та кросплатформеністю програмного засобу. Побудова програмного рішення поділяється на такі основні кроки:

- селекція характерних медіа-даних для певного кола обраних явищ, які планується розпізнавати;
- розробка прототипу згорточної нейромережі з двома прихованими шарами та визначення кількості класів об'єктів, які планується розпізнавати;
- власне створення нейромережі з допомогою високопродуктивної системи GPGPU на платформі Laptop;
- імплементація програмного застосування як на платформі Laptop так і на мобільній платформі для використання нейромережі для обробки медіа-даних в реальному часі, які збираються за допомогою вбудованих камери в режимі HD 30FPS та мікрофону.

В ході дослідження було проведено серію експериментів з покращення узагальнюючої здатності нейромережі та пришвидшення її роботи на мобільному пристрої за рахунок попередньої обробки медіаданих та мінімізації кількості операційних ресурсів, які використовуються.

Висновки та результати. Результатами роботи є кросплатформений програмний комплекс на мові програмування Python 3, який включає в себе програмні застосування для обробки відео та аудіо інформації в реальному часі на Laptop та мобільній платформі. Під час експериментів з нейромережею помічено, що використання вікна Хана для попередньої обробки аудіосигналів для навчання підвищувало узагальнюючу здатність нейромережі за рахунок пригнічення шумів. Також досягнуто 3-х кратного зменшення часу тренування нейромережі для обробки аудіо з використанням графічного адаптеру в порівнянні з використанням центрального процесору та 8-ми кратного зменшення часу для обробки відео. Вдалося досягти плавної роботи нейромережі на мобільній платформі.

Апаратні засоби. В роботі застосовувалися: CPU Intel Core i7 6700HQ, GPU Nvidia GTX970M з 1280 CUDA ядрами, Mobile CPU Cortex-A15, Mobile GPU PowerVR SGX544MP3.

Література. 1. D. Dietrich, B. Heller and B. Yan, Data Science and Big Data Analytics // John Wiley and Sons. – 2016. – р. 435. 2. Погорельский С.Д., Бойко Ю.В., Трибрат М.И., Грязнов Д.Б. Анализ методов повышения производительности компьютеров с использованием графических процессоров и программно-аппаратной платформы CUDA / Математичні машини і системи, 2010, №1. – С. 40–54.

Шулькевич Т.В., Селін Ю.М.

КПІ ім. Ігоря Сікорського, Київ, Україна

Математичний апарат для інтелектуального аналізу і прогнозування нелінійних нестационарних процесів. Досвід застосування

У доповіді описується математичний апарат, що може бути використаний для аналізу даних і для прогнозування нелінійних нестационарних процесів різної природи. Описано методи прихованих марковських моделей і лінгвістичне моделювання. Наведено результати чисельних експериментів по застосуванню математичного апарату.

Інтелектуальні методи аналізу даних, за визначенням, – це процес пошуку в «сирих» даних раніше невідомих, нетривіальних, практично корисних і доступних інтерпретації знань, необхідних для прийняття рішень в різних сферах людської діяльності [1]. Інтелектуальний аналіз даних – термін, який застосовується для опису отримання знань в базах даних, дослідження даних, обробки зразків даних, очищення та збору даних. Це процес виявлення кореляції, тенденцій, шаблонів, зв'язків і категорій [2]. Інтелектуальний аналіз даних розвивається на базі таких наук, як прикладна статистика, розпізнавання образів, штучний інтелект, теорія баз даних тощо.

Процес автоматичного пошуку прихованих закономірностей або взаємозв'язків між змінними в інтелектуальному аналізі даних ділиться на завдання класифікації, моделювання і прогнозування з використанням статистичних і математичних методів.

Дані, які отримуються для розрахунку показників різної природи, зазвичай мають вигляд часових рядів, і є безліч різних методів для моделювання та прогнозування цих часових рядів.

Але безліч існуючих методів не гарантують, що вони покривають всі можливі варіанти розвитку ситуації. Всі ці методи є 2-етапними: спершу аналізується ряд, далі обирається певний метод і робиться прогноз. Що робити, коли параметри ряду змінюються? Тому в доповіді пропонується розробка математичного апарату, що можна використовувати для прогнозування обох видів процесів різної природи у випадку зміни характеристик ряду.

Пропонуються до розгляду деякі методи, що можна застосовувати для використання в подібних завданнях. До них можна віднести наступні методи:

- приховані марківські моделі (ПММ);
- лінгвістичне моделювання (ЛМ).

Приховані марківські моделі. Для використання ПММ при розпізнаванні процесів необхідно розв'язати три задачі.

1. Якщо задані послідовність O_i спостережень і модель $\lambda = (A, B, \Pi)$, то як ефективно обчислити $P(O_i/\lambda)$ – ймовірність такої послідовності при заданих параметрах моделі?
2. Якщо задані послідовність спостережень O_i і модель $\lambda = (A, B, \Pi)$, то як визначити відповідну послідовність внутрішніх станів?
3. Якщо задані послідовність спостережень O_i , то як визначити параметри моделі $\lambda = (A, B, \Pi)$, виходячи з критерію максимізації $P(O_i/\lambda)$?

Тут $\lambda = (A, B, \Pi)$ – скорочений запис прихованої марковської моделі. Де

A – матриця ймовірностей переходів з одного елемента послідовності до іншого;

B – розподіл ймовірностей спостережування станів;

Π – розподіл ймовірностей початкового стану (зазвичай приймають $\Pi = 1$).

Всі три задачі мають відповідні алгоритми розв'язків. Але для аналізу нелінійних, нестационарних процесів потрібно вирішення першої задачі.

Це завдання оцінки моделі, яке полягає в обчисленні ймовірності того, що модель відповідає заданій спостережуваний послідовності. Очевидно, що рішення цього завдання дозволить

ефективно розпізнавати елементи з деякого обмеженого набору, маючи готові, вже навчені моделі цих елементів. Для цього обчислюється ймовірність генерації такої послідовності спостережень для кожної з наявних у базі даних моделей. Модель того, що з найбільшою ймовірністю генерує таку послідовність спостережень, і є шуканим результатом [3].

Таким чином, якщо у нас стоїть питання вибору найкращої моделі з набору вже існуючих, то рішення першої задачі дає нам відповідь на це питання.

Лінгвістичне моделювання. При лінгвістичному моделюванні відбувається заміна часового числового ряду спостережень на символічний ряд:

$L : F(t_i) \rightarrow \langle e_i \rangle$, де

$F(t_i)$ – часовий числовий ряд;

$\langle e_i \rangle$ – символічний ряд.

У загальному випадку лінгвістична модель динамічного процесу складається з наступних елементів:

$\langle D, I, L, G \rangle$, де

D – сукупність тимчасових рядів динамічного процесу і рядів, похідних від вхідних даних;

I – спосіб і правила інтервалізації;

L – морфізм відображення інтервального представлення низки на певний алфавіт;

G – відновлена граматики динамічного процесу.

Згідно етапів побудови лінгвістичної моделі вихідна задача буде розбита на такі підзадачі:

- підзадача отримання різницевих рядів;
- підзадача інтервалізації;
- підзадача лінгвістизації;
- підзадача побудови матриці переходів.

Отриману послідовність символів аналізують на наявність граматичних конструкцій. На виході отримуємо список таких граматичних конструкцій з ймовірностями їх наявності у досліджуваному процесі, а також матрицю ймовірностей переходу з символу в символ. Прогноз надається у вигляді спектру ймовірностей.

Запропонований математичний апарат відрізняється стійкістю отримуваних результатів і забезпечує підвищення якості прогнозів відповідних показників.

Запропоновано методику побудови математичних моделей нелінійних нестационарних процесів у вигляді лінгвістичного опису та прихованих марківських процесів, яка відрізняється відсутністю етапів оцінювання структури і параметрів моделі.

Наведені методи є універсальними як з боку виду отриманої інформації, так і з боку наявності в цій інформації нелінійностей і нестационарності. Але мають загальний недолік всіх статистичних методів – відсутність історичної інформації [4].

Автори в ході проведених досліджень провели безліч чисельних експериментів з багатьма числовими рядами і отримали практичне підтвердження придатності описаних методів для аналізу і прогнозування нелінійних нестационарних процесів різної природи.

Література. 1. U. Fayyad, G. Piatetsky-Shapiro, P. Smyth. (1996.) From Data Mining to KnowledgeDiscovery in Databases. AI Magazine. №17(3). 2. В.Л. Плєскач, Т.Г. Затонацька. (2011) Інформаційні системи і технології на підприємствах : підручник – К. : Знання. 3. 3. Селін Ю.М., Баклан І.В. (2013) Математичний апарат для прогнозування часових рядів економічного та екологічного типів, що можуть бути піддані зовнішнім впливам // Вєстн. Херсонського нац. ун-та. – №2(47). – С.315–318. 4. Т.В. Shulkevych, I.V. Baklan, Yu.M. Selin. (2016) Data Mining Mathematical Apparatus for Forecasting of Nonlinear Nonstationary Processes of Various Nature // Системні технології. Регіональний міжвузівський збірник наукових праць. – Випуск 6 (107). – Дніпро. – С. 151–158.

Plenary talks

System analysis of
complex systems of
various nature

Intelligent systems for
decision-making

High Performance
Computing and
Microsystems Engineering

Progressive information
technologies

3

High Performance Computing and Microsystems Engineering



Section 3

High Performance Computing and Microsystems Engineering

1. High performance computing: Supercomputing, distributed, cloud, Grid, and quantum computing, high performance systems and networks.
2. Microsystems engineering and design.

Секция 3

Высокопроизводительные вычисления и разработка микросистем

1. Высокопроизводительные вычисления: Суперкомпьютерные, распределенные, облачные, Грид- и квантовые вычисления, высокопроизводительные системы и сети.
2. Разработка и проектирование микросистем.

Секція 3

Високопродуктивні обчислення та розробка мікросистем

1. Високопродуктивні обчислення: Суперкомп'ютерні, хмарні, Грід- та квантові обчислення, високопродуктивні системи та мережі.
2. Розробка та проектування мікросистем.

Pechurin N.K.,¹ Kondratova L.P.,² Pechurin S.N.¹

¹National aviation university, Kyiv, Ukraine; ²Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

Leontiev's model for Internet of things information loading estimation

The estimation (forecasting) of the (future) information load on the Internet of things [1] is a necessary condition for the successful development of this relatively new object of information technology, in particular, for the effective distribution of efforts to create an information and telecommunications infrastructure.

In the context of a global computer network, it is not only the actual information load itself that is estimated from terminal devices (network information sources), but at least their number, i.e. load on the part of the technical support of the terminal nodes of the network seems to be very complicated (one of the authors [2] recalled the difficulties, the beginning of the 70s of the last century, evaluation of a similar indicator by the forces of the VNII POU – scientific research institute for a computer network of the scale of the USSR created for so-called OGAS under the scientific supervision of acad. V.M. Glushkov).

The report is devoted to the consideration of the approach to the estimation and forecasting of traffic intensity generated by information sources in the global computer network (things): it is suggested to use the indicator, which is mediated by the information load on the network, but formed by trustworthy organizations [3, 4]. This indicator is GDP at purchasing power parity, indirectly representing the volume of things (smart sensors and actuators) that can (potentially) generate an information load on the network. This indicator (in conjunction with GDP, GNP, etc.) is calculated annually for all countries, whose subscribers can compose the global computer network of things.

The main error, in the context of our problem, which the chosen indicator brings in, is due to the combination (folding) in it of two components: the annual volume of the produced goods (i.e. things that make up the computer network of things hardware) and services (noise in our view).

It is proposed, instead of the pending assumption of a direct linear relationship between GDP at purchasing power parity and annual production of things, to introduce, as was done, for example, in [5], the assumption of connections established by a system of algebraic balance equations in the form:

$$(E - A(t)) \cdot (X(t))^T - (R(t))^T = 0,$$

$$x_i(t) = x_{ith}(t) + x_{is}(t),$$

where the values of the components $A(t)$, $X(t)$, $R(t)$, in contrast to [5], depend on time, and the study of the model itself can be carried out using parametric programming methods.

To estimate (forecasting) the GDP at purchasing power parity, for example, it can be used the instrumentation group AR, MA, ARMA. A real example of using the proposed method of assessing the information load on the Internet of things is given.

References. 1. Kevin Ashton. That 'Internet of Things' Thing. In the real world, things matter more than ideas. – RFID Journal (June, 22, 2009). 2. Zaychenko Yu.P., Pechurin N.K. Algorithm of mathematical programming for determining the rational structure of CC network // Automated systems of data processing and management. – K.: IC of Academy of Sciences of Ukrainian SSR, 1975. – 4 p. 3. The World Factbook / Central Intelligence Agency. – Electr.resource. – www.cia.gov. 4. World Trade Organization. – Electr.resource. – stat.wto.org. 5. Pechurin N.K., Kondratova L.P., Pechurin S.N. Reliable balanced UAV computing cluster and V.Leontyev model 19 International Scientific and Technical Conference "System Analysis and Information Technologies SAIT-2017", May 22–25, 2017, Kyiv, Ukraine. – Kiev, 2017. – P. 202–203.

Sergiyenko A.M., Hasan M.J., Sergiyenko P.A.

Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

Square root calculations in FPGA

Introduction. The square root function $\sqrt{x}$ is important elementary function in the scientific calculations, digital signal and image processing [1]. The artificial neural nets need this function as well [2]. At present, the field programmable gate arrays (FPGAs) are expanded for solving the problems, where the function $\sqrt{x}$ calculations are of demand. There are different IP cores of the function $\sqrt{x}$, which are proposed by the FPGA manufacturers and third-party companies [3]. But these IP cores were designed decades ago and they usually don't take into account the features of the new FPGA generations. Therefore, they need improvements. In the presentation, an improved algorithm of the function $\sqrt{x}$ is proposed, which is suitable for the FPGA implementation.

CORDIC-type algorithms. The CORDIC-type algorithm of the elementary function calculation derives a single exact digit of the result in each computation step. The well-known CORDIC algorithm of the $\sqrt{x}$ calculations consists in the following. It calculates the function $\operatorname{atanh}(x/y)$. But the side result is the function $K\sqrt{x^2 - y^2}$, and by the substitution $x = A + 0.25, y = A - 0.25$, we get $x_n = K\sqrt{A}$ [3]. The disadvantages of this algorithm are additional multiplication to the coefficient $1/K \approx 1.207$, and repeating some iterations for the algorithm convergence.

More constructive algorithm is the CORDIC-like algorithm of the function $\sqrt{x}$ calculation [4], which is based on the following relations. For each number $x \in [0.25; 1.0]$ the coefficients $a_i \in [0; 1]$ are found so

$$x \prod_{i=1}^n (1 + a_i 2^{-i})^2 \approx 1.0; \Rightarrow 1/\sqrt{x} \approx \prod_{i=1}^n (1 + a_i 2^{-i}) \Rightarrow \sqrt{x} \approx x \prod_{i=1}^n (1 + a_i 2^{-i}). \quad (1)$$

The algorithm is the following:

```
x[0] = x; y[0] = x;
for(i = 0, i < n, i++) {t = x[i] + 2^(-i)*x[i];
    q = t + 2^(-i)*t;
    if (q < 1) {x[i+1] = q; y[i+1] = y[i] + 2^(-i)*y[i];} // a[i]=1
    else {x[i+1] = x[i]; y[i+1] = y[i];} // a[i]=0
}
```

The result is $\sqrt{x} \approx y[n]$.

Modernized algorithm. The most delay in the considered algorithm gives the twofold addition of the shifted data. These stages of addition can be substituted by a single stage:

$$q = (x_i + 2^{-i}x_i) + 2^{-i}(x_i + 2^{-i}x_i) = x_i + 2^{-i+1}x_i + 2^{-2i}x_i.$$

Because modern FPGAs perform the three input adder as a single stage of the six input look-up tables (LUTs), then such computations can be implemented for a single clock cycle without additional time and hardware overheads.

The algorithm analysis shows that when i reaches the limit $n/2$, then the most $i-1$ significant bits of x_i become equal to a one by any x_0 , and i most significant bits of y_i are exact digits of the result. Therefore, the rest of resulting bits can be calculated after analysis and computation the difference $1 - x_i$.

Consider $\varepsilon_1 = 1 - x_{n/2}$, and $\sqrt{x} = \varepsilon_x + y_{n/2}$. Then, to get the exact result, the correction ε_x is derived from the value ε_1 , and it is added to the approximated result. Due to (1),

$$\varepsilon_1 = 1 - x \prod_{i=1}^{n/2} (1 + a_i 2^{-i})^2; \varepsilon_x = \sqrt{x} - x \prod_{i=1}^{n/2} (1 + a_i 2^{-i}).$$

Then

$$z = \sqrt{x} \prod_{i=1}^{n/2} (1 + a_i 2^{-i}); \varepsilon_1 = 1 - z^2 = (1 + z)(1 - z); \varepsilon_x = \sqrt{x}(1 - z).$$

Consider $z \approx 1$, then $\varepsilon_1 \approx 2(1 - z)$; $\varepsilon_x \approx \sqrt{x}\varepsilon_1/2 = y_{n/2}(1 - x_{n/2})/2$. The result is $y_n = y_{n/2} + y_{n/2}(1 - x_{n/2})/2$. So, in order to obtain a refined result, the correction is added to the approximate result $y_{n/2}$. To do this, a subtraction and a multiplication should be taken. Moreover, due to the fact that ε_1 and ε_x have the zeroed most significant bits, then the multiplication is performed at half bit width. The resulting algorithm looks like the following:

```
x[0] = x; y[0] = x;
for (i = 0; i < n/2; i++) {
    q = xi + 2-(i+1)*x[i] + 2-2i*x[i];
    if (q < 1.0) {x[i+1] = q; y[i+1] = y[i] + 2-i*y[i]; }
    else {x[i+1] = x[i]; y[i+1] = y[i];}
}
y = y[n/2] + y[n/2]*(1.0 - x[n/2])/2;
```

In modern FPGA the two and three input adders have the same hardware volume and speed. So, the modified algorithm provides the speed-up approximately in four times comparing to the initial algorithm.

Experimental results. The derived algorithm was described by VHDL as the IP core. It was compiled for Xilinx Spartan-6 FPGA for various input and output data bit widths. Fig. 1 and Fig. 2 show the relation of the hardware costs in the number of LUTs, and the maximum clock frequency on the bit width for the combinatorial and pipelined networks of this IP core, respectively. It should be noted that the modules with a bit width up to 32 additionally have a single multiplication block DSP48, and the rest of them have four such blocks.

For comparison, Fig. 1, 2 show the characteristics of the IP cores offered by Xilinx Inc. In general, the proposed IP core loses to the branded one. But its advantages are that it is free and it can be configured for arbitrary input and output bit width, and for any FPGA type. In addition, the proposed IP core has a lower latent delay.

For example, for 24 bits, its latent delay is only 15 cycles versus 24 cycles of the competitor. This means that when implementing the floating-point calculations, the proposed module provides greater performance.

Conclusion. The modified CORDIC-like algorithm for deriving the square root function is proposed. The algorithm is distinguished by the minimized number of steps, which is proportional to the given data and result bit width. The algorithm is described in VHDL and is intended for the FPGA implementation. It is the most effective during its implementation in the floating point square root module.

References. 1. R. Woods, J. McAllister, G. Lightbody, and Y. Yi "FPGA-based Implementation of Signal Processing Systems" J. Wiley and Sons, Ltd., Pub. 2008. 364 p. 2. "FPGA Implementations of Neural Networks". A.R. Omondi, and J.C. Rajapakse, Eds. Springer. 2006. 360 p. 3. K. Yoshikawaa, N. Iwanagaa, and A. Yamawaki "Development of Fixed-point Square Root Operation for High-level Synthesis". Proc. 2nd Int. Conf. on Industrial Application Engineering. 2014. pp. 16–20. 4. T.C. Chen "Automatic computations of exponentials, logarithms, ratios and square roots". IBM J. Res. and Develop. 1972. №4. pp. 380–388.

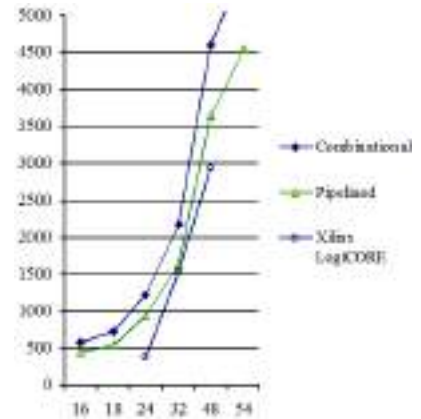


Figure 1. Hardware volume of the $\sqrt{x}$ IP core

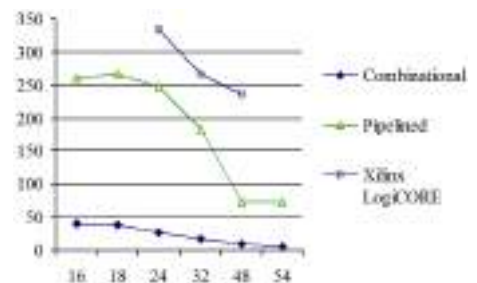


Figure 2. Clock frequency, MHz, of the $\sqrt{x}$ IP core depending on the bit width

Vinogradov Ju.N., Sergiyenko A.M., Quadir S.H.

Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

Minimized FIR Filter Design Implemented in FPGA

Introduction. The field programmable gate array (FPGA) is widely used for the high-speed digital signal processing. FPGA architecture is adapted for the effective finite impulse response (FIR) filter implementation. For this purpose the Xilinx FPGAs contain the DSP48 blocks, each of them is intended for a single filter stage calculation in the pipelined mode. But the filter length is limited by the FPGA volume and by the number of the DSP48 blocks in a single column of the chip. So, the length of the FIR filter which is generated by the Xilinx Coregen tool for the Spartan-6 devices is limited by the numbers from 8 to 48 [1]. And the more this number the more expensive the chip is. Besides, when FPGA is used only for the filtering, then the configured hardware like look-up tables (LUTs), registers is underloaded and is used ineffectively.

There are many methods of the multiplier-less FIR filter implementation. The most of them use the constant coefficient multipliers (CCMs), which have the minimized hardware volume. They are widely used in FPGAs for decades. But no one of them uses the hardware multipliers like the DSP48 block [2–4]. In the presentation, a new method of the FIR filter design is proposed which utilizes both hardware multipliers and CCMs. It provides both the increased filter length and minimized hardware providing the high throughput.

Filter structure. The FIR filter, which is intended for the Xilinx FPGA implementation, has the well-known systolic structure [1]. The synchronous dataflow graph (SDF) of the k -staged systolic structure is illustrated by Fig. 1,a, and its symbol is shown in Fig. 1,b. Here x_i , y_i are the input and output data, the circle, triangle, and bar represent addition, multiplication to the coefficient and delay to a single clock cycle, respectively. This graph is mapped to the respective structure by the one-to-one mapping providing the high pipelined computations with the maximized clock frequency. So, such SDF represents the filter structure as well.

The FIR filter coefficient set is named as an impulse response $H(n)$. It is proposed to implement the FIR filter from three units, which calculate the convolution to the left $H_1(n_l)$, middle $H_2(n_m)$, and right $H_3(n_r)$ subsets of the coefficients, $n = n_l + n_m + n_r$. The respective SDF is shown in Fig. 2.

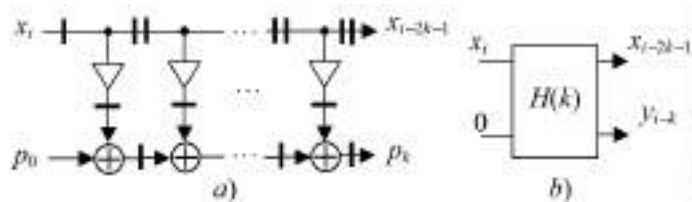


Figure 1. SDF of the systolic FIR filter (a), and its symbol (b)

In the most of cases, the bit width of the coefficients $H_1(n_l)$, and $H_3(n_r)$ is much less, than the bit width of the coefficients $H_2(n_m)$. Therefore, it is favorably to implement the filter part, marked in Fig. 2 as $H_2(n_m)$, using the DSP48 blocks, and others parts using CCMs. In the last situation, the filter stage can be implemented as the small tree of adders with the small bit width. Such a tree can have both high throughput and small hardware volume.

Due to this filter schema, the third filter unit must have the increased adder bit width to preserve the low level of the truncation errors. This can decrease the filter performance dramatically, because the long adder, based on LUTs, is much slower than the respective adder in the DSP48 block. To minimize this disadvantage, the improved SDF is proposed, which is illus-

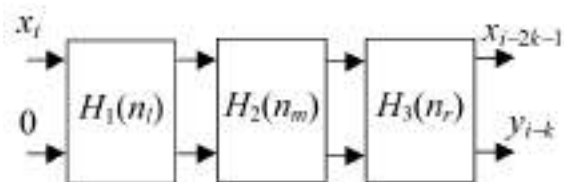


Figure 2. Modified SDF

trated by Fig. 3. The proposed filter contains two parts $H_1(n_l)$, and $H_3(n_r)$ with the equal bit width, which is much less than the bit width of the DSP48 unit. The filter result is formed in the separate adder with the increased bit width. As a result, the designed FIR filter contains n_m DSP48 blocks, which can be much less than the filter length n .

Experimental results. The design of FIR filters of the order of $n = 35$ for the input, output data, and coefficient bit width 16 was considered. The CCMs were built using the canonical signed digit coefficient representation [2–4], and the adder tree. The FIR filter design experience shows that about a half of the 16-bit filter coefficients can be represented in the canonical signed digit form using only four nonzero digits. For example, the filter stage, which performs the multiplication of x to 0.000000101101110_2 and addition of the sum of products p_i from the previous filter stage, looks like the following:

$$p_{i+1} = p_i + 2^{-6}x - 2^{-8}x - 2^{-11}x - 2^{-14}x.$$

The modern FPGAs, which have the six input LUTs, support the design of the high-speed three input adders. So, such CCM with the addition of the sum of products is implemented in the two staged pipelined network of LUTs providing the maximum clock frequency. Under these conditions, the first and third filter units (see Fig. 2, 3) have both minimized hardware volume and high throughput.

The DFGs in Fig 1, 2, 3 were described by the VHDL language. Then, the filter projects were synthesized for the Xilinx Spartan-6 FPGA. The results of the synthesis for some low pass filter of the 33-rd order are shown in Table 1. To estimate the whole hardware volume Q , it was considered, that an 18-bit multiplier in the FPGA requires 200 equivalent LUTs for its implementation.

The Table 1 analysis shows that the filter structure based on DSP48 blocks, provides the maximum sampling clock frequency f_C at the cost of the hardware volume. The second combined structure gives lower value f_C , but it provides much less hardware volume and the better ratio of throughput to hardware volume f_C/Q .

Conclusion. A method of the systolic type FIR filter design for FPGA is proposed, which utilizes both built-in multipliers and application specific constant coefficient multipliers. The method provides both the increased filter length and minimized hardware volume providing the high throughput. The derived filter structure has approximately in two times fewer hardware multipliers than the usual filter generated by the tool like Xilinx Coregen. The resulting filter is represented by the VHDL program and therefore, it can be effectively implemented in FPGA of any type and manufacturer.

References. 1. “Spartan-6 FPGA DSP48A1 Slice User Guide”. UG389, v1.2, Xilinx Inc. May 29, 2014. – 46 p. 2. U. Meyer-Baese, “Digital Signal Processing with Field Programmable Gate Arrays”. Springer, 4-th Ed. 2014, 930 p. 3. A. Sergiyenko, V. Vasyliencko, O. Maslennikov, “FIR filter soft core generator”, Prace IV Konferencji Krajowej “Reprogramowalne układy cyfrowe”, RUC’2001. Szczecin, Poland. 2001. – P. 167–172. 4. M. Kumm, “Multiple Constant Multiplication Optimizations for Field Programmable Gate Arrays”. Springer, 2016. – 206 p.

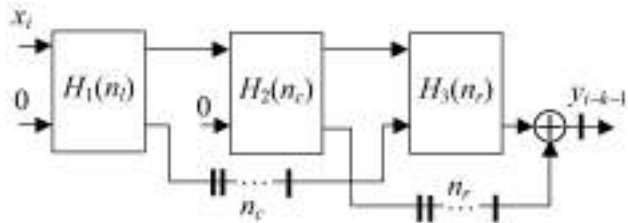


Figure 3. Improved SDF

Table 1. Parameters of the FIR filters implemented in Xilinx Spartan-6 FPGA

Structure	LUTs	Registers	DSP48	Q	f_C , MHz	f_C/Q
Systolic	0	0	33	33	390	11.8
Modified	772	1538	17	20.9	146	7.0
Improved	914	1639	17	21.6	267	12.4

Галета П.О.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Застосування технології Blockchain у Mesh мережах на прикладі розподіленого DNS-сервісу

Сьогодні активного розвитку набули peer-to-peer технології, поступово відвойовуючи свою частину ринку у традиційних систем з централізованою структурою. Широкого розповсюдження набули рішення у сфері файлообміну (BitTorrent, Gnutella, eDonkey, KaZaA і т.п.), електронних платіжних систем (Bitcoin, Litecoin і т.п.). Дані тенденції не обійшли стороною і комп'ютерні мережі. Так, Mesh – є переосмисленням класичного підходу до створення мереж у контексті однорангових систем. Через свою архітектуру, у процесі організації Mesh мереж з'являються нові, нехарактерні для централізованих мереж, проблеми, для вирішення яких може бути застосована технологія Blockchain. Однією з таких проблем є проблема створення розподіленого DNS-сервісу.

Аналіз проблеми. Mesh-мережі є peer-to-peer технологією, тому на відміну від централізованих мереж, використання в них класичних засобів для організації збереження даних (Mongo, PostgreSQL, Cassandra, і т.п.) не є можливим. Іншою ж проблемою у організації систем заснованих на Mesh технології є необхідність у досягненні консенсусу у системі. Наприклад, розглядаючи цю проблему у контексті побудови DNS-сервісу, необхідно досягнути виконання умови “У системі не можуть існувати два сервіси з однаковою адресою”.

Технологія Mesh. Mesh мережа (Сітчаста мережа) — топологія мережі, в якій кожен вузол (називається вузлом меш) передає дані в мережі [2]. Всі вузли співпрацюють у розподілі даних в мережі. Mesh-мережа – це однорангова мережа, пристрої якої (Mesh-станції, MP, Mesh-Points) володіють функціями маршрутизатора і здатні використовувати різні шляхи для пересилки пакету. Ця технологія стає особливо необхідною за відсутності дротової інфраструктури для з'єднання станцій. В цьому випадку пакети пересилаються від однієї Mesh-станції до іншої до досягнення шлюзу з дротовою мережею. Для більшої надійності станція може мати більш, ніж одну сусідню Mesh-станцію. Кожен вузол в ній має такі ж повноваження як і всі інші, грубо кажучи – всі вузли в мережі рівні. Основними перевагами Mesh мереж є легкість у розгортанні, низька ціна користування, анонімність та захищеність у мережі. Основними сферами використання Mesh мереж є IoT (комунікація між пристроями), військовий сектор (швидке розгортання мереж в умовах бойових дій чи стихійного лиха) та громадські мережі (мережі у віддалених районах).

Технологія Blockchain. Блокчейн – це журнал з фактами, що реплікується на кілька комп'ютерів, об'єднаних в мережу рівноправних вузлів (P2P) [1]. Фактами може бути що завгодно, від фінансових операцій та до підписання контенту. Члени мережі – анонімні особи, звані вузлами. Всі комунікації всередині мережі використовують криптографію, щоб надійно ідентифікувати відправника і одержувача. Коли вузол хоче додати факт в журнал, в мережі формується консенсус, щоб визначити, де цей факт повинен з'явитися в журналі. Цей консенсус називається блоком. Тобто блокчейн можна розглядати як (слабо) синхронізовану базу даних, що реплікується стільки ж раз, скільки вузлів в мережі, або як суперкомп'ютер, утворений комплексом всіх CPU / GPU вузлів, що входять до нього. Цей суперкомп'ютер можна використовувати для зберігання і обробки даних, тобто також, як можна використовувати віддалений API. Відмінність тільки в тому, що не потрібно створювати бекенд, можна бути впевненим, що дані надійно захищені і обробляються в мережі належним чином.

Створення розподіленого DNS сервісу. Для побудови DNS сервісу буде використаний Ethereum [3] – платформа для створення децентралізованих онлайн-сервісів на базі блокчейна. Ця платформа була обрана через наявність у ній контрактів- міні-програм, що розташовуються в середині блоків ланцюга блокчейну, та дозволяють змінювати стан даних у блокчейні лише при виконанні деяких умов. Кожен контракт несе в собі міні-базу даних, надає методи для зміни її даних. Оскільки контракти реплікуються по всіх вузлах, то і їх бази даних теж. Кожен

раз, коли користувач викликає метод з контракту і, відповідно, змінює дані, ця команда реплікується і повторюється всією мережею. Це дозволяє створити розподілений консенсус для виконання обіцянок.

Основними вимогами до DNS сервісу є можливість виконання наступних операцій:

1. Можливість реєструвати нову адресу
2. Можливість отримувати ір адресу, знаючі символічну адресу
3. Можливість передавати права власності на адресу іншому користувачу

У якості одиниці, що зберігається у блокчейні, використана key/value пара, де у якості ключа використовується символічна адреса, а у якості значення – пара ір адреса, ідентифікатор власника адреси. Для виконання перших двох вимог поставлених перед DNS сервісом використовується Ethereum API, що дозволяє виконувати запис та читання до блокчейну. Нова адреса вважається зареєстрованою, за умови, що вона ще не була присутня у ланцюгу, а власником адреси стає користувач системи, з облікового запису якого відбувся запис. Для досягнення третьої вимоги використовуються смарт-контракти, що змінюють власника адреси у разі, якщо з облікового запису покупця відбувся грошовий переказ на обліковий запис власника адреси із сумою, зазначеною власником у контракті.

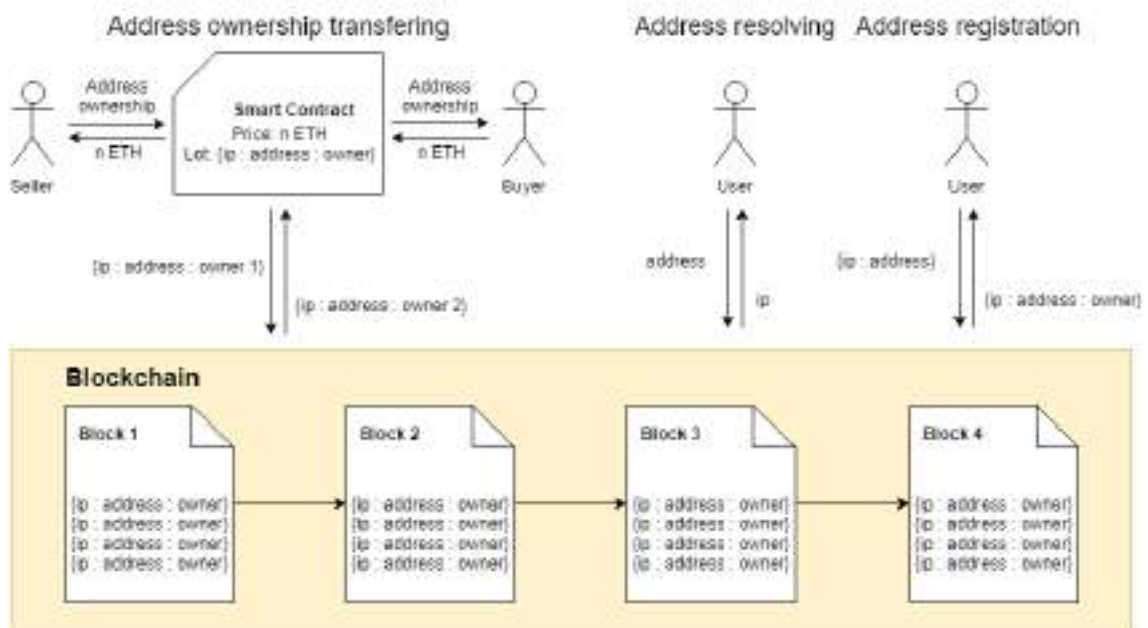


Рис. 1. Діаграма взаємодії користувачів з DNS сервісом на основі блокчейну

Висновки. Застосування технології Blockchain у Mesh мережах представляється зручним і надійним рішенням. Адже у ряді задач, серед яких і задача створення DNS сервісу, технологія Blockchain вирішує основні проблеми організації сховища (необхідність у децентралізованості, захищеність від підробки), а більшість її недоліків (повільність додавання нових записів, надмірна реплікація) несуттєві у даному контексті.

Література. 1. Marco Iansiti and Karim R. Lakhani The Truth About Blockchain // Harvard Business Review, January–February 2017, p. 118–127. 2. Everyone is a node: How Wi-Fi Mesh Networking work by Jerry Hildenbrand // IEEE Journal on Selected Areas in Communications, 2016, p.26–33. 3. Ethereum – blockchain APP platform [Електронний ресурс] //Режим доступу: – <https://www.ethereum.org>. Дата доступу 30.03.2018.

Івченко Д.А.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Використання Neo4j і GraphQL на базі GraphDB для створення мікросервісу

Вступ. Реляційні бази даних були потужними програмними додатками з 80-х років і залишаються зараз. Вони зберігають високо структуровані дані в таблицях з зумовленими стовпцями певних типів і безліччю рядків одного і того ж типу інформації і, частково завдяки жорсткості їх організації, вимагають від розробників додатків строго структурувати дані.

У реляційних базах даних посилання на інші рядки і таблиці вказуються шляхом посилання на їх (первинні) атрибути ключів через стовпці зовнішнього ключа. Це можна виконати з обмеженнями, але тільки тоді, коли посилання ніколи не є обов'язковим. Вибірка обчислюється під час запиту, зіставляючи первинні і зовнішні ключі багатьох рядків таблиць, які підлягають об'єднанню. Ці операції розраховані на обчислення і пам'ять, і мають експонентну вартість.

Якщо ви використовуєте зв'язки «багато-до-багатьох», вам необхідно представити таблицю JOIN (або таблицю об'єднань), яка містить зовнішні ключі обох таблиць-учасників, що ще більше збільшує витрати на операції об'єднання. Ці дорогі операції зазвичай вирішуються шляхом денормалізації даних, щоб зменшити кількість необхідних вибірок. Хоча не кожен варіант використання підходить для такого типу точних моделей даних, в минулому відсутність життєздатних альтернатив і велика підтримка реляційних баз даних ускладнювали створення альтернатив.

Від реляційних баз даних до графових. Зв'язки першокласних мереж у моделі графових баз даних, на відміну від інших систем управління базами даних, які вимагають від нас встановлення зв'язків між об'єктами, що використовують спеціальні властивості, такі як зовнішні ключі. Об'єднавши прості абстракції вузлів і зв'язків в пов'язаних структурах, графові бази даних дозволяють нам створювати складні моделі, які тісно пов'язані з нашою проблемною областю. У певному графові бази даних схожі на наступне покоління реляційних баз даних, але з підтримкою першого класу для «зв'язків» або з неявними сполуками, зазначеними за допомогою зовнішніх ключів в традиційних реляційних базах даних.

Кожен вузол в моделі графових баз даних безпосередньо і фізично містить список взаємопов'язаних записів, які представляють його зв'язок з іншими вузлами. Ці взаємопов'язані записи організовані за типом і напрямком, і можуть містити додаткові атрибути. Всякий раз, коли ви запускаєте еквівалент операції JOIN, база даних просто використовує цей список і має безпосередній доступ до пов'язаних вузлів, що усуває необхідність в дорогому обчисленні пошуку / зіставлення.

Зміст графової бази даних. Дуже просто, графова база даних – це база даних, призначена для обробки зв'язків між даними як першокласна мережа в моделі даних.

Ми живемо в взаємозв'язаному світі. В ньому немає ізольованих частин інформації, але багато, пов'язаних структур навколо нас. Лише база даних, яка спочатку підтримує зв'язки, здатна ефективно зберігати, обробляти і запитувати з'єднання. У той час як інші бази даних обчислюють їх під час запиту через дорогі операції JOIN.

Доступ до вузлів і зв'язків у графовій базі даних – це ефективна неперервна операція, яка дозволяє швидко зчитати мільйони зв'язків в секунду на вузол. Незалежно від загального розміру вашого набору даних, графові бази даних перевершують управління високопов'язаними даними і складними запитами. Маючи тільки шаблон і набір вихідних точок, графові бази даних досліджують зв'язки навколо початкових точок – збір та його узагальнення інформації з мільйонів вузлів і зв'язків – залишаючи мільярди поза периметром пошуку недоторканими.

Модель властивостей графа. Якщо ви коли-небудь працювали з об'єктною моделлю або діаграмою зв'язків сутностей, модель з мітками властивостей буде здаватися вам знайомою.

Вузли – це об’єкти на графіку. Вони можуть містити будь-яку кількість атрибутів (пари ключ-значення), так звані властивості. Вузли можуть бути позначені ярликами, які представляють їхні різні ролі. На додаток до контекстуалізації властивостей вузла і зв’язків мітки можуть також служити для прикріплення інформації про індекс метаданих або обмеження до певних вузлів.

Зв’язки забезпечують направлення, іменовані, семантично релевантні зв’язку між двома вузовими об’єктами. Зв’язок завжди має напрямок, тип, початковий вузол і кінцевий вузол. Як і вузли, зв’язки також можуть мати властивості. У більшості випадків відносини мають кількісні характеристики, такі як ваги, витрати, відстані, рейтинги, інтервали часу або стійкість. Оскільки ссылки зберігаються ефективно, два вузла можуть ділитися будь-яким числом або типом зв’язків, не жертвуючи продуктивністю. Зверніть увагу, що, хоча вони спрямовані, зв’язки завжди можуть ефективно переміщатися в будь-якому напрямку [1].

Будівельні блоки властивостей графа. У базі даних графа є одне основне узгоджене правило: «Жодних зламаних зв’язків». Оскільки зв’язки завжди мають початковий і кінцевий вузол, ви не можете видалити вузол, не видаляючи також пов’язані з ним зв’язки. Ви також можете завжди припускати, що існує відношення ніколи не вкаже на неіснуючу кінцеву точку.

Вступ до Neo4j. Neo4j є відкритою базою даних NoSQL з відкритим вихідним кодом, яка забезпечує сумісний з ACID транзакційний сервер для ваших додатків. Розробки ведуться з 2003 року, він був загальнодоступним з 2007 року. Вихідний код, написаний на Java і Scala, доступний на GitHub [2].

Neo4j використовується сьогодні тисячами компаній і організацій практично у всіх галузях, включаючи фінансові послуги, уряд, енергетику, технології, роздрібну торгівлю і виробництво. Сотні розробників і архітекторів в цих галузях – це сертифіковані професіонали Neo4j.

Деякі особливості роблять Neo4j дуже популярним серед розробників, архітекторів і адміністраторів баз даних:

- Cypher, мова декларативних запитів, аналогічний SQL, але оптимізована для графіків. Тепер використовується іншими базами даних, такими як графік SAP HANA Graph і Redis через проект openCypher.
- Дозволяє масштабувати до мільярдів вузлів на середньому апаратному забезпеченні.
- Гнучка схема властивостей графів, яка може адаптуватися з плином часу, дозволяючи згодом матеріалізувати і використовувати нові зв’язки для «швидкого доступу» до вузлів.
- Драйвери для популярних мов програмування, включаючи Java, JavaScript, .NET і Python.

Базова реалізація. Описання структур для графу буде проводитися на мові Cypher [3].

Задача, яку буде вирішено: збереження медичних даних про пацієнта (особисті дані та результатів аналізів). На рис. 1 наведено код на мові Cypher, який необхідний для створення записів про пацієнта в графовій базі даних.

```

1 CREATE (Patient1:PATIENT_INSTANCE { was_born:'1988', has_location:'Kyiv', has_name:'Dima'})
2 CREATE (PatientData:PATIENT_DATA { has_timestamp: '2018-02-28-12:00'})
3
4 CREATE (Patient_data_1_01:PATIENT_DATA_ELEMENT { has_value: '88', of_type: 'BLOOD_PRESSURE_DIASTOLIC'})
5 CREATE (Patient_data_1_02:PATIENT_DATA_ELEMENT { has_value: '120', of_type: 'BLOOD_PRESSURE_SISTOLIC'})
6
7 CREATE
8   (PatientData)-[:HAS_ELEMENT]->(Patient_data_1_01),
9   (PatientData)-[:HAS_ELEMENT]->(Patient_data_1_02),
10  (PatientData)-[:CREATED_BY]->(Patient1);

```

Рис. 1. Створення записів про пацієнта в графовій базі даних

Напишемо вибірку, де вказуємо, які саме ті ноди, які ми хочемо отримати.

```
1 MATCH (e:PATIENT_DATA_ELEMENT)-[:HAS_ELEMENT]-(d)-[CREATED_BY]->(p:PATIENT_INSTANCE)
2 RETURN d, e, p;
```

Рис. 2. Створення записів про пацієнта в графовій базі даних

Результуючий граф зображено на рис. 3.



Рис. 3. Результуючий граф

Рішення з використанням graphql. На початку 2018 року компанія Neo4j випустила розширення під назвою neo4j-graphql [4], яке дозволило виконувати запити в графову базу даних за допомогою GraphQL формату [5].

Далі наведено повний перелік його функцій:

- Автогенерування всіх типів запитів для кожного елементу згідно до схеми;
- Кожна властивість елемента доступна, як query аргумент;
- Надано можливість сортувати результати при їх отриманні;
- Також є вбудована курсорна пагінація;
- Наявна валідація даних відповідно до схеми.

Опишемо схему бази використовуючи типи GraphQL.

```
1+ type PATIENT_INSTANCE {
2   name: String!
3   timestamp: String!
4   born: String!
5   data: [PATIENT_DATA] @relation(name:"CREATED_BY", direction:IN)
6 }
7 type PATIENT_DATA {
8   created: String!
9   elements: [PATIENT_DATA_ELEMENT] @relation(name:"HAS_ELEMENT", direction:IN)
10 }
11 type PATIENT_DATA_ELEMENT {
12   value: String!,
13   type: [PATIENT_DATA_ELEMENT_TYPE]!
14 }
15
16 enum PATIENT_DATA_ELEMENT_TYPE {
17   BLOOD_PRESSURE_DIASTOLIC
18   BLOOD_PRESSURE_SISTOLIC
19 }
20
```

Рис. 4. Типи елементів з описами полів, що в них наявні

Згенеруємо графову базу за допомогою команди: `CALL graphql.idl('schema-text')`. Далі отрим-

маємо базову схему в вигляді графу, запустивши команду: `call graphql.schema()`. Результат показано на рис. 5.



Рис. 5. Графове представлення бази

Автогенерація мутацій. Додатково було згенеровано мутації у відповідності до схем. Кожна з мутацій повертає оновлену статистику. Так для нашої схеми будуть доступні наступні мутації.

```

createPATIENT_INSTANCE(name: String!, timestamp: String! born: String!) : String
updatePATIENT_INSTANCE(name: String!, timestamp: String! born: String!) : String
deletePATIENT_INSTANCE(name: String!) : String
  
```

Рис. 6. Отримані мутації

Проста мутація, яку вже можна буде викликати на сайті за допомогою Apollo Client, буде виглядати наступним чином:

```

mutation {
  createPATIENT_INSTANCE(name: "Denial", born: "1994")
}
  
```

Рис. 7. Мутація для створення пацієнта

Висновки. В даній статті було розглянуто переваги графових баз даних над реляційними. Спроектвана та реалізована графова база даних на основі новітніх технологій: GraphDB, Neo4j, Cypher, GraphQL, Apollo Client, React. Отримано повноцінний сервіс для роботи з графовою базою даних за допомогою GraphQL запитів.

Література. 1. Графова база даних https://en.wikipedia.org/wiki/Graph_database. 2. Веб-сайт Neo4j <https://neo4j.com>. 3. Опис формату Cypher <https://neo4j.com/developer/cypher/>. 4. Репозиторій neo4j-graphql <https://github.com/neo4j-graphql/neo4j-graphql>. 5. Опис концепції GraphQL <https://en.wikipedia.org/wiki/GraphQL>.

Кирюша Б.А., Колінько А.М.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Взаємодія мультиагентної системи IoT девайсів з виділеним сервером

Існує 2 підходи для регулювання роботи IoT девайсів: архітектура з виділеним сервером та архітектура однорангових агентів. Вибір підходу значно залежить від цілей створення системи, наявних ресурсів, вимог до масштабованості системи, її безпеки, надійності тощо.

Цілі створення системи. Розподілені IoT системи створюються для

1. Збору і аналізу даних з навколишнього середовища;
2. Обробки даних з різних елементів в реальному часі;
3. Делегування серверу обчислень;
4. Віддаленого керування елементами системи.

Збір даних зазвичай відбувається давачами підключеними до мікроконтролеру що потім передає інформацію на виділений сервер, до хмари чи іншим мікроконтролерам.

Збір і аналіз даних з навколишнього середовища. В першому випадку (архітектура з виділеним сервером) наявність підключення до мережі Інтернет є необов'язковою. За відсутності підключення система може представляти собою локальну мережу. Це підвищує захищеність системи від зовнішнього втручання, але зменшує гнучкість та вимагає окремого компоненту (наприклад, шлюзу) для взаємодії з системами IoT девайсів. Якщо ж підключення до мережі наявне, система може бути легко масштабована з додаванням як окремих IoT девайсів, так і налаштуванням взаємодії з іншими системами [1]. Також за наявності підключення до мережі Інтернет можна зберігати резервні копії в хмарі.

Використання хмарного сховища даних також спрощує роботу з великими обсягами даних за допомогою окремих баз даних та таких технологій як Hadoop та Spark [1].

За архітектури однорангових агентів, власне, агентами можуть виступати мікроконтролери. В цьому випадку до мікроконтролерів ставляться вищі вимоги до наявності обчислюваних ресурсів, оскільки в даному випадку мікроконтролери повинні не тільки збирати і відправляти дані, а ще і отримувати запити від інших мікроконтролерів, відправляти відповіді на запити, відрізняти дублікати повідомлень, опрацьовувати втрату повідомлень (напр. повторювати запит за відсутності відповіді по timeout чи за алгоритмом плаваючого вікна) і т.д. [2]. Тобто обчислення, які виконував сервер, тепер перерозподіляються між агентами. Це робить систему надійнішою (при відмові роботи одного мікроконтролеру його функції можуть бути частково передані іншим якщо двачі підключені до шини, з якої зчитують дані декілька мікроконтролерів), але така архітектура дещо складніша у налаштуванні.

Обробка даних в реальному часі. Для обробки даних в реальному часі ставлять підвищені вимоги до надійності та швидкодії системи в цілому, проте немає вимоги гарантованої доставки повідомлення, адже виміри, отримані протягом попереднього часового інтервалу, вже не настільки актуальні, як виміри нинішнього часового інтервалу. Внаслідок підвищених вимог до швидкодії системи використання однорангової архітектури обмежене, і може застосовуватись лише в випадку невеликої кількості агентів (інакше може виникнути хаотична передача запитів/відповідей та значна розсинхронізація між віддаленими мікроконтролерами). В даному випадку можна використовувати архітектуру з виділеним сервером чи організовувати передачу даних до хмари по високошвидкісному каналу.

Делегування обчислень. Делегування обчислень передбачає делегування необхідних обчислень до виділеного серверу, хмари чи grid-системи, отримання та аналіз результатів обчислення. Можлива структура з виділеним сервером, який одночасно слугує для збору запитів та даних для обчислення від мікроконтролерів, формування та передачі задачі до grid-середовища, отримання результатів обчислення, їх аналіз та передавання відповідних повідомлень до IoT девайсів.

Віддалене керування. Віддалене керування системою IoT девайсів може бути як основною ціллю створення системи (напр. для керування актуаторами), так і додатковою функцією однієї з вищеперечислених структур (в випадку основної структури з архітектурою однорангових агентів буде утворюватись структура зі змішаною архітектурою). Реалізація даної функції передбачає наявність виділеного сервера підключено до мережі Інтернет, що може керувати мікроконтролерами та має інтерфейс для віддаленого доступу.

Табл. 1. Порівняльна таблиця

	Збір і аналіз даних з навколишнього середовища	Обробка даних в реальному часі	Делегування обчислень	Віддалене керування
Архітектура з виділеним сервером	Підтримується	Підтримується	Підтримується	Підтримується
Архітектура однорангових агентів	Підтримується	Частково підтримується	Не підтримується	Утворює змішану архітектуру
Вимоги до надійності	Варіюються (залежить від природи даних)	Високі	Середні	Середні
Вимоги до безпеки	Варіюються (залежить від природи даних)	Високі	Середні	Високі

Гібридна архітектура. Архітектура з виділеним сервером та одноранговими агентами – два крайні підходи. В першому випадку мікроконтролери виконують функції для передачі даних серверу і реакції на отримані від серверу повідомлення. Тобто до обчислювальної здатності мікроконтролерів ставляться мінімальні вимоги як до керованих IoT девайсів, а необхідний аналіз і обчислення проводить виділений сервер. З іншого боку, за однорангової архітектури всі обчислення та аналіз виконуються самими агентами, які також повинні керувати поширенням повідомлень в мережі, і, як наслідок, мікроконтролери повинні мати достатню обчислювальну потужність для виконання таких функцій. Таким чином кожен агент володіє певною необхідною інформацією про систему.

Однак можливі архітектурні рішення не обмежуються лише цими двома підходами: можливі також безліч перехідних варіантів.

- виділити серед однорангових агентів декілька “ключових” агентів, що будуть виконувати специфічні функції (напр. отримувати повідомлення від групи інших агентів, формувати пакет та відправляти його до хмари);
- розділити агентів по групах за функціями (напр. виділити мікроконтролери за регіоном чи показником для вимірювання: температура, тиск, освітленість тощо) та виділити цим групам обчислювальні вузли, які і будуть взаємодіяти з виділеним сервером;
- організувати опитування випадковим чином обраних агентів виділеним сервером для моніторингу роботи системи;
- організувати передачу агентами високопріоритетних повідомлень виділеному серверу про виміри системи, що знаходяться близько до межі допустимих значень; тощо.

Література. 1. Fremantle, P. A reference architecture for the internet of things. [ebook] – 2015. Available at: https://wso2.com/download/getfile/wso2_whitepaper_a-reference-architecture-for-the-internet-of-things.pdf [Accessed 8 Apr. 2018]. 2. Singh, M. and Chopra, A. (2017). The Internet of Things and Multiagent Systems: Decentralized Intelligence in Distributed Computing. 2017 IEEE 37th International Conference on Distributed Computing Systems (ICDCS).

Круш І.В., Михалько В.Г.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Адаптивний точковий пошук з використанням методів машинного навчання

За останнє десятиліття спостерігається експоненційне зростання кількості даних у мережі Інтернет, яке буде продовжуватись і надалі. Перед системами, котрі зберігають та обробляють ці дані, виникає велика кількість технічних викликів, в тому числі можливості швидкого пошуку та вставки даних [1].

Класичним підходом до вирішення задачі швидкого пошуку є використання індексних структур даних. Зокрема, для діапазонного пошуку найкраще себе зарекомендували В-дерева, які дозволяють виконувати цю операцію за $O(\log n)$. Коли необхідно виконати точковий пошук, то зазвичай використовують хеш-таблиці, які працюють за час $O(1)$. Внаслідок того, що індексні структури є важливими елементами СУБД, вони були значно покращені і оптимізовані за минулі десятиріччя. Незважаючи на це, у них є один суттєвий недолік – вони не враховують реальний розподіл даних [2].

Альтернативним підходом до пошуку даних може бути використання методів машинного навчання для побудови векторних просторів даних та організація швидкого пошуку у них. За допомогою спеціальних алгоритмів можна представити запит як елемент простору даних та побудувати алгоритм пошуку n найближчих до нього елементів простору. Безумовно, можна об'єднати два підходи – використати методи машинного навчання для прискорення або заміни класичних індексних структур даних.

Значною проблемою застосування методів машинного навчання є те, що вони обчислювально дуже інтенсивні, адже базовані на операціях множення матриць. На процесорі складність цієї операції перевищує $O(n^2)$, однак задача легко паралеліється. Відповідно, множення матриць вигідніше виконувати на GPU. Також, якщо взяти до уваги значний розвиток баз даних, що працюють на GPU (до прикладу, MapD Core), то стає зрозуміло, що це дає широкі можливості застосування підходів машинного навчання у тісній інтеграції із СУБД [2, 3].

В останні роки почали з'являтися перші наукові роботи по використанню машинного навчання для заміни класичних індексних структур. Зокрема, в 2017 році дослідники з Google показали, що такі підходи навіть при використанні звичайних процесорів можуть давати кращі результати [4].

Розглянемо застосування такого підходу для точкового пошуку. Класична індексна структура для точкового пошуку – це хеш-таблиця. При цьому, основною її проблемою є колізії ключів – тобто ситуації, коли для декількох різних ключів призначається один слот в таблиці. Базуючись на парадоксі днів народжень, можна порахувати, що очікувана кількість конфліктів для звичайних хеш-таблиць дорівнює 33% [4, 5].

Двома основними стратегіями вирішення проблеми колізій є:

- використання списків;
- виконання додаткових обчислень для визначення вільного слота.

При цьому, при використанні обох цих стратегій ефективність пошуку і вставки зменшується [2].

До вирішення цієї задачі можна підійти інакшим чином: натренувати модель за допомогою методів машинного навчання таким чином, щоб вона для різних ключів ставила у відповідність унікальні позиції в таблиці (рис. 1).

Одним із методів, яким цього можна досягти – побудувати кумулятивну функцію розподілу

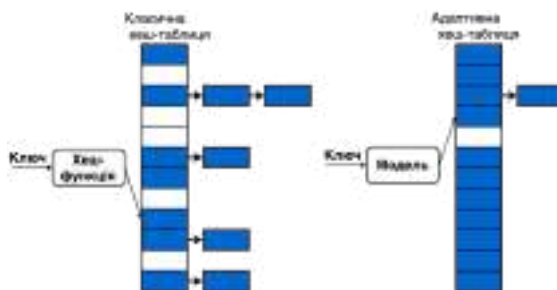


Рис. 1. Відмінність роботи класичного хеш-індексу від навченого хеш-індексу

ключів і натренувати модель, яка вивчить цю функцію. Для цього можна використовувати різні моделі, починаючи від простої лінійної регресії і закінчуючи глибокими нейронними мережами.

Наведемо приклад псевдокоду для тренування моделі – двошарової нейронної мережі з активацією у вигляді функції Leaky ReLU, що дозволяє додати у модель невелику кількість негативних активацій, які працюють як регуляризатор. Цільовою функцією буде крос-ентропія, що намагатиметься розподілити дані з N елементами рівномірно у N слотів.

```
model = Sequential(
    Linear(input_size, hidden_size),
    LeakyReLU(),
    Linear(hidden_size, hidden_size * 2),
    LeakyReLU(),
    Linear(hidden_size * 2, out_size),
    Softmax())
loss_fn = CrossEntropyLoss()
optimizer = optim.Adam(model.parameters(), lr=learning_rate)
model.train()
```

Модель, що наведена вище, є доволі простою. Внаслідок цього в ній наявні декілька проблем. По-перше, для того, щоб точковий пошук працював правильно, цю модель необхідно перенавчити (overfit) на заданих даних. По-друге, такі моделі можуть вивчити загальний розподіл, але в них будуть виникати проблеми на “останній милі”, тобто під час визначення точної позиції в таблиці. Тому для того, щоб модель повністю вивчила розподіл даних, в ній має бути доволі багато параметрів [4].

Один із способів покращення моделі – зробити її ієрархічною: спочатку тренувати верхні шари, які будуть відтворювати загальний розподіл даних, а потім рекурсивно переходити до нижніх шарів. При цьому, якщо дані в нижніх шарах не сильно піддаються навчанню, можна використовувати і більш прості алгоритми для локального пошуку – наприклад, звичайний перебір даних [4].

Інший підхід до тренування – використання машинного навчання без учителя. В цьому випадку для цільової функції порядок ключів в таблиці неважливий, так як її значення буде залежати тільки від кількості колізій. Цей підхід є складнішим для реалізації, але він може дозволити моделі вивчити дані за допомогою меншої кількості параметрів.

Варто зауважити, що при використанні описаних вище підходів на звичайних процесорах буде складно досягти пришвидшення точкового пошуку. Однак, якщо більш важливим є оптимальне використання пам'яті, то в цьому випадку випадку методології пошуку, засновані на використанні машинного навчання, можуть стати у нагоді.

В цілому, навіть на сучасному етапі розвитку апаратного забезпечення, реалізація адаптивного пошуку з використанням технологій машинного навчання часто є виправданою, адже такий підхід дозволяє значно краще враховувати розподіл даних. А зважаючи на те, що в найближчий час очікується значне зростання продуктивності GPU та TPU, використання методологій адаптивного пошуку на базі машинного навчання може стати більш ефективним в порівнянні з класичними індексними структурам.

Література. 1. P. Ffoulkes InsideBIGDATA guide to the intelligent use of big data on an industrial scale. InsideBIGDATA, Massachusetts, 2017. 2. H. Garcia-Molina, J.D. Ullman, J. Widom Database Systems: The Complete Book. Prentice Hall Press, Upper Saddle River, NJ, 2008. 3. T. Mostak. An overview of MapD (massively parallel database). MIT Technical Report, 2013. 4. T. Kraska, A. Beutel, E.H. Chi, J. Dean, N. Polyzotis. The Case for Learned Index Structures [Online]. Available: <https://arxiv.org/abs/1712.01208>. 5. F.H. Mathis A Generalized Birthday Problem. SIAM Review. Society for Industrial and Applied Mathematics, 1991.

Лозінський А.П.

Міжнародний науково-навчальний центр інформаційних технологій та систем НАН та МОН України, Київ, Україна

Багаторівнева архітектура хмарної платформи

Розглядається результати дослідження технологій реалізації основних властивостей хмарних обчислень. Пропонується багаторівнева архітектура хмарної платформи, що відображає ієрархію організації хмарних обчислень від рівня апаратних вузлів до прикладних задач.

Подібно до переходу підприємств на початку цього століття від експлуатації власних електрогенераторів до промислової електромережі, сучасні підприємства розглядають використання хмарних обчислень замість власних серверних та ЦОД [1].

З метою досліджень технологій хмарних обчислень, автором реалізовано модель локальної хмарної платформи (Модель), що здатна постачати хмарні послуги від інфраструктурних рішень до завершених програмних середовищ [2]. Модель досліджується на предмет технологій та методів реалізації основних властивостей хмарних обчислень [3,4]. У її побудові використано програмний продукт з відкритим кодом OpenStack [5].

Багаторівнева архітектура. З метою структурного представлення внутрішнього устрою та функціонування Моделі автором розроблено *багаторівневу архітектуру хмарної платформи* (Архітектура) у складі семи рівнів (рис. 1).

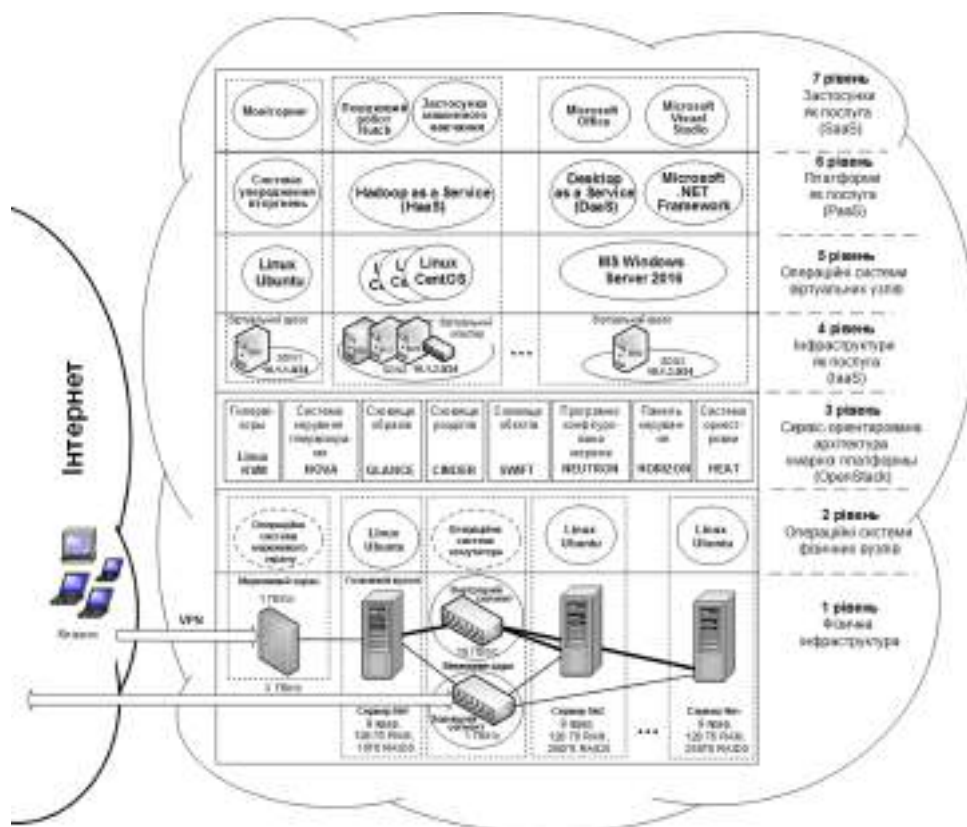


Рис. 1. Багаторівнева архітектура хмарної платформи

Дослідження методів реалізації в Моделі основних властивостей хмарних обчислень виявило, що вони об'єднують в собі наступні теоретичні та інженерні досягнення різних галузей інформаційних технологій: багатоядерні мікропроцесори, швидкості мережевого обладнання

10 Гбіт/с та більше, технології віртуалізації, бази даних, сервіс-орієнтовану архітектуру, композиції сервісів, web-технології, програмно-керовані мережі, автоматизація керування ЦОД, автономні обчислення, принципи GRID- та utility-computing.

Розшарування цього різноманіття технологій на сім функціональних рівнів дозволило аналізувати кожен з них окремо та виявляти наявність впливу та функціональних взаємозв'язків між ними. Кожен рівень служить платформою для побудови іншого рівня, розташованого вище. Архітектура доповнює ієрархію хмарних послуг IaaS, PaaS та SaaS такими загальноприйнятими поняттями як апаратний вузол, операційна система, сервіс та прикладна задача. Приведене на рис. 1 наповнення рівнів Архітектури подано як приклад, відповідно до реалізації створеної автором Моделі. Наповнення функціональних рівнів буде відрізнятися у відповідності до поставлених перед розробником хмарної платформи задач.

Розгляд першого фізичного рівня Архітектури наглядно демонструє можливість необмеженого горизонтального нарощування апаратної платформи, що створює у споживачів хмарних послуг враження відсутності обмежень в ресурсах. Цьому враженню також сприяє виокремлена в окремий третій рівень сервіс-орієнтована архітектура OpenStack, яка обмежує видимість користувачами хмарних послуг нижніх трьох рівнів Архітектури та реалізує об'єднання переваг різних технологій в рамках єдиної хмарної платформи.

Переваги хмарних обчислень. Нижче наведено перелік функціональних можливостей, що доступні споживачам та постачальникам хмарних послуг на сучасному етапі розвитку завдяки об'єднанню переваг різних інформаційних технологій в хмарних обчисленнях.

Переваги користувачів хмарних послуг: а) самообслуговування; б) керування хмарними послугами засобами API; в) оренда віртуальних та апаратних вузлів; г) створення інфраструктурних рішень; д) керування потужністю вузлів; е) надійність зберігання даних.

Переваги постачальників хмарних послуг: а) ефективне навантаження ресурсів технологіями віртуалізації; б) гнучкість керування сервіс-орієнтованою архітектурою; в) високий договірний рівень обслуговування засобами сервісу міграції; г) централізоване керування.

Об'єднання в рамках хмарних обчислень передових технологій, прозорість та гнучкість керування інфраструктурними рішеннями є стимулом постійного впровадження нових функціональних можливостей. Сучасний етап розвитку хмарних платформ залежить від рівня досконалості технологій, що входять до їх складу.

Подібно промисловій електромережі, хмарні обчислення є перспективним напрямком розвитку як інструмент стимулювання сфери інформаційних технологій на державному рівні. Шляхом створення промислового постачання хмарних послуг можливе централізоване обслуговування потреб підприємств державного та комерційного сектору в обчислювальних засобах (послуги IaaS), середовищах обробки даних (послуги PaaS), прикладних програмах (послуги SaaS).

Висновки. Запропонована багаторівнева архітектура хмарної платформи структурує розуміння будови платформ хмарних обчислень. Хмарні обчислення орієнтовані на впровадження передових наукових та технічних розробок. Хмарні платформи мають бути ефективним інструментом стимулювання розвитку сфери інформаційних технологій.

Література. 1. Gens F. The 3rd Platform: Enabling Digital Transformation. November 2013, p. 13. Available: <http://www.idc.com/>. 2. Лозинский А.П., Симахин В.М., Урсатьев А.А. Моделирование технологий обработки больших данных на локальной облачной платформе // Управляющие системы и машины. – 2017. – №3. – С. 6–19. 3. Гриценко В.И. Cloud Computing и облачная модель предоставления ИТ-услуг / В.И. Гриценко, А.А. Урсатьев // КБТ. – 2013. – Вип. 171. – С. 5–19. 4. Mell P., Grance T. The NIST Definition of Cloud Computing. National Institute of Standards and Technology. NIST Special Publication 800-145. Available: <http://nvlpubs.nist.gov/nistpubs/Legacy/SP/nistspecialpublication800-145.pdf>. 5. Open Stack. Open source software for creating private and public clouds. Available: <http://openstack.org/>.

Назарова І.А.

ДВНЗ “Донецький національний технічний університет”, Покровськ, Україна

Масштабованість паралельних розподілених обчислень на основі ізоефективного аналізу

В докладі представлено результати досліджень, присвячених аналізу масштабованості та ефективності паралельних обчислень для багатопроцесорних систем з розподіленою пам'яттю при розв'язанні задачі Коші для систем звичайних диференціальних рівнянь (СЗДР). Проаналізовано характер поведінки динамічних характеристик розроблених методів: часу паралельної реалізації алгоритму, загальних накладних витрат на паралелізм, прискорення, ефективності використання паралельного обладнання, оцінки для максимального значення ступеню паралелізму та кількості застосованих процесорів.

Розробка паралельних додатків в якості своєї мети може мати не тільки зменшення часу виконання, а й забезпечення можливості вирішення завдань більшої розмірності. Здатність паралельного алгоритму ефективно використовувати процесори при збільшенні складності розрахунків є важливою характеристикою паралельних обчислень і називається масштабованістю. Паралельний алгоритм є масштабованим, якщо при зростанні числа процесорів він забезпечує збільшення прискорення при збереженні постійного рівня ефективності використання процесорів [1]. Відомі такі підходи до оцінки масштабованості паралельного алгоритму, як масштабування прискорення, ізошвидкісна міра, рівні умови продуктивності тощо [1, 2]. У даній роботі досліджується масштабованість різних паралельних алгоритмів для паралельних архітектур за умови збереження ефективності при збільшенні розміру проблемної задачі із зростанням числа процесорів, у якості методу оцінки масштабованості обрано математичний апарат ізоефективного аналізу. На основі такого підходу в одному рівнянні лаконічно фіксуються характеристики паралельного алгоритму, методу і параметри паралельної архітектури, на якій реалізовано алгоритм.

Дослідження масштабованості паралельних обчислень проілюстровано на прикладі розв'язання задачі Коші на основі блокових однокрокових схем з вбудованими методами оцінювання локальної апостеріорної похибки для керування кроком інтегрування [3, 4]. Блокові багатоточкові методи вирішення динамічних задач є особливо актуальними, бо добре узгоджуються з архітектурою паралельних обчислювальних систем і не вимагають обчислення значень в проміжних вузлах, що значно підвищує ефективність розрахунків. Дані методи володіють достатніми характеристиками стійкості і є по своїй суті паралельними, оскільки дозволяють отримувати розв'язок одночасно в декількох точках сітки інтегрування.

На відміну від явних методів вирішення СЗДР, реалізація альтернативних засобів оцінки апостеріорної локальної похибки на основі блокових методів пов'язана із рядом особливостей:

- немає відповідних послідовних аналогів, отже, потрібно розробити і обґрунтувати метод оцінки локальної похибки;
- зміна кроку інтегрування можлива лише після виконання обчислень у всіх вузлах поточного k -го блоку;
- при умові незадовільної оцінки локальної похибки практично усі обчислення для точок блоку виявляються даремними (тільки деякі звернення до правої частини СЗДР можуть бути використані знову).

По-перше, у дослідженні використано ідею вкладених форм, що була запропонована для оцінки локальної похибки чисельного рішення СЗДР методами типу Рунге-Кутта. Таким чином, для однокрокових блокових багатоточкових методів розроблено вкладені методи на основі двох різних підходів:

1. комбінації незалежних формул різних порядків точності;
2. комбінації спеціально підібраних формул різних порядків точності на основі методу послідовного підвищення порядку Кривої.

По-друге, для оцінки похибки на кроці використана технологія локальної екстраполяції

Річардсона-Ромберга з блоковим опорним методом невисокого порядку.

Для аналізу ступеню масштабованості розроблених паралельних алгоритмів і їх програмних додатків проводилось дослідження варіації загального часу паралельного виконання на паралельній архітектурі з метою отримання такого числа процесорів p , при якому T_p буде мінімальним, а прискорення S та ефективність E – максимальними. Основні результати дослідження масштабованості отримані у випадках, коли є можливість побудувати так звану функцію ізоефективності: $m = f(p)$, де m – вхідний розмір задачі, p – число процесорів використаної паралельної системи. Для цього необхідно, спочатку отримати аналітичний вираз для загальних накладних витрат на паралелізм:

$$T_0 = p * T_p - T_1,$$

де T_0 – загальні накладні витрати на паралелізм;

T_p – час виконання алгоритму на паралельній архітектурі;

T_1 – час виконання алгоритму на однопроцесорній системі.

Потім записати основне рівняння ізоефективного аналізу:

$$T_1 = K * T_0,$$

де K – є коефіцієнт масштабованості, що обчислюється як $K = E/(1 - E)$.

Звідси, визначається функція ізоефективності $m = f(p)$, тобто залежність, на основі якої можна обчислити, як необхідно збільшувати розмір задачі при збільшенні кількості процесорів для підтримки постійної ефективності: $m = f(p, top, tw, ts, E, r/k, TF)$, де top – флоп, ts – латентність, tw – час на передачу одного слова, r/k – порядок методу/число точок в блоці, TF – часова складність правої частини СЗДР. Для отримання пріоритетних областей застосування двох або більш алгоритмів для розв'язання однієї задачі треба попарно порівняти їх накладні витрати на паралелізм. Новизна запропонованої методики оцінки якості і масштабованості паралельних чисельних методів полягає в підвищенні ефективності вирішення задачі Коші за рахунок підбору оптимальної комбінації алгоритм/архітектура. Застосування математичного апарату ізоефективного аналізу дозволило перевести дослідження масштабованості паралельних алгоритмів при варіюванні чисельних динамічних параметрів завдання, методу і паралельної системи з області багаторазового експерименту для великих розмірностей в аналітичну область. Практична значимість оцінки масштабованості полягає в побудові оптимальних паралельних алгоритмів для заданих умов реалізації, а також у формуванні областей пріоритетного використання різних алгоритмів для вирішення однієї і тієї ж задачі. Розробка паралельних алгоритмів здійснювалася за ієрархічною декомпозиційною методикою із застосуванням графових моделей. Визначення теоретичних динамічних характеристик паралелізму виконувалося за допомогою пакета Mathematica (Wolfram Research Inc.), програмна реалізація здійснювалася при використанні стандарту передачі повідомлень MPI на мові C++. Тестування алгоритмів проводилося на базі MPICH-2.2, однією з найбільш відомих реалізацій стандарту MPI для OS Windows. Всі тимчасові характеристики алгоритмів визначалися з використанням функцій бар'єрної синхронізації. Для скорочення часу виконання колективних операцій обміну застосовувалися алгоритми оптимальної покоординатної маршрутизації.

Література. 1. Gupta A. Analyzing scalability of parallel algorithms and architectures / A. Gupta, V. Kumar // Technical Report TR 91-18. Computer science Department, University of Minnesota, 1993. To Appear in Journal of Parallel and Distributed Computing. 2. Gupta A., Kumar V. Scalability of parallel algorithm for matrix multiplication // Technical report TR-91-54, Department of CSU of Minneapolis, 2001. – P. 1–211. 3. Назарова И.А. Масштабируемость параллельных алгоритмов и архитектур при численном решении задачи Коши на основе явных разностных схем / И.А. Назарова, Л.П. Фельдман // Наукові праці ДонНТУ. Серія «Інформатика, кібернетика та обчислювальна техніка». – 2017. – № 1(24). – С. 100–107. 4. Назарова І.А. Дослідження масштабованості паралельних алгоритмів на основі ізоефективного аналізу / І.А. Назарова, А.О. Чорна // Наукові праці ДонНТУ. Серія «Інформатика, кібернетика та обчислювальна техніка». – 2015. – №2(21). – С. 56–62.

Орехов О.А.¹, Орехова Н.А.²

¹КПІ ім. Ігоря Сікорського, ФІОТ, Київ, Україна; ²Інститут кібернетики ім. В.М.Глушкова НАН України, Київ, Україна

Порівняльний аналіз Blockchain-архітектури на основі дерев Меркле та протоколу PHANTOM на основі архітектури BlockDAG

Досліджуються два підходи до побудови розподілених смарт-контрактів та розглядаються аспекти виконання CAP-теореми для blockchain-протоколу та його BlockDAG-узагальнення PHANTOM.

Bitcoin-протокол та CAP-теорема. Блокчейн або послідовність блоків транзакцій [1] – відповідним чином організована розподілена постійно зростаюча база даних. Основним засобом утримання її розміру від швидкого розростання є технологія дерев Меркле, або “Мерклінг” [2]. При цьому (внаслідок децентралізованої природи цієї технології) питання узгодженості (C – consistency) даних, розташованих у різних вузлах (що позначається як P – Partition tolerance), згідно з відомою CAP-теоремою або теоремою Брюера [3], може бути вирішено тільки за рахунок збільшення проміжку часу, після якого вони стають доступними (A – availability). Як зауважено в [4], можливим є втілення обох сценаріїв – як PA, так і PC, але тільки за умови реалізації у вигляді відповідної архітектури розподіленої системи. Оскільки блокчейн-системи будуються з метою підтвердження певних фактів в умовах забезпечення цілісності даних і недопущення їхньої компрометації, більш нагальною є побудова рішення PC, приклад якого наведено в [5], із “сильною ймовірнісною узгодженістю” на базі “консенсусу Накамото”.

PHANTOM-протокол. На відміну від Bitcoin-протоколу, за яким будується дерево із посиланням на найдовший ланцюг серед існуючих, запропонований у [6] PHANTOM-протокол виконує структурування блоків у вигляді орієнтованого ациклічного графу (blockDAG – block Directed Acyclic Graph). Для вирішення проблеми розв’язку конфлікту транзакцій автори пропонують покласти відповідальність на взаємопов’язані “чесні” вузли, що співпрацюють в щільній кооперації. Такий підхід, на їхню думку, надасть змогу виявити “одинаків” і відповідним чином попередити атаку на цілісність з боку останніх.

Особливості реалізації PHANTOM-протоколу. Перехід від дерева до орієнтованого ациклічного графу визначає необхідність обробки можливих конфліктів між блоками, що не знаходяться на спільних ланцюгах і знаходяться у невизначеному стосовно один одного підпорядкуванні, для чого в роботі [6] вводиться поняття “антиконусу” (‘anticone’), а також робляться оцінки щодо теоретичної складності задач, що виникають. Незважаючи на те, що в загальному випадку задача пошуку максимального розміру кластеру DAG-графу є NP-складною, автори стверджують, що лінійно-складні “жадібні” алгоритми демонструють здатність знаходити асимптотично прийнятне рішення.

Чисельний експеримент. Для перевірки моделі було використано програму, описану в [7], в яку були внесені відповідні зміни. Метою експерименту було визначення принципової можливості використання запропонованого підходу. Було перевірено найбільш очевидні реалізації “жадібних” алгоритмів. Подальші дослідження планується спрямувати на вивчення поведінки PHANTOM-алгоритму у залежності від співвідношення між “чесними” акторами, що діють узгоджено, і “одинаками”. Робиться припущення, що ефективність алгоритму при збільшенні кількості блоків буде суттєво залежати від описаних в роботі [8] стратегій, які оберуть “одинаки”, та питомої ваги “одинаків” у blockDAG-графі.

Література. 1. <https://bitcoin.org/bitcoin.pdf>. 2. <https://blog.ethereum.org/2015/11/15/merkle-in-ethereum/>. 3. <http://www.julianbrowne.com/article/brewers-cap-theorem>. 4. https://www.goland.org/blockchain_and_cap/. 5. https://www.goland.org/why_block_chains_are_strongly_consistent/. 6. <https://eprint.iacr.org/2018/104.pdf>. 7. <https://hackernoon.com/learn-blockchains-by-building-one-117428612f46>. 8. Lefebvre, V.A. (2001). Algebra of Conscience. Dordrecht, London: Kluwer Academic Publisher. ISBN: 978-90-481-5751-8.

Останчук Я.М.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Стратегія розвитку CRM систем як сервісу у хмарі Microsoft Azure

Вступ. Досить багато компаній має відносини з їхніми клієнтами за допомогою введення таблиць Excel або використовуючи десктопні CRM (Customer Relationship Management) системи. Проте з того часу, як в 1993 році Том Сібел запустив додаток для автоматизації продажів і в подальшому став використовувати його для отримання інформації про клієнтів, яка стала першою CRM, багато чого змінилося [1–3]. На сьогодні, враховуючи нові тенденції, такі як технології блокчейн, використання штучного інтелекту у автоматизації процесів, інтеграція IoT, використання e-commerce та багато іншого вимагає розміщення CRM системи у хмарі [4]. Тому виходячи з нових тенденцій та зі сторони зручності, давайте розберемо стратегію побудови CRM системи у хмарі Microsoft Azure.

Вибір платформи. Для побудови CRM системи треба обрати досить правильну стратегію, де дані будуть захищені та будуть використані нові технології. Тому будемо використовувати Microsoft Azure PaaS (Platform as a Service), так як надається багато готових послуг, таких як балансування навантаження, backup, azure storage, аналітика, адміністрування та реплікація даних, а також має багато сервісів які дозволяють реалізувати нові тренди CRM систем.

Архітектура. Для початку треба розглянути типову архітектуру CRM у хмарі (рис. 1), де видно, як трафік проходить спочатку через traffic manager, який дозволяє нам боротися з Ddos атаками, а також змінювати трафік у центри обробки, які знаходяться ближче або через CDN (Content Delivery Network), який дозволяє не навантажувати головний сервер, а дістати наші дані напряму. Далі трафік проходить через load balancer, який розділяє трафік між декількома серверами чи службами і потім до нашого веб інтерфейсу, де клієнт може виконувати свою роботу.

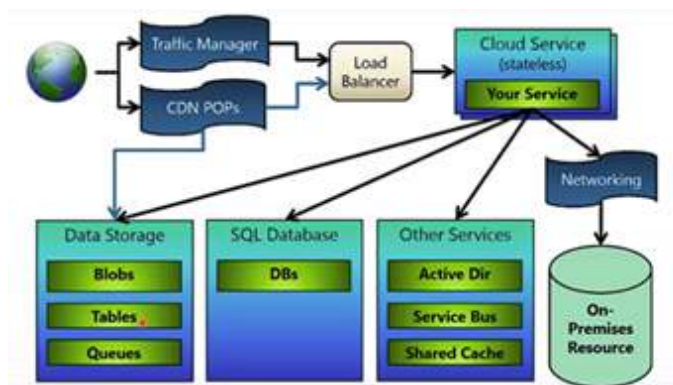


Рис. 1. Приклад Cloud Service у хмарі

Виконання коду. Хмарна CRM у хмарі Microsoft azure складається з таких частин, як веб інтерфейс з яким користувач взаємодіє (Web-role), веб-сервіси, які можуть представляти собою APi, а також веб-сервіси, які постійно працюють у хмарі (Worker-role) та обробляють запити, виконуючи постійну роботу, наприклад обробляючи повідомлення з черги або перевіряючи якісь дані, такі як курс валют чи акцій.

Збереження даних. Для збереження даних можна використовувати як реляційну базу даних таку як MySQL або SQL Azure, яка по суті являється Microsoft SQL Server у хмарі, так і за допомогою Azure Storage, яке включає у себе: Blobs, Tables, Queues. Ми використовуємо Blobs для збереження великих даних, таких як фото або відео файли, а також документи, які в свою чергу можна розмістити у BlockBlobs або PageBlobs, де BlockBlobs може містити інформацію до 200 Гб, а PageBlobs до 1 Тб. Для збереження даних, які треба масштабувати або мати до

них швидкий доступ, ми використовуємо Tables, яка являє собою NoSQL, де доступ можна отримати через пару ключ-значення. Для забезпечення передачі повідомлень між сервісами ми використовуємо Queues, які зберігаються до 7 днів. Також Azure має базу даних Azure Cosmos DB, яка являє собою багатомодельну глобальну розподілену службу баз даних [5].

Аналітика. Для аналітика Azure надає такі сервіси, як Stream Analytics, які дозволяють виконувати аналітику для Інтернет речей (IoT) у хмарі, а також приймати тисячу потоків аналітичних даних (рис. 2) [6]. Ще використовується Data Lake Analytics, яка використовується для запуску завдань аналізу великих даних, що обробляють великі масиви даних. Управління та іншої документації, в якій показано, як створювати й адмініструвати завдання пакетної, оперативної та інтерактивної аналітики, а також створювати запити на мові U-SQL. Використання машинного навчання (Azure Machine Learning), що використовується для аналізу даних.

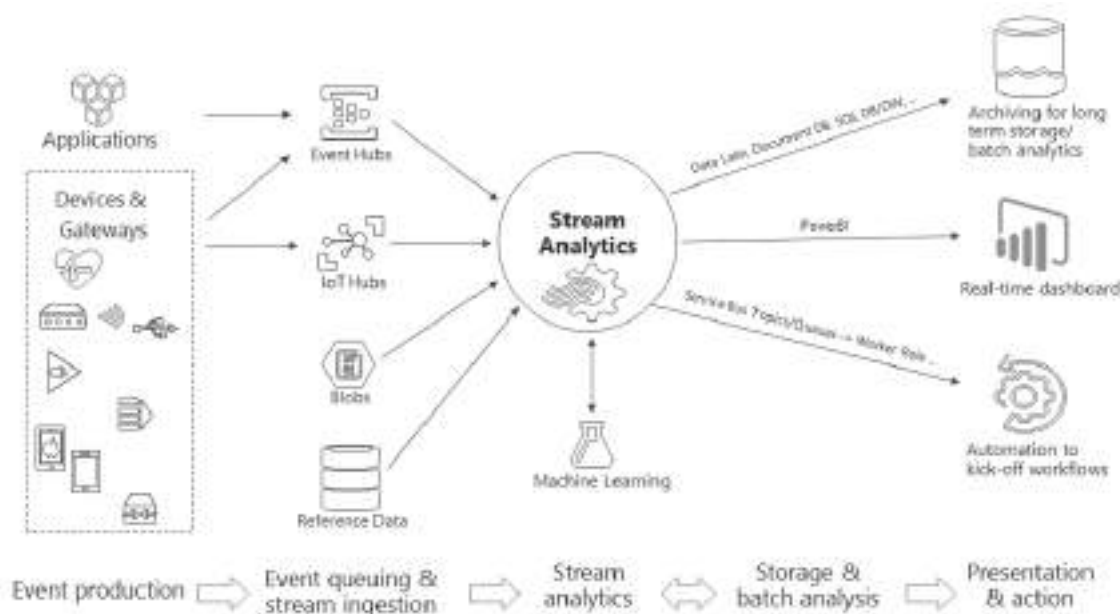


Рис. 2. Інтегрування Stream analytics

Висновок. Таким чином ми задаємося питанням, які переваги нам надає PaaS Microsoft Azure порівняно з простою IaaS? Головною перевагою є заощадження часу, а також коштів, через те що нам не потрібні витрати на системних адміністраторів, яких треба мінімум чотири, а також гарантує надійність збереження даних і швидкість доступу до них, завдяки механізму зберігання даних Microsoft Azure (Blobs, Tables, Queues), можливості яких були описані вище. Треба зазначити, що також на платформі ми маємо багато готових сервісів, які полегшують нашу роботу, що було зазначено у розділі аналітика. Проте, треба зазначити, що світ не стоїть на місці та кожні півроку ми бачимо нові тренди у сфері інноваційних технологій, які за допомогою хмарної платформи стає втілити набагато легше, а найголовніше швидше.

Література. 1. Handbook of Service Science / Springer, 2010 — 720 p. 2. Maglio Paul, Kieliszewski Cheryl, Spohrer James. Handbook of Service Science. — Pittsburg: Springer US, 2010. — 720 p. 3. Succeeding through Service Innovation. <http://www.ifm.eng.cam.ac.uk/ssme/>. 4. Article By Manohar Chapalamadugu CRM trends <http://www.destinationcrm.com/Articles/Web-Exclusives/Viewpoints/6-CRM-Trends-to-Watch-in-2018-122717.aspx>. 5. Збереження даних в Microsoft azure <https://docs.microsoft.com/ru-ru/azure/#pivot=products&panel=storage>. 6. Аналітика в Microsoft azure <https://docs.microsoft.com/ru-ru/azure/stream-analytics/stream-analytics-introduction>.

Сапсай Т.Г., Довганюк А.О.

КПИ им. Игоря Сикорского, ФПМ, Киев, Украина

О последовательном диагностировании многопроцессорных систем

При разработке сложных систем управления на базе отказоустойчивых многопроцессорных систем (ОМС) организация самодиагностирования является одним из важнейших условий, что позволяет системе обнаруживать неисправные модули без дополнительного технического обслуживания, а также перераспределять нагрузку неисправных модулей на работоспособные. У ОМС есть ограничения на количество неисправных процессоров, которые сможет обнаружить система и при этом остаться работоспособной. В докладе рассматриваются случаи $T \leq 4$.

В качестве модели неисправностей выбрана модель Препарата-Метца-Чена, в которой процессоры обозначаются вершинами графа, а диагностические связи – дугами. Так же следует отметить результаты тестирования u_{ij} i -м процессором j -го: если i, j – исправны, $v_{ij}=0$; если i – исправен, j – неисправен, $v_{ij}=1$; если i – неисправен, $v_{ij} = x$, где x принимает значение 0 или 1 независимо от состояния тестируемого процессора [1].

Метод предполагает использование диагностического графа, который назван конструктивно регулярным и позволяет определить состояние всех модулей за исключением одного или двух процессоров в некоторых случаях [2]. Связи графа – однотипны, вершины имеют 2 входящие и 2 исходящие дуги. На количество процессоров n , скачки i и j накладываются некоторые ограничения: $n \neq ki$, $n \neq kj$, $i \neq kj$, $j \neq ki$, где $k > 1$ – целое.

Используются следующие понятия:

- 0-цепочка – цепочка вершин, которые связаны 0-дугами;
- вес модуля – количество модулей вместе с данным, которые подходят к нему 0-дугами.

В отличие от описанного в [2] метода, в котором предлагается параллельное и независимое тестирование, ниже описан способ последовательного взаимотестирования процессоров, т.е. в один момент времени может тестироваться только одна пара процессоров. После получения результата тестирования, проводится анализ, на основе которого либо определяется состояние модуля/модулей, либо выбирается следующая пара процессоров. Процесс диагностирования заканчивается определением состояния системы за исключением случаев возникновения одной или двух неопределенностей.

В докладе определяется нижняя граница количества тестовых проверок, при котором можно определить состояние системы – $T + 1$, где T – наибольшее количество неисправных модулей, которое допускается в системе.

Необходимые условия и ограничения:

- исправным является $(T + 1)$ -й процессор в 0-цепочке, неисправный определяется по результатам тестирования исправным;
- каждая установленная неисправность уменьшает величину T на 1.

Литература. 1. Preparata F.P., Metze G., Chien R.T. On the connection assignment problem of diagnosable systems // IEEE Trans. Comput. 1967. V. C-16. P. 848–854. 2. Romankevich V.A. Self-testing of multiprocessor systems with regular diagnostic connections / V.A. Romankevich // Automation and Remote Control. – 2017. – Vol. 78, Issue 2. – P. 289–299.

Слинько М.С.

Київський національний університет ім. Тараса Шевченка, факультет радіофізики, електроніки та комп'ютерних систем, Київ, Україна

Формалізоване проектування застосувань для графічних відеоадаптерів

Майбутнє високопродуктивних обчислень лежить в галузі систем з масовим паралелізмом. Прикладом таких систем є відеоадаптери, що використовуються як прискорювачі для вирішення неграфічних задач. Втім, використання цих модулів унеможливує використання алгоритмічного етапу проектування. У роботі розглянуто метод дослідження високопродуктивних систем, який ґрунтується на поєднанні апаратів транзиційних систем та мереж Петрі.

Відеоадаптери компанії NVidia можуть виділяти до 655360 потоків на вирішення задачі (в архітектурі Pascal). Дослідження і виправлення помилок проектування в такому контексті є неефективним у ручному режимі. Ефективним методом алгоритмічного проектування і дослідження паралельних алгоритмів є використання алгебр алгоритмів, які дозволяють сформулювати формули (схеми) алгоритмів, що залежать від різних параметрів, якими можуть виступати програмно-апаратні платформи, парадигми паралельного програмування тощо. В роботі пропонується використання апарату транзиційних систем, що дозволяють створювати різнорівневі абстракції, та можуть бути зведені до мереж Петрі (МП), котрі підтримують темпоральне моделювання.

В роботі проведено декомпозицію архітектури обчислень GPGPU на прикладі NVidia CUDA та досліджується планування виконання інструкцій. На апаратному рівні GPU від NVidia складаються з набору поточкових мультипроцесорів. Мультипроцесор виділяє групи по 32 потоки, які називаються варпами. Далі кожен мультипроцесор працює за принципом SIMT, тобто всі потоки в межах варпу в один момент часу виконують одну і ту саму інструкцію. В результаті теоретичних досліджень процесів виконання GPGPU-застосувань виділено наступні підсистеми:

1. Підсистема, що моделює абстрактну інструкцію. В роботі враховуються 2 типи інструкцій – арифметична та інструкція доступу до пам'яті. Втім, час доступу до різних видів пам'яті GPU відрізняється на порядки, тому в реальній моделі може бути доцільним виділення додаткових типів інструкцій, та, відповідно, додаткових місць та переходів.
2. Підсистема, що моделює варп, якому виділено набір інструкцій, та який послідовно їх виконує. Кожен варп містить набір інструкцій, що мають бути виконаними всіма його потоками. За моделлю SIMT інструкції надходять послідовно, а кожна з них виконується одночасно всіма потоками варпу.
3. Підсистема, що моделює виконання блоку потоків на поточковому мультипроцесорі. Основною структурною одиницею в даному процесі за означенням є планувальник варпу. Основні дії на цьому етапі включають вибір варпу та інструкції і забезпечення ексклюзивного доступу до обчислювальних ресурсів на час виконання кожної пари варп-інструкція. Ця пара дій має ітеративно виконуватися, поки всі інструкції не буде виконано на всіх варпах блоку.

Кожну підсистему представлено у вигляді окремої транзиційної системи (ТС), а процес виконання GPGPU-застосування – у вигляді синхронного добутку трьох ТС [1]. Апарат транзиційних систем було обрано як дискретну модель досить загального типу. За множиною обмежень синхронного добутку побудовано мережу Петрі (МП), що моделює сумісну роботу всіх підсистем (рис. 1). Зображення добутку ТС у вигляді мережі Петрі дало можливість застосовувати методи аналізу властивостей МП для аналізу властивостей добутку ТС. Семантика добутку ТС і семантика МП, яка його зображує, узгоджуються в тому сенсі, що послідовність глобальних переходів являє собою глобальну історію добутку ТС тоді і тільки тоді, коли вона є допустимою послідовністю спрацювань переходів в МП.

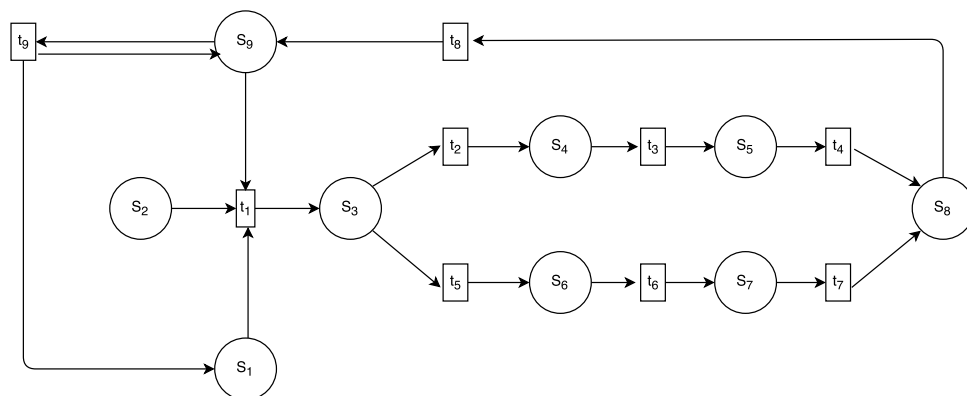


Рис. 1. Мережа Петрі, що представляє узагальнену модель виконання GPGPU-застосування

В даній мережі Петрі варто окремо описати наступні місця та переходи:

1. Місце S_1 , що моделює планувальник варпу, який обирає інструкцію для виконання потоками варпу.
2. Місце S_2 , що містить токени-інструкції.
3. Місце S_9 , що слугує маркером закінчення виконання поточної інструкції.
4. Перехід t_1 , що ініціює початок виконання інструкції.
5. Перехід t_9 , що ініціює наступну ітерацію роботи планувальника варпу після закінчення обробки поточної інструкції.

Інші елементи МП на рис. 1 моделюють виконання інструкції та інтеграцію з поточним варпом та планувальником.

Концептуальні помилки бажано знаходити ще на етапі проектування. В роботі досліджується структурна обмеженість та живучість моделі. В роботі розглянуто живучість системи при початковій розмітці $M_0 = (1, 1, 0, 0, 0, 0, 0, 0, 1)$, та кінцевій розмітці $M_k = (1, 0, 0, 0, 0, 0, 0, 0, 1)$. За допомогою алгоритму TSS [2] знайдено розв'язки рівняння стану МП і показано, що всі переходи покриваються позитивними розв'язками рівняння стану, тобто система є живою.

Табл. 1. Розв'язки рівняння стану МП для заданих початкової та кінцевої розміток

t_1	t_2	t_3	t_4	t_5	t_6	t_7	t_8	t_9
1	1	1	1	0	0	0	1	1
1	0	0	0	1	1	1	1	1

Для дослідження структурної обмеженості розглянуто розв'язки рівняння стану МП з транспонованою матрицею інцидентності, і доведено структурну обмеженість моделі.

Висновки. Показано, що проектування застосувань в гетерогенних середовищах виконання, зокрема з використанням GPU-модулів, потребує формалізованого математичного апарату та методів представлення алгоритму задля їх верифікації. Проведено декомпозицію процесу виконання застосування в архітектурі NVidia CUDA. Як приклад використано узагальнену модель цього процесу, створену в апараті транзиційних систем. Представлено можливості дослідження характеристик моделі, утвореної поєднанням апаратів транзиційних систем та мереж Петрі. Як приклади, розглянуто структурну обмеженість та живучість моделі.

Література. 1. Бойко Ю.В., Кривий С.Л., Погорілий С.Д. та ін. Методи та новітні підходи до проектування, управління і застосування високопродуктивних ІТ-інфраструктур. – К.: ВПЦ «Київський університет». – 2016. – 447 с. 2. Кривий С.Л. Лінійні діофантові обмеження та їх застосування. – Букрек: Чернівці. – 2015. – 224 с.



This page was left blank intentionally

Plenary talks

System analysis of
complex systems of
various nature

Intelligent systems for
decision-making

High Performance
Computing and
Microsystems Engineering

Progressive information
technologies

4

Progressive information technologies



Section 4

Progressive information technologies

1. Support (mathematical, algorithmic, linguistic, informational-organizational, technical, software) of the control systems, information processing, and development technologies.
2. E-commerce.
3. Information security and guarding.
4. Data and knowledge bases as the environment of information support for control and design.

Секция 4

Прогрессивные информационные технологии

1. Обеспечение систем управления (математическое, алгоритмическое, лингвистическое, информационно-организационное, техническое, программное), обработка информации и технологии их создания.
2. Электронная коммерция.
3. Информационная безопасность и защита информации.
4. Базы данных и знаний как среда информационной поддержки управления и проектирования.

Секція 4

Прогресивні інформаційні технології

1. Забезпечення систем управління (математичне, алгоритмічне, лінгвістичне, інформаційно-організаційне, технічне, програмне), обробка інформації та технології їх створення.
2. Електронна комерція.
3. Інформаційна безпека та захист інформації.
4. Бази даних і знань як середовище інформаційної підтримки управління та проектування.

Basyuk T.M.

National University “Lviv polytechnic”, Lviv, Ukraine

Popularization of Internet resources by using “featured snippets”

The task of popularizing any online resource is reflected in the TOP of search results. According to the latest data, it is considered that the location on the first page of the publication is insufficient, as about 30% of transitions falls to the first position, and 50% – to the first three [1]. In other words, there is no significant difference whether the resource is located on the fifteenth position of search results or on the fiftieth, as the probability of its viewing will be extremely bare. Research shows that this situation has arisen for several reasons: the use of mobile devices (users simply “do not see” all results due to restrictions, which are determined by their size); search psychology (if the user repeatedly finds the necessary information in the first three positions of search results, then he/she creates a subconscious relationship that indicates that the next time the best results will be on them) [2]. The listed above reasons led to extremely high demands on the quality of work of both seo-optimizers and copywriters, requiring them to use a variety of mechanisms to increase the visibility of the resource for targeted queries.

The described problem is complicated by the situation of implementation of Google’s search technology of voice control, which requires finding the best answer to the user’s questions. In order to solve this problem, technologies of pertinent relevance have begun to be applied. The feature of which is not only the search for relevant information, but also the satisfaction measure with it. That is, to what extent the latter corresponds to the expectations of the user. As a specific solution to this problem, “featured snippet” began appearing in search results, entering into which guarantees a quantitative jump in search traffic. The answer block is a text that is automatically displayed from an internet page with the most relevant information, according to the search engine. The distinctive feature of the latter is the design (radically different from the rest of the search results) and the location (displayed above the search results in the so-called zero position). Advanced “featured snippets” greatly increase the trust of the company, do not require additional payment, and facilitate the attraction of search traffic. In this case, the content of the resource in the snippet is much more effective in the process of attracting visitors than displaying the first position of search results [1].

The analysis showed that for today there is no single technology of entering into the extended featured snippet. In particular, having entered into the featured snippet in one of the search results, at the next result, the search engine may decide that another resource better matches the display parameters. Given that the actual task is to analyze known types of featured snippets and formulate recommendations for popularizing resources with their use. There are the following types of features snippets: Paragraph snippets – includes a description as a paragraph of text and may contain an image; List snippets – includes elements with step-by-step instructions; Table snippets – Used to structure data.

The study conducted by using the Getstat statistics service showed that the most popular type of featured snippet is the display in the form of paragraph snippet (about 82%). In view of this, recommendations were made to popularize the resource by entering featured snippets in the Google search engine:

1) *The use of question-answer strategy.* An experimental study showed that most queries entering into the feature snippet are answers to user questions. That is, the search engine first of all searches for the answer to the questions on the content of the resource that is included in the top search engine. Given that one of the conditions for the promotion of the reference resource is the clear statement of questions in the relevant subject area and full answers to them. In this case, the content should be relevant and comply with the requirements of registration.

2) *Analysis of Competitive Resources.* This recommendation is extremely important in the case of an expanded competitive environment and is realized by a plurality of resources. Among the

common ones is the statistical service Semrush, which not only searches for advanced featured snippets, but also allows analyzing snippets that entered into the explored Internet resource. The analysis of the search queries of competitors, after which the latter receive advanced featured snippets, will allow both carrying out operations for the optimization of content of the Internet resource and the selection of keywords.

3) *Answers to related questions within one article.* According to the research carried out by using the Ahrefs Seo-analysis tool, it turned out that as soon as the featured snippet refers to the defined Internet resource, the content of this page is likely to enter into a specific snippet by related queries. In view of this, the structuring of the content in such a way as to answer as closely as possible the user's questions within the same page is the author's essential condition.

4) *The optimal number of words.* The analysis conducted by means of Semrush service showed that the content length in featured snippets varies from forty to fifty words. This means that, ideally, every paragraph of the text of an article should consist of the corresponding number of words.

Taking into account the given recommendations, the general algorithm of the access of the Internet resource to the advanced featured snippet was formed.

The main stages of the algorithm.

1. Finding keywords related to the resource. It is carried out at the stage of constructing a semantic core and is to search for information queries related to the resource being studied.
2. Study of related queries. It is implemented by an increase in the information value of the page by applying the described recommendations for the formation of content.
3. Add links to other resource pages. This action can increase the weight of the page and its visibility by the search engine [4].
4. Selection of the format of advanced featured snippet. Although the study showed that the most popular paragraph snippet, however, the specificity of the subject area can make significant adjustments both in favor of table representation and list.
5. Content publishing and testing. The final stage is to monitor the changes introduced. If the created pages do not enter into the advanced featured snippets, it is necessary to re-analyze the actions of competitors (recommendation 1) and improve the content of the Internet resource.

Conclusion. The study conducted showed that the advanced featured snippet is the highest ranking that, on the one hand, promotes both the increase of site conversion and increase in the number of visitors, while on the other hand – it enables bypassing competitors without financial investments in advertising. In view of this, the provided recommendations and formed stages create a methodological basis for the popularization of the Internet resource with the display of content in the advanced featured snippet.

Further research will be aimed at creating mathematical and algorithmic support for the popularization of Internet resources by using “featured snippet”.

References. 1. Petrescu P. Google Organic Click-Through Rates in 2014. Available online at: <https://moz.com/blog/google-organic-click-through-rates-in-2014>. 2. Basyuk T. The Popularization Problem of Websites and Analysis of Competitors. In: Shakhovska N., Stepashko V. (eds), *Advances in Intelligent Systems and Computing II. CSIT 2017. Advances in Intelligent Systems and Computing*, vol 689. Springer, Cham pp. 54–65. IEEE (DOI:10.1007/978-3-319-70581-1-4). 3. Schwartz B. New study on featured snippets within Google reveals growth optimization techniques. Available online at <https://searchengineland.com/new-study-featured-snippets-within-google-reveal-growth-optimization-techniques-275714>. 4. Basyuk T. Innerlinking website pages and weight of links. *Proceedings of the XII International Scientific and Technical Conference “Computer science and information technologies CSIT-2017”, 2017*, pp.12–15, IEEE (DOI: 10.1109/STC-CSIT.2017.8098725).

Kalinovsky Ya.A.¹, Boyarinova Yu.E.^{1,2}, Khitsko Ya.V.², Sukalo A.S.³

¹*Institute for Information Recording of the NAS of Ukraine, Kyiv, Ukraine;* ²*Igor Sikorsky Kyiv Polytechnic Institute, Kyiv, Ukraine;* ³*National University of Water and Environmental Engineering, Rivno, Ukraine*

Principles of constructing algorithms for processing digital signals using hypercomplex number systems

One of the most common tasks when processing digital signals is the convolution of numerical sequences. Since the complexity of convolution calculation rapidly increases with the length of collapsed arrays and their dimension, it is very important to develop such convolution algorithms that require as fewer real operations as possible. There are a large number of methods for the rapid calculation of convolution [1]. In this paper, we consider algorithms for performing convolution based on the transition to hypercomplex spaces. The foundations of this approach were developed by the authors [2,3].

Let it be necessary to perform one-dimensional linear convolution of two numerical sequences of dimension 2^n . We will consider these numerical sequences as components of hypercomplex numbers belonging to some hypercomplex number systems (HNS) Γ_1 dimension 2^n . The product of these numbers will contain the pair products of the components of the convolution sequences. However, they will be combined in amounts not in the same composition as necessary for organizing convolution components. In addition, the number of real multiplications with multiplication of hypercomplex numbers in the general case is equal to 2^{2n} , this is the same as in the direct calculation of convolution.

There are two problems: the first is the reduction in the number of real operations when multiplying the hypercomplex numbers; The second is the organization of the choice of paired products of convolution components. The solution of these two problems allows one to synthesize such convolution algorithms, which by the number of operations are more efficient than direct calculation algorithms for convolution.

To solve the first problem, we can go over to such a HNS isomorphic to the original one, the multiplication table of which is weakly filled. Such pairs of HNS exist [4,5]. The transition between them requires only the operations of adding real numbers. The solution to the second problem depends on the particular type of used HNS.

To convince the principle of reducing the number of multiplications, the convolution of numerical sequences – $\{x_0, x_1\}$ and $\{y_0, y_1\}$. Convolution will have three counts:

$$z_{-1} = x_0 y_0; z_0 = x_0 y_1 + x_1 y_0; z_1 = x_1 y_1. \quad (1)$$

To perform (1), you need to perform 4 real multiplications. We will consider the elements of sequences $\{x_0, x_1\}$ and $\{y_0, y_1\}$ as components of hypercomplex numbers belonging to a 2-dimensional HNS W with a multiplication table $e_1 * e_1 = e_2 * e_2 = e_1; e_1 * e_2 = e_2 * e_1 = e_2, \dots$. Let $X = x_0 e_0 + x_1 e_1; Y = y_0 e_0 + y_1 e_1$. Then

$$XY = (x_0 y_1 + x_1 y_0) e_0 + (x_0 y_0 + x_1 y_1) e_1. \quad (2)$$

The right-hand side of (2) contains all the pair products necessary to form the convolution (1). Moreover, the component at e_0 completely coincides with the count z_0 , and the component e_1 is equal to the sum of the remaining samples. But the number of multiplications in (2) is also 4, as in (1). Thus, it is necessary to reduce the number of multiplications and separate the samples x_{-1} and x_1 . To reduce the number of multiplications, it is advisable to go from W system to W_1 system with multiplication table $f_1 * f_1 = f_1; f_2 * f_2 = f_2; f_1 * f_2 = f_2 * f_1 = 0$. To translate numbers from W to W_1 , we use the isomorphism operator of these systems [2,4,5]. Then $X_1 = (x_0 + x_1) f_0 + (x_0 - x_1) f_1; Y_1 = (y_0 + y_1) f_0 + (y_0 - y_1) f_1$. Such a transition requires

an additional 4 addition. However, considering that the core of the convolution is constant, the transition of the elements W_1 of which is done in advance, we can assume that the transition from the W system to the W_1 system requires only 2 additions.

The multiplication of numbers in accordance with the multiplication table W_1 :

$$X_1Y_1 = \alpha_0f_0 + \alpha_1f_1 = (x_0 + x_1)(y_0 + y_1)f_0 + (x_0 - x_1)(y_0 - y_1)f_1, \quad (3)$$

and requires 2 multiplications.

The reverse transition from W_1 system to system in accordance with [3–5] will have the form:

$$XY = e_0(\alpha_0 + \alpha_1)/2 + e_0(\alpha_0 - \alpha_1)/2, \quad (4)$$

which also requires 2 addition. Division by 2 is a short operation, requiring only shift of the register. It can be performed in advance of the convolution kernel. Therefore, it can be ignored. Equating the right-hand sides of (3) and (4), we obtain:

$$x_0y_1 + x_1y_0 = (\alpha_0 + \alpha_1)/2; x_0y_0 + x_1y_1 = (\alpha_0 - \alpha_1)/2. \quad (5)$$

If one of the terms on the left-hand side of (5) is calculated in advance, for example, x_0y_0 then the convolution components (2) will be calculated by the following formulas:

$$z_{-1} = x_0y_0; z_0 = (\alpha_0 + \alpha_1)/2; z_1 = (\alpha_0 - \alpha_1)/2 - x_0y_0, \quad (6)$$

which requires 3 multiplications and 5 additions. In this case, it is not large, but as the length of the convolution sequences increases, as shown by the authors of [4], it increases significantly, as shown in the table.

Table 1. Quantity of multiplication

№	n	Quantity of multiplication	Decrease the quantity of multiplications	In percentages
1	2	4	1	25
2	4	16	5	31.3
3	8	64	19	30

To synthesize algorithms for large n , pairs of isomorphic large-size HNS are required, respectively. The methods for creating such HNS were considered in the authors' works [4]. They are based on Theorems 1 and 2 of this paper on the isomorphism of HNS with doubling Grassmann-Clifford. The conducted researches have shown that the use of methods of hypercomplex numerical systems allows to significantly reduce the number of material operations in the algorithms of performing linear convolution. In this case, the choice of the rational structure of the algorithm depends on the dimension of the collapsible arrays.

References. 1. Bleihut R. Fast algorithms for digital signal processing / R.Bleihut – M.: Mir, 1989. – 449 p. 2. Kalinovsky Ya.A. Hypercomplex number systems and fast algorithms for digital information processing / Ya.A.Kalinovsky, D.V.Lande, Yu.E. Boyarinova, Ya.V. Khitsko – K.: IPRI NASof Ukraine, 2014. – 130p. 3. Kalinovsky Ya.A. Structure of the hypercomplex method for fast calculation of linear convolution of discrete signals/Kalinovsky Ya.A. // Реєстрація, зберігання і обробка даних. – 2013. – Т. 15, № 1. – P. 31–44. 4. Kalinovsky Ya.A. High-dimensional isomorphic hypercomplex number systems and their use to improve computational efficiency / Ya.A. Kalinovsky, Yu.E. Boyarinova. – K.: Infodruk, 2012. – 183p. 5. Kalinovsky Y.A. The basic principles and the structure and algorithmically software of computing by hypercomplex number/ Y.A. Kalinovsky, Y.E. Boyarinova, A.S. Sukalo, Y.V. Khitsko – arXiv preprint arXiv:1708.04021, 2017.

Kasyanchuk I. V.

Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

Education System to Improve Remembering of Information

Review of learning management systems and searching the way how to improve it referring to brain training.

Abstract. Learning Management System (LMS) is a software for delivery educational material for students, tracking their progress, reporting, contacting them etc. Many kinds of LMS take place: webbed and installed, free and commercial, integrated and not [1]. All of them consist of several phases to learn the material, such as reading text (or watching video), checking learning progress with system using tests and student gets his appraisal. After that learner need to make decision by himself what to do with his result – continue learning another material or repeat it again.

The disadvantage of this approach is that user have to learn material once and nothing motivates to repeat it once more. Consequently, user usually forgets learned material without regular training. So, there is need to implement learning system which required user to repeat learning material systematically.

Learning Concept. Remembering information in our brain consists of several phases. After material familiarization, it appears in short-term memory. But without systematically repeating this material can be forgotten. To exclude it, we need to store learned information in long-term memory, occasionally renew material. Afterwards, it will be not difficult to recollect it [2].

The main advantage of proposed LMS is that user's data need to be repeated in determined time intervals. To achieve that, system sends push-notifications every time, when interval expired based on **spaced repetition** learning technique [3]. According to this technique, repeating stage comes after some period of time, and this period will be increased in the next time [4]. The system automatically scales this period using the previous tests results and machine learning algorithms.

There are several types of information that LMS can be worked with: *text* and *terms*. *Text* is logical sequence of sentences, with sense and structure. *Terms* are pairs of word with their definition. In this work we consider the LMS with terms data type only.

The learning process configures to user. In the *terms* data can be scaled with such of parameters: *data size* and *data order*.

One batch of data that user will be learned in one time is called *iteration*. Iterations come when some period of time is passed and user needs to repeat words.

Data size is a count of terms that user can learn or repeat for one iteration. When system determines if user have successfully learned N words, it tries to increase count to $N + 1$ words. Otherwise, it decreases this count to $N - 1$ words, so size of information goes down.

Another learning parameter is *learning order*. This is an order in which information shows to user. Learning order depends on how much time has passed from last word been repeated and the last result of test.

Mathematical Model. The mathematical model describes algorithm work based on common symbols.

Dictionary of terms looks like as set of pairs “word-translation” (1).

$$W = \{ \text{“word1”} \rightarrow \text{“translation1”}, \text{“word2”} \rightarrow \text{“translation2”} \} \quad (1)$$

The set of words which were learned at iteration i defined as Q_i .

$$Q_i = \{q_1, q_2, \dots, q_{n_i}\} \quad (2)$$

The set of words which were learned at iteration i and user has passed the test (confirmed his knowledge) defined as T_i .

$$T_i = \{t_1, t_2, \dots, t_{n_i}\}, T_i \subset Q_i \quad (3)$$

Other words, that is not learned by user yet, store in set F_i .

$$F_i = Q_i \setminus T_i \quad (4)$$

Also, the scaling parameters defined as n_k – data size (count of words) on k -th iteration, and, consequently, Q_k – the dictionary on k -th iteration.

For dictionary Q_{k+1} it can be used the algorithm below.

1. For each word w_i from dictionary:
 - Find the iteration number l_i which contains the current word from the end. Also, check the test result in this iteration and save to b_i (that should be equals 1 if user learned word or 0 otherwise). If number of iteration is not found (user has not reviewed this word yet), then $l_i = 0, b_i = 0$.
 - Remembering factor $r_i = l_i/n$, where n – the total number of iterations.
 - Add tuple $\{w_i, r_i, b_i\}$ to the result vector U .
2. Sort the vector U by b_i field and r_i field after. It reorders vocabulary data accordingly with how long time was passed after the word have repeated, and words that was not recalled by user will be ordered before another words.
3. The first n_{k+1} words (where n_{k+1} is the optimal count of words, that was calculated by machine learning algorithm) is the words that should be reviewed in $(k + 1)$ -th iteration.

$$Q_{k+1} = \{U_0, U_1, \dots, U_{n_{k+1}}\} \quad (5)$$

This algorithm is the example of system which will be automatically configured to user learning capabilities.

Auto-Configured Learning System. Prediction of mentioned above parameters can be embedded in any existing learning management system. Developing of this system is a part of this work.

The primitive learning process contains three phases: review, testing and result viewing. **Review** is simple reading of terms with the purpose to renew information in user short-term memory [5]. It depends on *size of information* and *dictionary* parameters. The second phase called **testing** represents tasks with learned words to check user knowledges. The result of this phase will be added to database for further processing and prediction. The **result viewing** is the phase where user can see points of the test and list of words that he need to repeat. Learning is the cyclic process, and it finishes when user has no errors in testing.

The interface of system must be comfortable and intuitive, user should not be distracted by unclear UI elements and focus on learning only.

Conclusions. Learning Management Systems are becoming more and more popular and one of problems of these systems is that they usually cannot automatically configured to user learning capabilities. In this work we reviewed methods of configuration based on test results and the algorithms which enable adaptation to user on example with learning foreign words.

Auto-configured algorithm analyzes all previous attempts to learn the list of terms (pairs key-value, word-meaning etc.) and calculates parameters for the next iteration. Each attempt consists of two main parts: words review and test.

Spaced repetition technique allows user to learn the words by intervals, which cannot be too small to avoid brain overload, and cannot be too large in order to exclude information forgetting.

References. 1. Types of Learning Management Systems. [Online]. Available: <http://learning-management.financesonline.com/types-learning-management-systems/>. 2. Alan D. Baddeley “Human Memory: Theory and Practice”, Psychology Press Ltd, East Sussex, UK, 1997. 3. John R. Anderson, Robert Milson “Human Memory: An Adaptive Perspective”, American Psychological Association Inc., Vol. 96, No. 4, pp. 703–719, 1989. 4. Spaced Repetition. [Online]. Available: http://en.wikipedia.org/wiki/Spaced_repetition. 5. Hikaru Takeuchi, Atsushi Sekiguchi, Yasuyuki Taki, Satoru Yokoyama, Ryuta Kawashima “Training of Working Memory Impacts Structural Connectivity”, Journal of Neuroscience 3 March 2010, 30 (9) 3297–3303.

Kopp A.M., Orlovskyi D.L.

National Technical University “Kharkiv Polytechnic Institute”, Kharkiv, Ukraine

An approach to measure similarity of business process models

Introduction. At higher levels of BPM (Business Process Management) maturity, organizations tend to accumulate considerable amounts of business process models [1]. Thus, business process model repositories may contain hundreds or even thousands models, represented using various modeling notations [2]. Since business process modeling technique is used by organizations to describe knowledge about their activities, the problem of store, share, and reuse of organizational knowledge, represented using business process models, becomes relevant.

In this paper a similarity measure, used to retrieve business process models from a repository in order to their further reuse in a business process continuous improvement cycle according to BPM concept, is proposed. The problem of similar business process models retrieving has been earlier considered in studies [1, 3, 4], which propose several measures, based on label similarity, structural similarity, behavioral similarity etc.

Details. Earlier we have proposed using of the knowledge representation model RDF (Resource Description Framework) to describe business process models, developed in organization. RDF is based on “subject-predicate-object” statements, which are convenient for machine processing [5]. A set of such statements may be represented as a marked directed graph.

Proposed RDF Schema, used to describe a business process model as the RDF graph, includes classes and properties, such as:

1. *FlowObject* – class, which includes process flow objects, such as functions (*Function*), processes (*Process*), events (*Event*), and gateways (*Gateway*), related using *triggers* property.
2. *OrganizationalUnit* – class, which includes organizational units, such as departments (*Department*) and roles (*Role*), related to functions and processes using *executes* property.
3. *SupportingSystem* – class, which includes supporting IT-systems, related to functions and processes using *usedBy* property.

Therefore, maintenance of a repository of business process models, represented as RDF graphs, may allow various possibilities, such as store and retrieve knowledge about organizational activities, and its further reuse to design new or improve existing business processes. Improvement of an existing organizational business process, according to BPM concept, assumes selection of its design variants and further transformation using obtained recommendations.

Thus, business process models, similar to an existing business process model, should be retrieved from the business process model repository. Hence, the similarity measure of two business process models $BPModel_1$ and $BPModel_2$, represented using RDF graphs, is proposed:

$$\begin{aligned}
 &BPModelSim(BPModel_1, BPModel_2) = \\
 &= \alpha_1 \cdot \left[\beta_{lab} \cdot Sim(U_1, U_2) + \beta_{str} \cdot Sim \left(\bigcup_{u \in U_1 \cap U_2} executes_1(u), \bigcup_{u \in U_1 \cap U_2} executes_2(u) \right) \right] + \\
 &+ \alpha_2 \cdot \left[\beta_{lab} \cdot Sim(S_1, S_2) + \beta_{str} \cdot Sim \left(\bigcup_{s \in S_1 \cap S_2} usedBy_1(s), \bigcup_{s \in S_1 \cap S_2} usedBy_2(s) \right) \right] + \quad (1) \\
 &+ \alpha_3 \cdot \left[\beta_{lab} \cdot Sim(F_1, F_2) + \beta_{str} \cdot Sim \left(\bigcup_{f \in F_1 \cap F_2} triggers_1(f), \bigcup_{f \in F_1 \cap F_2} triggers_2(f) \right) \right],
 \end{aligned}$$

where $\alpha_i, i = \overline{1, 3}$ – weights represent importance of organizational units, supporting IT-systems, and process flow objects similarity respectively, $\alpha_1 + \alpha_2 + \alpha_3 = 1$; β_{lab} and β_{str} – weights represent importance of label (unique resources used in both RDF graphs) and structural (unique statements used in both RDF graphs) similarity respectively, $\beta_{lab} + \beta_{str} = 1$; U – a set of organizational units;

S – a set of supporting IT-systems; F – a set of process flow objects; Sim – a similarity measure of two sets, e.g. Jaccard index J ; $executes(u)$ – a set of processes, executed by the organizational unit $u \in U$; $usedBy(s)$ – a set of processes, supported by the IT-system $s \in S$; $triggers(f)$ – a set of process flow objects, triggered by the object $f \in F$.

Proposed similarity measure of two business process models (1), outlined above, is normalized $BPMModelSim(x, y) \in [0, 1]$, symmetric $BPMModelSim(x, y) = BPMModelSim(y, x)$, and reflexive $BPMModelSim(x, x) = 1$, where x is existing business process model and y is a model, retrieved from the business process model repository.

Result. As an example, two models (fig. 1), which describe goods order business process in EPC (Event-driven Process Chain) notation, were translated to RDF graphs. Similarity of these models was calculated using the proposed measure (1).

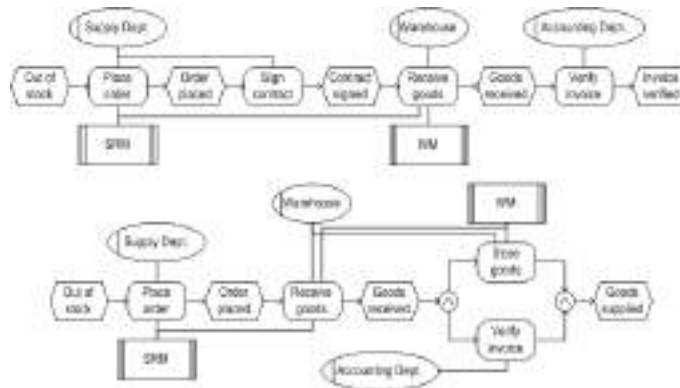


Figure 1. Sample pair of business process models (SRM – supplier relationship management system, WM – warehouse management system)

According to the obtained results, considered business process models are totally similar ($BPMModelSim$ is 1) if only organizational units or supporting IT-systems are considered as important ($\alpha_1 = 1$ or $\alpha_2 = 1$) to compare only by the label closeness ($\beta_{lab} = 1$). By the other side, comparison of these models only by the structural closeness ($\beta_{str} = 1$) if only process flow objects are considered as important ($\alpha_3 = 1$), has shown that considered models are not that similar ($BPMModelSim$ is 0,38). With considering that organizational units, supporting IT-systems, and process flow objects are equally important ($\alpha_1 = 0,33$, $\alpha_2 = 0,33$, and $\alpha_3 = 0,34$) to compare both by the label and structural closeness ($\beta_{lab} = 0,5$ and $\beta_{str} = 0,5$), we have defined that these models are quite similar ($BPMModelSim$ is 0,69).

Conclusion. In this paper we have proposed the similarity measure between business process models. In contrast with already known measures [1,3,4], it combines both label and structural approaches to define similarity of business process models, described using knowledge representation model RDF. Besides, proposed measure allows considering similarity not only by the process flow, but also by the whole process environment, including organizational units and supporting IT-systems.

References. 1. M. Dumas, L. Garcia-Banuelos, and R.M. Dijkman, “Similarity search of business process models”, Bulletin of the IEEE Computer Society Technical Committee on Data Engineering, vol. 32, no. 3, pp. 23–28, 2009. 2. Z. Yan, R. Dijkman, and P. Grefen, “Business process model repositories – Framework and survey”, Information and Software Technology, vol. 55, pp. 380–395, 2012. 3. R. Dijkman et al., “Similarity of business process models: Metrics and evaluation”, Information Systems, vol. 36, no. 2, pp. 498–516, 2011. 4. B. Van Dongen, R. Dijkman, and J. Mendling, “Measuring similarity between business process models” // Seminal Contributions to Information Systems Engineering, pp. 405–419, 2013. 5. RDF. [Online]. Available: <https://www.w3.org/RDF/>.

Kurylenko O.M.

Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

Credit scoring models development

Introduction. In the modern world, the practice of drawing upon a credit is widespread amongst individuals for supplying their specific needs. But since crediting is an unexpected and risky activity, banks expend much effort to avoid money losses. For that purpose, credit scoring models are built, that are tools for optimization of the process of crediting individuals. By means of these tools, banks can assess the borrower's portfolio based on the previous crediting history the bank already has.

In this paper, the approaches to analyzing input data and building models of crediting individuals are described [1].

Input data. As the input data, the dataset consisting of 11 features and 150000 records was taken. The features included both the client personal and credit history information:

- total balance on credit cards and personal lines of credit except for real estate and no installment debt divided by the sum of credit limits;
- age of borrower in years;
- number of times borrower has been 30–59 days past due but no worse in the last 2 years;
- monthly debt payments, alimony, living costs divided by monthly gross income;
- monthly income;
- number of open loans and lines of credit;
- number of times borrower has been 90 days or more past due;
- number of mortgage and real estate loans including home equity lines of credit;
- number of times borrower has been 60–89 days past due but no worse in the last 2 years;
- number of dependents in family excluding themselves;
- the feature, which indicates whether the person experienced 90 days past due delinquency or worse.

Data preprocessing. The most important part of building any model is data preprocessing, as this stage predefines the productivity of the future model. First of all, all records with two or more NA (missing) values were removed from the dataset, as they are regarded as which contain no sufficient information. Besides that, records with Income feature equal to zero were removed, as it's irrational to loan out to the individual with no income.

As there were still a lot of NA's left, they were replaced with the values received by the predictive method. Thus, the model for predicting missed values based on the existing values was built, as this approach helps to get more realistic data in the dataset as distinct from filling NA with just mean or median values.

Afterwards, the initial dataset was split into train and test subsets. This is needed to train the future models and assess their performance [2, 3].

Feature importance. After the data preprocessing stage, the feature importance of the received dataset was explored to obtain the feature ranking. This was done to understand which features are more important for the prediction. Thus, the feature ranking was obtained based on the mean decrease in accuracy method using the Random Forest classification model. This method directly measures the impact of each feature on the accuracy of the model what is done by permuting the values of each feature and measuring how much the permutation decreases the accuracy of the model [4].

As the consequence of this stage, the result showed in Fig. 1 was obtained.

Modelling. At the previous step there was built the Random Forest model and the probabilities of 90 days past due delinquency or worse were calculated and the neural network model for credit scoring with 2 hidden layers was built afterwards. In the neural network, Adam activation function was used in the hidden layers and Softmax activation function in the output layer.

One of the most important things on the stage of modelling is detecting the correct threshold so that both sensitivity and specificity are high. As it's supposed that loaning out to the individual

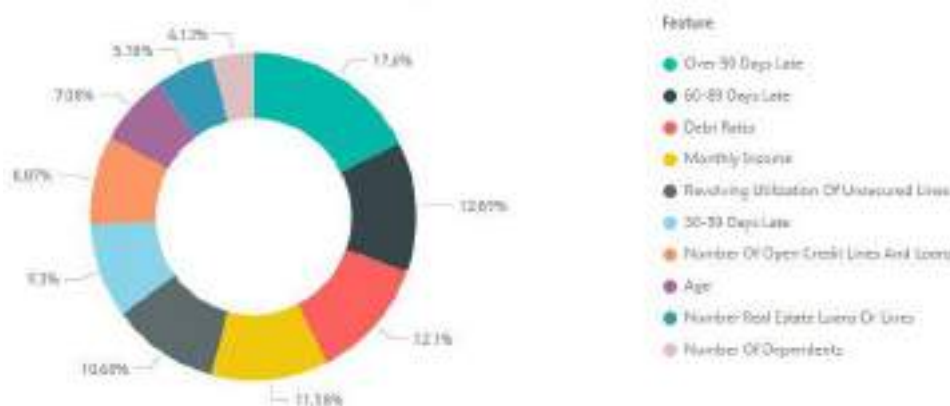


Figure 1. Feature importance

who won't give the money back is 5 times worse than not loaning out to the individual who will give the money back, the threshold must be chosen to provide the needed balance between specificity and sensitivity and give the maximum income and minimum losses for the bank.

The threshold was found to be 0.175 with specificity equals to 93.64% and sensitivity equals to 53.5% [3, 5].

Results. As the result 2 models were developed: Random Forest model and neural network model [3]. They were assessed by AUC metric and it appeared, that Random Forest gave a better result than neural network.

The result showed in Figure 2 was obtained.

Random Forest	0.86333
Neural Network	0.86111

Figure 2. AUC metrics for the models

Conclusions. The whole process of developing credit scoring models was done and 2 models were built as the result. Random Forest showed a better result than the neural network, but it must be noticed that neural networks work better on more data.

Because banks usually have a lot of historic data for training models and amount of the data increases year by year, it is obvious that it is preferable to use the neural network for the real-life data.

References. 1. Сиддики Н. Скоринговые карты для оценки кредитных рисков. Разработка и внедрение интеллектуальных методов кредитного скоринга / Н. Сиддики; пер. С англ. Е. Ильичева. – М.: Манн, Иванов и Фербер, 2014. – 268 с. 2. “MeasuringU: 7 Ways to Handle Missing Data”, MeasuringU.com, 2018. [Online]. Available: <https://measuringu.com/handle-missing-data/>. 3. Bielecki T.R. Credit Risk: Modeling, Valuation, Hedging / T.R. Bielecki, M. Rutkowski. – Berlin: Springer, 2002. – 500 p. 4. “Selecting good features – Part III: random forests | Diving into data”, Blog.datadive.net, 2018. [Online]. Available: <http://blog.datadive.net/selecting-good-features-part-iii-random-forests/>. 5. “Confusion Matrix in Machine Learning – GeeksforGeeks”. [Online]. Available: <https://www.geeksforgeeks.org/confusion-matrix-machine-learning/>.

Moroz A.M., Globa L.S.

Igor Sikorsky Kyiv Polytechnic Institute, Institute of Telecommunication Systems, Kyiv, Ukraine

Data Mining and its applications

Telecommunication companies generate amount of data (define as a big data). These data consist of call detail data, network data, and customer data. To process such a large amount of data telecom industry using Data Mining. In this research, we will try to discuss the methods of Data Mining and examine its main facets and concepts taking into account the specifics of data processing based on telecom operator infrastructure.

Introduction. Previously, applied statistics was widely used in the field of processing knowledge in telecom industry. The statistics evaluated, tested hypotheses, but gave rough and average results. Thanks to technical progress, people began to store huge amounts of information, which was heterogeneous, and naturally, had to be processed. The methods of statistics did not make it possible to predict and control processes and had a highly computational complexity from telecom operator point of view. The telecom operator could no longer solve the tasks by applying statistics. It is necessary to structure large amounts of data and to extract useful information from them in order to:

- analyze this data successfully in the future;
- the fuzzy knowledge base designing with the help of fuzzy logic;
- visualize these data for easy human perception.

The purpose of this paper is to identify the features of existing mathematical methods named new Data Mining used by telecom operators for solving problems arising from new information technology such as IoT, M2M and others. In more detail, we will consider the method of processing big data – Data Mining.

Nowadays, the Data Mining solutions are applied by telecom operators such as network operators and Internet providers with the aim to improve processing of different kind unstructured large data volumes for getting decisions that are more intelligent. Analytics enables telecom and Internet providers to increase economic effectiveness and efficiency in fields of service provisions significantly.

Data mining is the process of detecting in raw data previously unknown non-trivial but practically useful and accessible interpretation of knowledge necessary for decision-making in various spheres of human activity [1]. This technology helps to identify hidden links in databases of very big sizes. Since Data Mining has developed based on applied statistics, methods of artificial intelligence, database theory and so there are a lot of application for data processing. However, all this methods and algorithms are not very effective for big data processing [1].

Data Mining uses the concepts of averaging over a sample, the concept of templates and so on. However, there is no universal way to determine the meaning and connection of data among themselves, and additional studies are required for each subject area in order to identify these links [2].

The purpose of Data Mining is to find such models that can not to be found by the usual methods. There are two types of models: predictive and descriptive. Predictive models: positioned on a data set with known results [3].

These are classification models (describe the rules by which the description of an object can be attributed to one of the classes) and the sequence model (they describe the functions by which the change of continuous numerical parameters can be predicted).

Descriptive models: they pay special attention to the nature of the dependencies in the data set, the mutual influence of various factors, the construction of empirical models. There are easy for human perception.

Only the volume of information limits the scope of data mining technology. Therefore, managers and analysts who use this technology can get tangible advantages in the competition. Telecom-

munications – identifying the categories of customers with similar stereotypes of using services, consequently – the development of pricing policy, analysis of failures, prediction of peak loads.

Billing in telecommunications – complex of processes and solutions at communication enterprises responsible for collecting information on the use of telecommunications services, their tariffication, invoicing subscribers, processing payments.

Because Data Mining has developed at the junction of many areas, it is possible to reuse classes and methods of this technology: neural networks, decision trees, genetic algorithms, limited search algorithms, associative rules, cluster analysis and much more [4]. Moreover, methods of data mining allow to solve the problems of structural engineering design for innovative technical systems effectively in the telecom industry [1]. These methods have much in common with methods for solving problems of classification, diagnosis and pattern recognition. Nevertheless, one of their main distinctive features is the function of the regularities interpreting that form the basis rules for the objects inclusion in equivalence classes. Therefore, logical methods are becoming more common today. There is another important reason to determine the priority of logical methods. It lies in the complex systemic organization of the areas that constitute the application subject of modern information technologies.

The development of intelligent data processing methods for telecom companies is necessary for:

- Reducing the computational complexity of methods for processing big data for the services provision by operator to a subscriber with a given quality of service;
- Operator should to predict the risks that may arise from the operation of the telecommunications system;
- The operator should be able to identify faults in the system and find out the cause of their occurrence.

Some of the existing algorithms can be adapted to compute large distributed information arrays. At the same time, serious difficulties can arise with a visual representation of the results – because of the huge amount of information entering the input, the number of dissimilar reports at the output sharply increases. For their convenient presentation, new mathematical methods are needed, which are fundamentally different from report generators used for traditional storage.

Conclusion. In this overview the big data processing methods relevant for telecom operator were analyzed. There are lot of problems for data processing methods choosing:

- the problem of big data – need mathematical methods for structuring data and knowledge (in future work it's proposed to use ontology) and very important to find methods that do not have high computational complexity;
- the problem of prediction – proposed to use the knowledge base patterns and fuzzy logic.

All of the above problems can be eliminated by developing new methods for large data processing that will: successfully analyze these data, a fuzzy knowledge base designing, and visualize these data for easy human perception. The further work consists in studying methods of data structuring and fuzzy knowledge base designing for the telecom operator.

References. 1. Dr. Mehmet Ulema. (Professor of Computer Information Systems Manhattan College, Riverdale New York, USA) “Big Data and Telecommunications” [Electronic resource], 4th International Black Sea Conference on Communications and Networking, 2016. 2. Technologies of data analysis: DataMining, VisualMining, TextMining, OLAP training. allowance. 2 nd ed. / A.A. Barsegyan, M.S. Kupriyanov, V.V. Stepanenko, I.I. Kholod. SPb.: 2007. 59 p. 3. Educational-methodical complex on the discipline “Information Technologies” of the Institute of Economics and Management [Electronic resource]. URL: <http://www.sergeeva-i.narod.ru/inform/page11.htm>. 4. Chubukova I.A. Course DataMining [Electronic resource]. URL: <http://www.intuit.ru/departament/database/datamining/>.

Olefir O.S., Chernenko V.M.

Igor Sikorsky Kyiv Polytechnic Institute, Faculty of Applied Mathematics, Kyiv, Ukraine

System of logistics of material resources based on social networks

Introduction. The problem of finding the shortest path on the graph is one of the most important classical problems in the *theory of graphs*. To date, there are many algorithms for solving it. The significance of this task is determined by its various practical applications [1]. For example, GPS-navigators search for the shortest path between intersections [2]. In this article it is proposed to apply modifications of search algorithms on graphs to find the optimal way of delivering goods.

Comparison of algorithms. *Bellman-Ford's algorithm* can work with edges with negative weights, in contrast to the Dijkstra's algorithm. The most important difference, which gives the advantage of the Bellman-Ford's algorithm its relatively low complexity. It is $O(nm)$, where n is the number of vertices and m is the number of edges [3]. The reason for the higher velocity of the Bellman-Ford algorithm is that the vertices are considered only once, whereas in the Dijkstra's algorithm they are considered at each iteration to find the minimum vertex, which gives additional time for processing the algorithm.

The Floyd-Warshall's algorithm is the easiest in program implementation compared with the algorithms of Dijkstra's and Bellman-Ford's. The specified matrix, which is filled with the weights of the edges for each of the vertices. The advantage of the algorithm in the simplicity of its implementation, weakness – in the complexity is $O(n^3)$.

The complexity of the *Dijkstra's algorithm* depends on the method of finding the vertex V , as well as how to store unexpected vertices and how to update the tags. In the simplest case, when an infinite number of vertices is considered for a vertex with minimum $d[V]$, and for the storage of values of d , an array is used, the complexity of the algorithm is $O(n^2)$.

On cycles on neighboring vertices, the number of operations is to be proportional to the number of edges m (since each edge meets in these cycles exactly twice and requires a constant number of operations). Thus, the total complexity of $O(n^2 + m)$.

For *rarefied graphs*, non-visited vertices can be stored in a binary stack. As a key use the value of $d[i]$, then the time of removing the vertex U will be $\log(n)$, with the modification time $d[i]$ will increase to the $\log(n)$. Since the loop is executed in the order of n times, and the number of relaxations (changes in the labels) is not more than m , the complexity of such implementation is $O(n \cdot \log(n) + m \cdot \log(n))$.

If storing unpaired vertices to use a Fibonacci's heap [4], for which the removal takes place on average for $O(\log(n))$, and a decrease in the value of the average for $O(1)$, then the operating time of the algorithm will be $O(n \cdot \log(n) + m)$, however, the concealed constants in Asymptotic estimates of complexity are quite large and the use of Fibonacci heap is rarely feasible: conventional binary heaps in practice are more effective.

Conclusions. The basis of the developed system is the concept of social network, which today is a common means of communication between users. This allows you to interact with the system as simple and intuitive as possible. In the course of researching algorithms for finding the shortest path, two cases are considered: for dense and sparse graphs. Dense graphs represent a mathematical model for describing the network of roads in the middle of the metropolis, and the rare graphs – a network of roads between individual cities.

According to the research, it was found that the search for the shortest path for dense graphs is best suited to the Dijkstra algorithm using a binary heap, we obtain a decrease in the complexity to $O(n \cdot \log(n) + m)$.

References. **1.** Dijkstra, E. A note on two problems in connexion with graphs (1959). **2.** R. Hanneman M. Riddle Introduction to Social Networks Methods (2005). **3.** Cormen, Thomas H.; Leiserson, Charles E.; Rivest, Ronald L. Introduction to Algorithms (1990). **4.** Fredman, Michael Lawrence; Tarjan, Robert E. Fibonacci heaps in network optimization algorithms (1984).

Sergiyenko A.M., Qasim M.R.

Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

Modeling of the wave propagation in the solid bar

Introduction. The acoustic processes in solids are usually modeled in the computers using the finite difference method, which affords the huge computational resources. On the contrary, the digital waveguide (DWG) method provides the high-speed computations utilizing much less computation volume [1]. A lot of successful examples of the DWG implementation are shown in modeling and sound synthesis of the string and wind musical instruments [2,3]. But this model does not take into account the dispersion in the sound propagation.

In this work, a DWG method modification is proposed which provides the modeling of the sound dispersion in the solid bar.

DWG model basics. DWG is a digital model of the wave propagation in the ideal waveguide. It is based on the principles, described in the work [4], but adapted to the sound wave modeling. The forward $f_i = Rv_f$ and backward $b_i = -Rv_b$ waves are considered in the i -th point of the waveguide, where f_i and b_i are the pressure of the forward and backward (reflected) waves, respectively, R is the wave impedance, v_f, v_b are the particle velocities. For the solid cylinder with the cross-section A the impedance is equal to $R = \rho c/A$, where ρ is the matter density, c is the velocity of the longitudinal waves. The real pressure value is equal to $u = f_i + b_i$.

In the DWG model, the signals are sampled at a sampling frequency of F_s which is at least two times greater than the maximum considered frequency. Then the i -th section of the waveguide of the length L looks like two delay lines at $n = L/(c_o F_s)$ cycles, where c_o is the wave velocity.

The homogeneous parts of the waveguides are connected together by the adapters, also called as the scattering nodes. The adapter function is designed to satisfy the Kirchhoff's law. For the case of joining two waveguides:

$$v_b1 = rv_f1 + (1 - r)v_f2; v_b2 = (1 + r)v_f1 - rv_f2; \quad (1)$$

where $r = (R_2 - R_1)/(R_2 + R_1)$ is the reflection coefficient. Similarly, when three waveguides are connected, then the adapter node calculates the following formulas:

$$A_o = (2g_1 - 1)A + 2g_2B + 2g_3C; \quad (2)$$

$$B_o = 2g_1A + (2g_2 - 1)B + 2g_3C;$$

$$C_o = 2g_1A + 2g_2B + (2g_3 - 1)C;$$

where A, B, C are the waves entering the node, A_o, B_o, C_o are the waves leaving the node, g_1, g_2, g_3 are the specific impedances of the waveguides, and $g_1 + g_2 + g_3 = 1, g_i = R_i/(R_1 + R_2 + R_3)$.

DWG model improvement. Due to the theory, the primary longitudinal wave with the phase velocity c_p , and the secondary transverse wave with the velocity $c_s, c_s < c_p$ are propagated in the solid bars [5]. The waves of different types can be transformed into each other interacting with the media boundaries. Besides, the longitudinal waves have a dispersion, i.e. its velocity c_p depends on the wave length λ . The velocity c_p in a cylinder is approximated as follows [5]:

$$c_p = c_o(1 - \nu^2 p^2 a^2 / \lambda^2) = f(\lambda), \quad (3)$$

where ν is the Poisson ratio, which is equal to 0.29 for the steel, a is the cylinder radius. Therefore, the DWG model has to be corrected according to this formula.

The modified DWG model of a bar is illustrated by Fig. 1. This model consists of the left LA and the right RA adapters, the waveguide

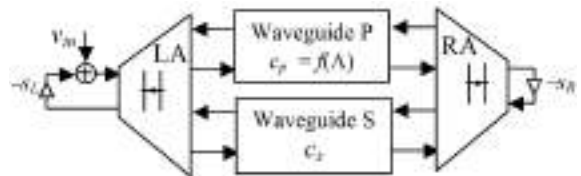


Figure 1. Structure of the solid bar model

P of longitudinal waves, and the waveguide S

of transverse waves. Each three port adapter compute the equations (2). Its port, which means the bar end, is connected to the network, which implements the wave reflection with the suppression factor s_L or s_R due to (1). The actuating signal v_{in} is fed to the model through an adder.

The waveguide S performs the delay to the given number of clock cycles and some wave attenuation. The waveguide P does the same but it has a set of channels. Each of them has the separate frequency band, delay, and attenuation depending on the respective wave length λ according to (3).

Experimental results. Up to eight channels of the waveguide P are implemented on the basis of the allpass filters described in [6]. Each channel has the dynamically tunable band pass filter, which is implemented as is shown in [7]. The resulting amplitude-frequency characteristic of the waveguide S is shown in Fig. 2. Here, the frequency f is measured in parts of the sampling frequency F_s , the characteristics of two neighboring channels are shown as well.

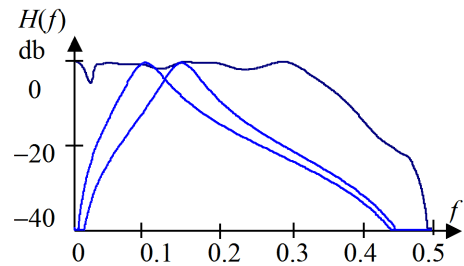


Figure 2. Amplitude-frequency characteristic of the waveguide P

The model was actuated by a narrow ultrasound impulse, and the velocity v was measured in the output of the adder node (see Fig. 1). The time diagram of it which represents the wave signal is shown in Fig. 3. This diagram shows that the waves are really dispersed after the propagation through the bar model.

The model is described by the VHDL language and therefore, it can be implemented both in the VHDL simulator and in FPGA. The model described in the style for the synthesis and configured in Xilinx Spartan-6 FPGA contains 2680 logic cells, 60 DSP48 blocks, 16 BlockRAMs, and provides the sampling frequency $F_s < 100$ MHz. This provides the modeling in real time.

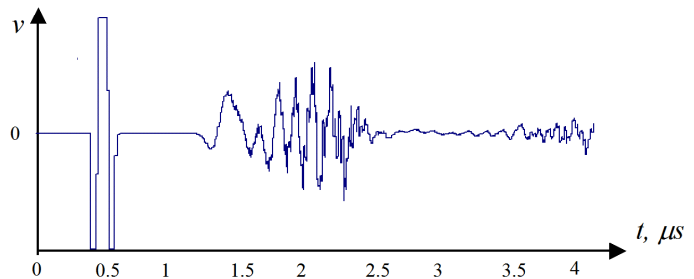


Figure 3. Time diagram of the wave signal in the solid bar model

Conclusion. A modified DWG

method for modeling the solids is proposed which takes into account the dispersed wave propagation. The method is based on the digital wave guides which delay is depending on the wave length. An example of the solid bar model described in VHDL shows the effectiveness of this model.

References. 1. D.R. Bergman “Computational Acoustics: Theory and Implementation”. J.Wiley and Sons, Ltd. 2018. – 296 p. 2. J.O. Smith III “Physical Modeling Using Digital Waveguides”. Computer Music Journal V.16. 1993. №4. – P. 74–91. 3. M. Karjalainen, and C. Erku “Digital Waveguides versus Finite Difference Structures: Equivalence and Mixed Modeling”. EURASIP Journal on Applied Signal Processing 2004. №7. – P. 978–989. 4. A. Fettweis “Wave digital filters: Theory and practice”. Proc. IEEE, V74, 1986. №2. – P. 270 – 327. 5. H. Kolsky “Stress Waves in Solids”. Dover Publications Inc. 2012. 6. P.A. Regalia, S.K. Mitra, and P.P. Vaidyanathan “The Digital All-Pass Filter: A Versatile Signal Processing Building Block”. Proc. IEEE. V.76. 1988. – 76, №1. – P. 19–37. 7. A.M. Сергиенко, Т.М. Лесик “Перестраиваемые цифровые фильтры на ПЛИС”. Электронное моделирование, Т. 32. 2010, №6. С. 47–56. (A.M. Sergiyenko, and T.M. Lesyk “Tunable digital filters in FPGA” Electronic Modeling. V.32. 2010. №6. – P. 47–56).

Tarasenko-Klyatchenko O.V., Tuparjeva V.A.

Igor Sikorsky Kyiv Polytechnic Institute, Faculty of Applied Mathematics, Kyiv, Ukraine

Algorithm of authentication of objects of the network on the basis of analysis of signal parameters on the physical level

Introduction. Due to the widespread use of wireless devices, the number of cases of their impact was increased by various attacks. Devices that operate in a flexible data environment have many features and advantages, but also drawbacks, among which the availability of potential attackers can be highlighted. Also, wireless networks are open to intruders from outside without physical connection, which affects the effectiveness of security methods in such a network.

Description. The authentication algorithm is based on tracking the characteristics of the channel of communication for each pair of transmitter-receiver. To estimate this characteristic, it is suggested to use the impulse response of the frequency band channel occupied by the transmitted signal. Let such an assessment be considered a radio rebate. The idea of the algorithm should be described as the following sequence of actions:

1. A database of current measurements for the legal users is created, the so-called measurement history, consisting of the records of the radio-reciprocal measurement of the channel between the receiver (j) and each transmitter (i) for $N - 1$ measurements [2]. All of them are written in the form of a matrix. To take into account the influence of random temporal characteristics of the channel and noise, [3] the average distance between the $N - 1$ measurement in (h) is calculated.
2. Another measurement of the radio redirection of the channel is carried out.
3. Calculation [4] of the minimum Euclidean distance between the current measurement and the measurement history.
4. Matching all values, (d). We compare the values of (d) with the boundary value γ , where ($\gamma \leq 0$). If distance, ($d \geq \gamma$), it is decided that the difference between the measured radio rebate of the channel and the history is due to the fact that the measured value belongs to another channel communication. Fix the wrong authentication.
5. If the distance is ($d \leq \gamma$), then we conclude that the position of the object is confirmed (not changed) and the new value (h) is included in the (H) database, and the oldest value is deleted. Next, the probability of correctly detecting that the transmitter's position has not changed and the probability of false alarms is determined. Let it be the permissible position of the transmitter, and i is the transmitter-infringer. It is necessary to determine if i is not an object i [3].

Conclusions. One of the most important tasks of information security is to ensure secure authentication of users. The study proposed a method of authentication, which is based on the improved history of radio-redirection of the channel in each pair of transmitter-receiver. To perform authentication, a database of radio channel redirects is created. Calculates the minimum Euclidean distance between the current measurement and the measurement history. This value is compared with the limit values and determines the existence of the user's access rights. Estimates are made of the probability of detecting that the transmitter's position has not changed and the probability of an access error for systems with one and more receivers.

References. 1. D.E. Denning and P.F. MacDoran. Location-based authentication: Grounding cyberspace for better security. Feb. 1996. 2. D. Ghogare, Swati P. Jadhav, Ankita R. Chadha, Hima C. Patil. Location Based Authentication: A New Approach towards Providing Security Shraddha. 3. Gus German, Quentin Spencer, Lee Swindlehurst, Reinaldo Valenzuela. Wireless indoor channel modeling: statistical agreement of ray tracing simulations and channel sounding measurements. 4. Kaya, A.O., Greenstein, L.J., Trappe, W. Characterizing indoor wireless channels via ray tracing combined with stochastic modeling. IEEE Transactions on Wireless Communications.

Terentiev O.M., Makohon R.O.

Igor Sikorsky Kyiv Polytechnic Institute, Institute for Applied System Analysis, Kyiv, Ukraine

Analysis of causality Bayesian network usage for data-mining purposes

Modern world is constantly progressing, so enormous amounts of information are being generated by various areas of the human life. And the speed if it's creation is only increasing, due to concepts like an internet of things, globalization and developments in the computer miniaturization. All this raw data contains useful information that may help improve performance of the analyzed subject. This confirms the relevance of the data-mining methods usage to find hidden relations within the dataset variables, determine yet unknown dependencies and patterns.

The purpose of this research is to analyze the causality Bayesian network usage as a data-mining tool for the real economical problem, namely the estimation of the person's creditworthiness, discover it's strengths and weaknesses.

As an algorithm for the causality Bayesian network creation the heuristic one was chosen. It is based on such values estimations as the MI (Mutual Information) value and the MDL (Minimum Description Length) one. This model was described in works of researchers like Zheng Y., Kwoh C.K. [1] and Heckerman D., Geiger D., Chickering D.M. [2].

The heuristic method was used as a way to drastically decrease the number of needed computational powers to build the accurate Bayesian network. The following table shows the amount of possible networks that machine learning algorithm may build out of number of nodes in a raw data.

Table 1. Exhausted search disadvantages

Number of nodes	Number of acyclic models
1	1
2	3
3	25
4	543
5	29281
6	3781503
...	..
15	$2.38 \cdot 10^{41}$
20	$2.34 \cdot 10^{72}$

The computer program builds the Bayesian network as an acyclic directed graph. Each node there stands for some variable in Bayesian sense. It may be either observable quantile or latent, unknown, hypothetic variables in the training data. Each edge stands for a conditional dependency of the children node from its parent one. So there is an opportunity to calculate all children probabilities, if parent's probabilities are known. The graph nature of the instrument makes the output results visually convenient for perception and analysis.

The heuristic method if based on the mutual information value between pairs of variables in a dataset (1).

$$MI(x^i, x^j) = \sum_{x^i, x^j} P(x^i, x^j) * \log \frac{P(x^i, x^j)}{P(x^i)P(x^j)}. \quad (1)$$

This value provides information about the order of building the optimal dataset. The method begins from the pair of nodes that has the biggest value of MI. There are three possible outcomes for each such pair: two ways of one variable being parent for the other one, and one opportunity for them not to be directly connected. So the method chooses the best outcome of these three ones, using the value of MDL, based on coding theory, suggested by Shannon. It states that there is a fixed minimal length of the code that encodes some message, that fully depends on the message information. So if the code is based on the incorrect perception of the message, its length would be

bigger than the optimal one. When applied to Bayesian networks, this function is determined as

$$L(j, g, x^n) = H(j, g, x^n) + \frac{k(j, g)}{2} * \log(n). \quad (2)$$

So having built Bayesian network by minimizing MDL on each step of the algorithm, optimal structure will be achieved.

The SAS Base programming language was used for building and learning of the analyzed Bayesian networks. No ready-made functions were used to ensure code's high performance and it's conformity with the explored learning method. The data set, which was used to create the network structure, was provided by the forecasting the probability of non repayment of loan competition, organized by <https://sascompetitions.ru>. It originally contained 34 variables in total of 1787571 observations.

For the accuracy checking purposes more than 10 different test data sets were analyzed by the coded Bayesian network. These examples were provided by SASHelp datasets. Final resulting network structures were validated by comparison with the corresponding GeNIe (<http://genie.sis.pitt.edu/>) and BayesiaLab (<http://bayesia.com/>) results. One of the analyzed examples contained 8 nodes, that would take over $7,83 * 10^{11}$ probable networks to examine for the exhausted search. The iterative heuristic search used 28 iterations to consider 120 different possible structures. The needed structure was built at the 81-th try on a 15-th iteration. Here is the graph representation of the program result:

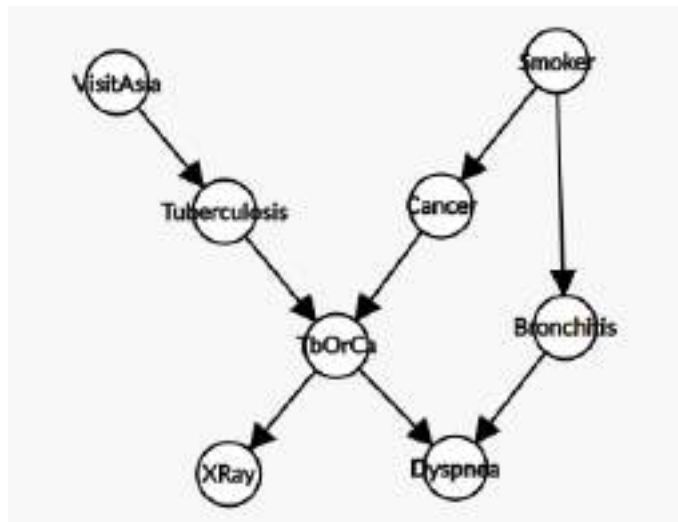


Figure 1. Example of program result network

It is important to state that implemented model allows building Bayesian network using training data, unlimited by the size of the data set and the number of the input variables. And the accuracy of implemented algorithm was equal to the one of the compared program products.

Conclusion. The causality Bayesian network building program product was developed, that allows construction of the network structure, using training data. The accuracy and the performance of the implemented algorithm was successfully tested in comparison with the existing program solutions. The product had demonstrated it's ability to fulfill the described task.

References. 1. Zheng Y. and Kwoh C.K. "Improved MDL Score for Learning of Bayesian Networks", Proceedings of the international conference on artificial intelligence in science and technology, pp. 98–103, 2004. 2. Heckerman D., Geiger D., Chickering D.M., "Learning Bayesian Networks: The Combination of Knowledge and Statistical Data," Machine Learning, vol. 20, no. 3, pp. 197–234, September 1995. [Online]. Available: <https://link.springer.com/article/10.1007/BF00994016>. [Accessed March 30, 2018].

Арчвадзе Н.Н.¹, Пховелишвили М.Г.²

¹Тбилисский государственный университет им. И.Джавახишвили, Тбилиси, Грузия; ²Институт вычислительной математики им. Н.Муслишвили Грузинского технического университета, Тбилиси, Грузия

Применение параллельных данных для прогнозирования сложных процессов

Обычно при прогнозировании анализируются те явления, которые предшествуют прогнозируемому явлению. Из предшествующих явлений некоторые могут быть достаточными, другие – обязательными условиями, при этом для выполнения некоторого события нужны существования всех предшествующих явлений, а для некоторых – одно или несколько явлений. Массив явлений, по своей природе может быть последовательным или параллельным. Например, если рассмотрим задачу прогнозирования землетрясений, то на него влияют ряд важных свойств сейсмического процесса, таких как сейсмический цикл, фор шоковая и афтер шоковая активность, постепенный характер убывания сейсмических событий после сильнейших событий [1]. При этом сильным землетрясениям предшествуют либо аномальное поведение сейсмической активности, либо ее аномальное понижение, либо их колебания в различных пространственно-временных зонах [2]. В [3] по мнению авторов, основным источником землетрясения в ближайшем будущем может быть объявлено электромагнитное излучение Земли VLF/LF.

Общепризнанные физика–механические модели описания сейсмического процесса по сей день не существуют. Появление суперкомпьютеров, с мощными вычислительными процессорами, дают новые возможности в автоматизации процессов прогнозирования. В качестве модели сложных прогнозирования процессов нами предлагается динамическая матрица векторов с расширяющимися строками и столбцами. Каждый вектор соответствует отдельному процессу и он может обрабатываться на отдельном ядре или процессоре суперкомпьютера.

Предлагаемый нами метод прогнозирования может быть применен для сложно прогнозируемых явлений, которые при прогнозировании используют несколько методов прогнозирования с небольшой вероятностью. Рассматривая параллельные данные этих методов прогнозирования можно получить новый метод прогнозирования с большой вероятностью.

Также, этот метод применения параллельных данных должен обеспечить обнаружения таких закономерностей, которые до сих пор не открыты из-за ограниченных возможностей последовательных алгоритмов.

Определим термин «одновременные данные» указанного процесса. Это данные, которые помещены в каждой строке матрицы. Они могут быть разного типа, но вместе участвуют в процессе прогнозирования. С левой стороны матрицы – временная шкала (напр., 1 мин), а справа стороны – функции $f_1, f_2, \dots, f_n, \dots$. На временной шкале отмечена точка t_0 , момент времени, когда мы начали изучения процесса. Естественно, при желании или надобности дополнительной информации, матрицу можно расширить снизу с добавлением информации, например, из каталогов землетрясений. Справа, представлены функции, характеризующие сейсмический процесс. Например, функции f_1 и f_2 могут быть функциями, значения которых соответствуют электромагнитному излучению Земли VLF и LF для данного момента времени t_i , а f_3 – функция из каталога землетрясений и т.д. В этот список функций могут попасть такие функции, которые как будто не связаны с прогнозируемым прогнозированием, но все-таки имеют влияние, например, политические процессы, погода и т.д.

Литература. 1. Гульельми А.В. Форшоки и афтершоки сильных землетрясений в свете теории катастроф. Успехи физ. наук. 2015. – Т. 185. № 4. – С. 415–429. 2. Шуман В.Н. Сейсмический процесс и современные мониторинговые системы. Геофиз. журн. 2014. Т. 36. № 4. – С. 50–64. 3. М.К. Kachakhidze, N.K. Kachakhidze. Prediction Capabilities of VLF/LF Emission as the Main Precursor of Earthquake. GESJ: Physics 2015 | No.1(13) ISSN 1512-1461. Pp. 80–89.

Бакун С.А., Терентьев О.М.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Методи перетворення категоріальних змінних в числові

Вступ. На сьогоднішній день все більшого масштабу набуває практичне застосування методів знаходження закономірностей у великих обсягах даних та методів прогнозування, зокрема, побудова регресійних моделей. Переважна більшість таких алгоритмів дозволяють працювати лише з числовими ознаками для опису об'єктів, або процесів, що спостерігаються. Однак на практиці часто зустрічаються задачі з категоріальними змінними, значення яких позначають приналежність об'єктів до якоїсь категорії. Приклади таких ознак: національність, громадянство, професія тощо. Нагальною необхідністю є коректне перетворення текстових значень з фіксованого набору, тобто категорій, у числовий формат, для забезпечення можливості роботи з категоріальними змінними.

Метод фіктивних змінних. Найвідомішим та простим методом перетворення категоріальних змінних у числові є метод фіктивних змінних. Суть методу полягає у розбитті кожної категоріальної змінної на множину фіктивних бінарних змінних (табл. 1). При цьому одна з категорій зазвичай не кодується з метою забезпечення функціональної незалежності множини створених змінних. Таким чином, категоріальна змінна перетворюється в набір $k - 1$ бінарних змінних, де k – кількість взаємовиключних категорій даної категоріальної змінної [1].

Табл. 1. Приклад перетворення змінної “Освітньо-кваліфікаційний рівень”

Назва категорії	A1	A2	A3
Немає	0	0	0
Молодший спеціаліст	1	0	0
Бакалавр	0	1	0
Магістр	0	0	1

Недоліками даного методу є те, що метод не враховує ні характеристики розподілу категорій, ні будь-який взаємозв'язок з цільовою змінною. Ще одним недоліком такого підходу є суттєве збільшення кількості змінних при заміщенні кожної категоріальної змінної на множину фіктивних змінних. Більшість з алгоритмів не зможуть обробити отриману кількість даних на реальних задачах.

Метод використання ваг категорій. Альтернативою методу фіктивних змінних є метод на основі використання ваг категорій, або WOE (Weight of evidence). WOE вимірює статистичну значущість кожного класу змінної [2]. Якщо змінна, що спостерігається, є бінарною змінною, тобто приймає тільки два значення (наприклад, 1 — якщо певна ознака присутня; 0 — якщо вона відсутня), то WOE розраховується за формулою:

$$WOE = \ln \left(\frac{d_i^{(1)}}{d_i^{(2)}} \right), \quad (1)$$

де $d_i^{(1)}$ – відносна доля першого спостереження в i -й категорії;

$d_i^{(2)}$ – відносна доля другого спостереження в i -й категорії.

Даний метод дозволяє ставити у відповідність категоріальним значенням числові значення ваг категорій змінної (табл. 2).

Порівняльний аналіз. В рамках порівняльного аналізу методів перетворення категоріальних змінних у числові було побудовано дві регресійні моделі, для прогнозування кредитоспроможності фізичних осіб. Для побудови моделей було використано вибірку даних з німецького банку, що надає кредити фізичним особам. Набір даних містить 3000 записів по клієнтах, у яких вже закінчився строк кредитування, та включає в себе інформацію щодо 17 показників

Табл. 2. Приклад перетворення змінної “Ціль кредиту”

Назва категорії	Значення WOE
Автомобіль	-0.21
Побутова техніка	-0.11
Відпочинок	0
Покупки в магазинах	0.05
Меблі	0.19

(анкетних даних) по кожній особі, з яких 7 є категоріальними:

1. TITLE – стать особи;
2. STATUS – сімейний стан;
3. REGN – регіон;
4. PRODUCT – тип бізнесу;
5. NAT – національність особи;
6. CARDS – тип банківської карти;
7. RESID – тип проживання.

Як можна побачити з таблиць 3 та 4, побудована модель з використанням коефіцієнту WOE для перетворення категоріальних змінних в числові має кращі значення індексу GINI та загальної точності моделі ніж з використанням методу фіктивних змінних на навчальній та тестовій вибірках.

Табл. 3. Порівняльна таблиця характеристик якості моделей на навчальній вибірці

Модель	Загальна точність (CA)	Індекс GINI
Модель з використанням фіктивних змінних	0.85	0.65
Модель з використанням ваг категорій	0.87	0.67

Табл. 4. Порівняльна таблиця характеристик якості моделей на тестовій вибірці

Модель	Загальна точність (CA)	Індекс GINI
Модель з використанням фіктивних змінних	0.82	0.57
Модель з використанням ваг категорій	0.84	0.6

Висновки. Перетворення категоріальних змінних в числові є досить трудомісткою задачею. Застосування методу фіктивних змінних передбачає додавання великою кількості нових змінних у модель, що ускладнює модель і погіршує оцінку її якості. Було запропоновано альтернативний метод перетворення категоріальних змінних – на основі використання ваг категорій. В ході роботи було проведено порівняльний аналіз даних методів. В результаті отримано підтвердження, що метод на основі використання ваг категорій може застосовуватися як метод, альтернативний методу фіктивних змінних.

Література. 1. Hosmer D.W. Applied logistic regression / David W. Hosmer, Jr., Stanley Lemeshow. — 2nd ed. — Hoboken: John Wiley & Sons, Inc., 2000. — 375 p. 2. Siddiqi N. Credit risk scorecards: developing and implementing intelligent credit scoring / Naeem Siddiqi. — Hoboken: John Wiley & Sons, Inc., 2006. — 196 p.

Васильев В.І., Вішталъ Д.М., Любашенко Н.Д.

КПІ ім. Ігоря Сікорського, ФПМ, Київ, Україна

RDF-модель порталу підтримки розміщення і використання даних у Всесвітній мережі

В роботі досліджено сучасний стан проблеми розвитку єдиної семантичної мережі (Semantic Web). Обґрунтовується доцільність створення порталу для підтримки процесу обробки даних у Всесвітній павутині (WWW) за технологією Linked Open Data.

Формулювання проблеми. Семантична павутина (Semantic Web), як концепція розширення Всесвітньої павутини документів (World Wide Web) до Всесвітньої павутини даних, продовжує викликати інтерес у інтернет спільноти. Технології Linked Open Data (LOD, Зв'язані відкриті дані) та Semantic Web просуваються Консорціумом Всесвітньої павутини W3C (World Wide Web Consortium) [1]. Ці технології застосовують стандарти для опису інформаційних веб-ресурсів і зв'язування їх з іншими веб-ресурсами, таким чином, щоб не тільки люди, а й програми змогли знаходити і обробляти дані з багатьох різних джерел. При реалізації цього підходу практична цінність даних, зокрема, можливість отримання нових знань там, де вони навіть не були передбачуваними, суттєво зростає. В даній роботі було проведено аналіз практичних і теоретичних досягнень в області вказаних технологій. Виявилось, що цей інструментарій застосовано в реальних масштабних проектах [4]. За останні роки зріс об'єм семантичних веб-даних і, відповідно, виросла кількість платформ, порталів, інтернет-сервісів і прикладних застосувань на базі семантичних технологій (RDFa, Turtle, RDF Schema, OWL, SPARQL, Dbpedia, Virtuoso) [3,6]. Спільна мета, яка об'єднує ці розробки – створення єдиної глобальної бази даних.

Проблема полягає в тому, що для подальшого розвитку мережі даних залучення нових учасників (окремих людей, бізнес-структури, державні органи, суспільні організації тощо) стикається з певними перепонами, а саме: освоєння нової парадигми роботи з веб-даними вимагає певних ресурсів часу для теоретичного навчання та освоєння спеціалізованого програмного забезпечення. Зокрема, початківець може “загубитися” у веб-просторі, оскільки дана тематика супроводжується великою кількістю документів, які містяться на різноманітних сайтах у різноманітних формах і форматах. Як один з варіантів вирішення цієї проблеми пропонується розробка порталу підтримки процесу навчання та реалізації розміщення даних у Всесвітній павутині і їх подальшого використання.

Портал підтримки роботи з веб-даними. Автори пропонують системний підхід до просування і розвитку ідеї веб-простору даних, а саме, створення порталу підтримки розміщення і використання даних (рис. 1). Процес розміщення і використання даних розглядається як технологічний процес у віртуальному світі, який має підтримуватись як будь-який технологічний процес в реальному світі. Створення порталу забезпечить обслуговування різних категорій користувачів, які зможуть легко знаходити теоретичні дані по технології Linked Data, використати існуючі прикладні застосування, розмістити свої дані для відкритого доступу, розробити програму обробки даних тощо. На даний момент авторами пропонується застосувати типову архітектуру порталу (рис. 1) [2]. Для даного прикладу за технологією LOD портал інтегрує веб-ресурси музичного спрямування, подані в різних форматах, і тим самим забезпечує семантичний зв'язок даних з різних джерел.

RDF-модель порталу. Терміном Linked Data визначено підхід до розміщення даних в мережі, в основі якого є зв'язування даних через посилання, аналогічні гіперпосиланням в мережі веб-документів. Підхід базується на 4-х базових принципах [1,7]. Під реалізацію цих принципів Консорціум W3C затвердив стандартну модель даних RDF (Resource Description Framework) [1]. Модель дозволяє подати інформацію про ресурси будь-якої природи (документ, людину, фізичне явище, матеріальний об'єкт, абстрактне поняття).

RDF модель – це графова модель, яка є множиною “трійок” вигляду
<SUBJECT> <PREDICATE> <OBJECT>.



Рис. 1. Архітектура порталу Linked Data

Предикат вказує відношення між ресурсом-суб'єктом та іншим ресурсом-об'єктом. Унікальні ідентифікатори кожного елемента "трійки" забезпечують однозначну інтерпретацію.

Для проектування порталу пропонується розробити його модель RDF. Частина моделі описує етапи роботи з веб-даними: підготовка даних; розміщення даних в інтернеті, яке названо серіалізацією RDF-моделі; при потребі переведення даних з інших форматів у RDF формат, для чого існує відповідне програмне забезпечення; розробка інтернет-застосувань, які використовують розміщені дані; інше.

Фрагмент моделі для сутностей "Підготовка даних", "Стадія життєвого циклу", "Сайт" показано нижче у двох варіантах: перший надає семантичне навантаження графа, яке зрозуміле людині, а другий містить URIs (Uniform Resource Identifier), які оброблятимуться програмами. Для позначення предикатів використано онтології [5].

- 1) < Підготовка даних за технологією Linked Data > < subclass of > < Stage-of-Life > ;
< Підготовка даних за технологією Linked Data > < is Shown At > < Site > .
- 2) < <http://www.lodportal.com/prep> > < <http://www.w3.org/2000/01/rdf-schemasubClassOf> >
< <http://contextus.net/ontology/ontomedia/ext/common/traitStage-Of-Life> > ;
< <http://www.lodportal.com/prep> > < <http://www.europeana.eu/schemas/edm/isShownAt> >
< <http://https://www.w3.org/TR/2017/REC-dwbp-20170131/> > .

Перевагою запропонованої моделі порталу є наступне. Метадані про потрібні для нашої проблеми ресурси будуть подані в Інтернеті у форматі Linked Data. Отже, їх виявлення спеціалізованими програмами є гарантованим. Портал також забезпечить знаходження всіх доступних на момент запиту даних і доступ до них, якщо, очевидно, вони були розміщені за стандартами LOD. Отже, користувач може бути впевнений, що вся необхідна для роботи інформація знаходиться у нього під рукою. Крім того, поява нових пошукових результатів вплине на автоматизований розвиток самого порталу.

Отже, за функціональним наповненням і орієнтацією на будь-якого користувача портал можна використати як інформаційно-довідкову систему, як навчальну систему для підготовки спеціалістів в області Semantic Web, як технологічну підтримку для реалізації проектів розміщення та використання відкритих даних. Хоча йдеться про відкриті дані, портал буде корисним і для організацій, яким треба здійснювати інтеграцію даних з обмеженим доступом.

Література. 1. World Wide Web Consortium, <https://www.w3.org/>. 2. The Euclid learning materials, <http://euclid-project.eu>. 3. Bob DuCharme Learning SPARQL. Querying and Updating with SPARQL 1.1 / Bob DuCharme. — O'Reilly Media, 2013. — 386 p. 4. Data on the Web Best Practices. W3C Recommendation 31 January 2017. 5. <http://dublincore.org/metadata-basics/>. 6. The home of the U.S. Government's open data, <https://www.data.gov/>. 7. Tom Heath Linked Data: Evolving the Web into a Global Data Space. Synthesis Lectures on the Semantic Web: Theory and Technology / Tom Heath, Christian Bizer. — Morgan @ Claypool, 2011. — 136 p.

Видолоб А.В.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Забезпечення інформаційної безпеки системи інтернету речей

У інтернеті речей (Internet of Things, IoT) крім порушення конфіденційності традиційних мереж зв'язку (повторювання, підслуховування, викривлення інформації тощо.) [1], виникають проблеми з захистом споживчої складової. Вони обумовлені:

- відсутністю стандартів не тільки захисту, але і взаємодії;
- відсутністю в наші дні інтересу у виробників, як у першого ступеня реалізації [2].

Проблема з пристроями IoT полягає в тому, що вони мають незначну безпеку та практично незахищені. Вони, як правило, не запускають стандартні операційні системи, які підтримують зазвичай використовувані інструменти захисту інформаційних технологій або просто не мають достатньо пам'яті для них. Багато пристроїв також не має можливості застосовувати оновлення прошивки, що робить неможливим виправлення вразливостей у безпеці, коли вони виявляються.

Існуючі засоби захисту підприємств, які контролюють стан пристрою, не працюють на пристроях на базі IoT в межах мережі, адже вони не мають видимості на цих пристроях.

Переважає більшість нових схем безпеки IoT не працюють на захист пристроїв, вже встановлених у мережах. Комбінація постійно з'єднаних і не захищених пристроїв IoT – ідеальний набір можливостей для кібератакерів.

У попередні роки зібрані бот-мережі для DDoS атак в основному складали персональні комп'ютери. Тепер ці бот-мережі збираються за допомогою пристроїв IoT. IoT пристрої також використовуються для встановлення “бекдор” для кібератак, які не мають нічого спільного з DDoS. Для стандартних мережевих систем безпеки неможливо ідентифікувати ці бекдори, перш ніж вони можуть бути використані.

Однією з найпотужніших стратегій захисту IoT систем є сегментація мережі. Мікросегментація розміщується вбудованою у весь анклавний трафік у мережах. Це суттєво зменшує рух на схід-захід (бічний рух) в мережах.

Мікросегментація дозволяє призначати політику для всіх пристроїв і користувачів у мережі, які визначають, кому з ким дозволено говорити та до яких ресурсів допускається доступ. Ці політики ґрунтуються на користувачі, порту та IP-адресі, а також забезпечують автоматизацію, яка дозволяє використовувати їх навіть у найбільшому підприємстві. Мікро-сегментація зупиняє кібератакерів від використання пристроїв IoT як “рабів” у ботнеті або як бекдор для інших форм витончених атак.

Мікросегментацію можна також поєднувати з новими технологіями, такими як переміщення цільового кіберзахисту (Moving Target Cyberdefense, MTD), що може ще більше зменшити доступну поверхню нападу. Хоча сегментація зменшує поверхню атаки, зловмисник все ще здатний виконувати розвідку та викрадати інформацію про мережу, необхідну для цілей планування атак. MTD зупиняє всю розвідку противника від скомпроментованих кінцевих точок або інсайдерських загроз. Якщо вони не бачать це, вони не можуть знайти цього. IP-адреси для пристроїв IoT затуманені і невидимі нападникам. Це метод придатний для захисту всіх незахищених застарілих пристроїв IoT, які вже встановлені у мережі.

Як сегментація, так і MTD можна додати до поточної стратегії захисту, що дозволить суттєво знизити ризик того, що пристрої IoT будуть скомпрометовані. Ці технології дозволяють захистити поточну встановлену базу пристроїв IoT та розмістити практично будь-які пристрої IoT, які потрібно буде додати до мережі у майбутньому.

Література. 1. Соколов М.Н. Проблемы безопасности интернет вещей /М.Н. Соколов, К.А. Смолянинова, Н.А. Якушева // журнал “Вопросы кибербезопасности”. — 2015. — № 5(13). — С. 34. 2. Алексей Лукацкий Криптография в “Интернете вещей” [Електронний ресурс]. – Режим доступу: <http://www.slideshare.net>.

Виклюк Я.І., Гусак О.М.

ПВНЗ “Буковинський університет”, Чернівці, Україна

Шляхи підвищення ефективності протипожежного моніторингу лісу

Запропонована інформаційна технологія розширення функціональних можливостей неспеціалізованих безпілотників для виявлення лісових пожеж на ранніх стадіях запалювання, що дасть змогу підвищити ефективність протипожежного моніторингу лісу без встановлення додаткового обладнання і не вимагаючи від держави жодних витрат.

Створення нових та вдосконалення існуючих засобів розвідки осередків лісових пожеж допоможе подолати протиріччя між високим рівнем витрат на розробку високотехнологічних протипожежних засобів та їх недостатньою ефективністю. В дослідженні пропонується залучити до протипожежного моніторингу лісу безпілотники індивідуальних користувачів, активний розвиток яких ми спостерігаємо сьогодні [1]. Ідея залучення та можливого розширення інформаційно-технологічних можливостей безпілотників базується на аналізі їх інформаційно-технологічних характеристик та сфер застосування [2]. Активний розвиток безпілотників для вирішення задач моніторингу зумовлений рядом їх значних переваг над традиційними методами моніторингу:

- можливістю ведення спостереження за відсутності екіпажу на борту;
- достатньою тривалістю і дальністю польоту;
- маневреністю;
- спроможністю вести моніторинг у обмежених метеоумовах;
- відносно невеликою вартістю та малими витратами на експлуатацію;
- можливістю масового виробництва.

Пропонована інформаційна технологія дозволяє без встановлення додаткового високотехнологічного обладнання, лише за рахунок розширення наявних інформаційно-технологічних можливостей стандартних неспеціалізованих апаратів, власниками яких є індивідуальні користувачі, підвищити ефективність протипожежного моніторингу лісу.

Процес залучення безпілотників до системи моніторингу лісових масивів може бути реалізований наступним чином. Власнику безпілотника пропонується встановити інформаційну систему і взяти участь у спеціальній програмі співробітництва, у рамках якої безпілотники паралельно до своїх вузьких завдань виконують додаткову функцію – оповіщення про небезпеку займання лісу. Користувач безпілотника запускає його над лісом і по wi-fi здійснює керування за допомогою стандартного засобу телеметрії. Безпілотник в фоновому режимі сканує лісову поверхню і періодично надсилає інформацію про свої координати в ДСНСУ, де формується список усіх активних на даний момент безпілотників. У випадку виявлення осередку лісової пожежі, безпілотник засобами 3G зв'язку надсилає в ДСНСУ сигнал тривоги, що складається з оригінального цифрового зображення підозрілої ділянки, її обробленого зображення та її координат. Одночасно ДСНСУ передає сигнал тривоги та список активних в даному районі лісу безпілотників до місцевого оперативно-координаційного центру територіального органу ДСНСУ для перевірки. У випадку підтвердження сигналу про небезпеку, із центру на пристрій керування активним в даному місці безпілотникам надсилається прохання змінити курс для уточнення інформації про небезпеку.

Створення такого мобільного інформаційно-технологічного сервісу на основі добровільної соціо-комунікаційної ініціативи, доповнить існуючу систему протипожежного моніторингу лісу і допоможе підвищити її ефективність, оскільки дасть можливість виявляти лісові пожежі на ранніх стадіях запалювання і одержувати інформацію про них раніше ніж з офіційних джерел, при цьому не вимагаючи від держави жодних витрат.

Література. 1. Planeta hobby [Електронний ресурс]. – 2017. – Режим доступу до ресурсу: <https://modelistam.com.ua/radioupravlyaemye-modeli/>. 2. Світ дронів [Електронний ресурс]. – 2017. – Режим доступу до ресурсу: https://theoryandpractice.ru/posts/7834-peace_drone/.

Гіоргізова-Гай В.Ш., Шеренковський А.О.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Шлюзи в системах IoT

Вступ. Ринок Інтернету речей розглядається аналітиками як найбільш перспективний в найближчому десятилітті. Прогнозований загальний економічний ефект IoT вимірюється трильйонами доларів США [1]. Сьогодні велика кількість виробників пропонує компоненти для IoT рішень, серед яких важливе місце займають шлюзи. При чому, функції шлюзів і їх призначення у різних IoT проектах досить сильно відрізняються [1].

Постановка задачі. Загалом шлюзи IoT представляють собою мережеве обладнання, яке знаходиться на кордоні між ОТ (Operational Technology) – апаратно-програмними комплексами для контролю і управління фізичними процесами, та ІТ (Information Technology) – системами і мережами для створення, обробки, зберігання, забезпечення безпеки і обміну будь-якими формами електронних даних. Але на відміну від більш звичних для ОТ та ІТ компонентів (датчики, контролери, модеми, маршрутизатори, протоколи каналів зв'язку, IoT-платформи) аналітичні огляди для шлюзів практично відсутні. Так само, відсутні і критерії їх класифікації та порівняння, що суттєво ускладнює їх вибір при проектуванні IoT-систем.

Аналіз стандартів архітектури IoT. Зараз стандартизацією займається багато організацій. Серед найбільш впливових можна виділити:

- Міжнародний союз електрозв'язку (International Telecommunication Union, ITU): 193 країни і понад 700 членів по секторам і асоціаціям (науково-промислових підприємств, державних і приватних операторів зв'язку, радіомовних компаній, регіональних і міжнародних організацій) [4].
- Всесвітній форум IoT (IoT World Forum, IWF): IBM, Intel, Cisco, Samsung [5].
- Національний інститут стандартів і технологій міністерства торгівлі США [3].
- Консорціум індустріального Інтернету (Industrial Internet Consortium, IIC): SAP, IBM, Intel, Fujitsu, General Electric, Oracle [2].

Всі з перелічених організацій віддають шлюзу функцію створення локальних мереж для підключення неінтелектуальних “речей” і узгодження протоколів при взаємодії між ОТ та ІТ. Мережа може мати як стандартний TCP/IP стек протоколів і бути реалізована за допомогою Wi-Fi мережі або дротового підключення, так і бути не IP мережею та підключати датчики за допомогою таких технологій як Bluetooth, ZigBee або 6LoWPAN та ін. Також всі вони визначають функції безпеки, проте вони можуть відрізнятися, відтак у моделі ITU функції безпеки обмежені лише на рівні пристрою, у інших же моделях він забезпечує ще й безпеку на рівні мережі. Всі організації окрім ITU до ОТ відносять ще й функції аналітики даних від речей. У IWF шлюз класифікують, як мережевий пристрій, проте в архітектурі на суміжному рівні виділяються функції перефрмис/туманних обчислень. Cisco, котрий бере активну участь у IWF, явно позиціонує свої шлюзи, як пристрої, що проводять аналітику даних та реагують на події. ІІС для шлюзів виділяє функції аналітики краю, що по факту є додаванням рівню туманних обчислень (з моделі IWF) до можливостей шлюзу. В моделі національного інституту стандартів шлюз одразу називається агрегатором і являється частиною апаратного забезпечення, що займається дослідженням даних від датчиків. Він являється певним сервером, що знаходиться у безпосередній близькості до речей. Відтак шлюз перестає бути лише конвертором з одних протоколів у інші, він набуває здібностей аналізувати дані, оброблювати їх і зберігати/надсилати у більш стислій формі. Можливе реагування на показники певних датчиків або певні події у режимі реального часу, а не у часі транзакції, як при безпосередньому підключенні речей до хмари.

Критерії класифікації і порівняння. Оскільки в IoT одним з обов'язкових компонентів є хмарний сервер, а для його повноцінного розгортання, як зараз практично всі визнали, необхідна платформа, то можна виділити наступні найбільш поширені варіанти надання

виробниками своїх шлюзів:

- Виробник надає шлюзи (лінійки шлюзів) універсального призначення, платформу можуть надавати незалежні виробники (Intel, Panasonic Corp. of North America, Samsara).
- Виробник шлюзів також є і розробником платформи (Digital Six Laboratories, Eurotech).
- Виробник пристроїв і шлюзів постачає їх відповідними додатками, платформу надає інший вендор (Dell Technologies, Panasonic).

Ринок послуг та пристроїв IoT зазвичай поділяють на два великі сегменти: *промисловий* і *споживчий* сегмент, які відрізняються вимогами до ціни, надійності, безпеки, потужності, масштабованості. У цих рамках вендори можуть надавати як універсальні пристрої так і пристрої для конкретних вертикалей ринку в складі комплексних рішень (Cisco, Eurotech). Також важливо розуміти для яких цілей необхідний шлюз. Відтак ряд виробників пропонують пристрої із підтримкою перефінансування/туманних обчислень на базі власної платформи (Cisco, Davta Networks, Dell Technologies, Hewlett Packard Enterprise). Інші ж, просто постачають шлюзи разом із засобами для створення власних додатків первинної обробки даних, до якої зазвичай відносять діагностику стану об'єктів, обробку подій в режимі реального часу (Intel, Dell Technologies, Panasonic). В цьому випадку у якості критеріїв вибору слід звернути увагу на наступне:

- Надійна ОС (власна або стороння). ОС реального часу.
- Надання в комплекті зі шлюзом ПЗ для розробки IoT-додатків або готових додатків. ПЗ для розробки додатків має володіти ефективними подійними можливостями (тригери подій, обробники подій і т.д.), інтерфейсами прикладного програмування (API), комплектами розробки додатків для прототипування, налагодження, тестування і т.д.

У випадку, коли необхідно обрати шлюз, котрий буде займатися лише пересиланням даних від речей до хмари слід звернути увагу на наступні речі:

- Максимальна кількість пристроїв, з якими може взаємодіяти шлюз.
- Можливості збору даних з різних джерел, їх інтеграції, уніфікації представлення (протоколів і форматів даних) для взаємодії з зовнішнім сервером і пристроїв між собою.
- Перелік внутрішніх (з пристроями) і зовнішніх (з сервером) інтерфейсів.

У рамках подібної функціональності шлюзи можуть відрізнятися технічними характеристиками: обчислювальною потужністю, об'ємом пам'яті, виконанням (одні призначені лише для роботи у приміщеннях, інші ж здатні витримувати жорсткі зовнішні умови) та ін.

Висновки. При розробці IoT проекту потрібно на основі стандартних моделей створити модель архітектури власного проекту, що дасть можливість специфікувати основні компоненти системи, їх взаємодію і взаємозв'язок, визначити технічні вимоги до цих компонентів. Запропоновані критерії класифікації і порівняння допоможуть вибрати шлюзи необхідної функціональності і технічних характеристик за пропозиціями різних постачальників.

Література. 1. The Internet of Things: Mapping the Value Beyond the Hype. (2015, Червень). Офіційний сайт McKinsey Global Institute. Переверено 30.03.2018 за посиланням http://www.mckinsey.com/insights/business_technology/the_internet_of_things_the_value_of_digitizing_the_physical_world. 2. The Industrial Internet of Things. (2017, 31 січня). Офіційний сайт консорціуму індустріального Інтернету. Переверено 04.04.2018 за посиланням http://www.iiconsortium.org/II_C_PUB_G1_V1.80_2017-01-31.pdf. 3. Networks of 'Things'. (2016, Липень). Офіційний сайт національного інституту стандартів і технології міністерства торгівлі США. Переверено 04.04.2018 за посиланням <https://nvlpubs.nist.gov/nistpubs/SpecialPublications/NIST.SP.800-183.pdf>. 4. ITU: membership. (2018). Офіційний сайт міжнародного союзу електров'язку. Переверено 04.04.2018 за посиланням <https://www.itu.int/online/mm/scripts/mm.list>. 5. Transformation to a Next Generation IoT. (2015, 21 грудня). Офіційний сайт міжнародного форуму IoT. Переверено 04.04.2018 за посиланням <https://www.iotwf.com/resources/108>.

Гнатенко В.Ю., Ситников В.С., Ступень П.В.

Одесский национальный политехнический университет, Одесса, Украина

Электрическая модель с идеальными элементами для поиска максимального потока по транспортной сети

Известны методы решения задач линейного программирования с использованием электрических моделей постоянного тока [1, 2]. В [3] предложен метод расчета одной из задач линейного программирования – поиска кратчайшего пути во взвешенном ориентированном графе с применением электрической модели с идеальными электрическими элементами. Идею этого метода, позволяющую модифицировать аналоговую электрическую модель в цифровую, можно применить для решения задачи поиска максимального потока по транспортной сети.

В данных тезисах рассматривается развитие предложенной в [3] модели представления взвешенного ориентированного графа. Предлагается формировать и обрабатывать в процессе анализа список ветвей с присущими им характеристиками и параметрами [4].

Каждой ветви графа транспортной сети ставится в соответствие ветвь электрической схемы. Каждая ветвь схемы содержит диод, для регулирования потока через ветвь и источник тока J_b . Кроме того, обязательно должна быть еще дополнительная ветвь с диодом. Подключение этой дополнительной ветви между выбранными двумя узлами электрической схемы эквивалентно выбору истока и стока транспортной сети. Ветви в списке могут располагаться в любом порядке.

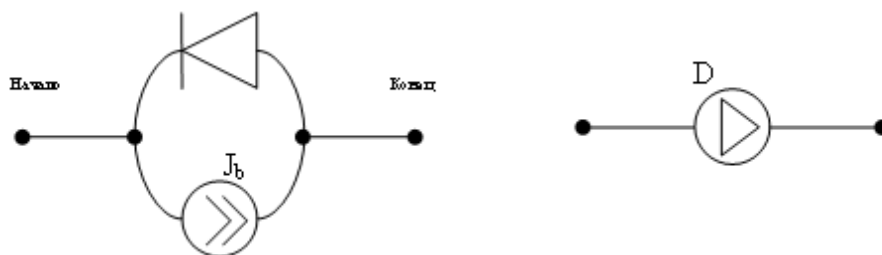


Рис. 1. Электрическая схема ветви и ее условно-графическое обозначение.

Ток источника тока J_b эквивалентен весу ветви транспортной сети. Каждый элемент списка ветвей графа содержит четыре поля:

- номер узла, с которым связано начало ветви;
- номер узла, с которым связан конец ветви;
- величина тока источника тока ветви;
- величина обратного тока ветви, протекающего через диод.

С учетом упомянутого состава ветвей электрической схемы уравнение каждой ветви описывается выражением:

$$I_i = J_{ib} - (\phi_{Ki} - \phi_{Ni})G_{bi} \quad (1)$$

$$G_{bi} = \begin{cases} 0, & \text{если } \phi_{Ki} - \phi_{Ni} \leq 0, \\ 1, & \text{если } 0 < \phi_{Ki} - \phi_{Ni} \leq U_{ib} \\ \frac{J_{ib}}{\phi_{Ki} - \phi_{Ni}}, & \text{если } \phi_{Ki} - \phi_{Ni} > U_{ib}, \end{cases}$$

где i – номер ветви, J_{ib} – ток источника тока ветви, U_{ib} – напряжение насыщения диода ветви, I_i – ток ветви, $\phi_{Ki} - \phi_{Ni}$ – разность потенциалов ветви, G_{bi} – проводимость диода i -й ветви.

Так как ветви с идеальными диодами не имеют сопротивления определенной величины, электрическая цепь не содержит элементов с конечной величиной сопротивлений, для расчета ее токов и напряжений неприменимы традиционные методы теории электрических цепей. Предлагается использовать метод установления, который заключается в анализе токов и напряжений во время протекания переходного процесса до статического состояния [5]. Так

как исходная электрическая цепь не содержит накопителей энергии и в ней в принципе невозможен переходный процесс, предлагается преобразовать исходную электрическую цепь в динамическую добавлением вышеупомянутых накопителей путем присоединения к каждому узлу схемы по емкости, другой конец каждой из которых соединить с базисным узлом, не принадлежащим данной схеме и общим для всех емкостей. В результате становится возможным протекание переходного процесса, по окончании которого токи емкостей станут равными нулю и не будут оказывать влияния на распределение токов и напряжений в схеме.

Следует отметить, что характер изменения напряжений во времени строго линейный, поэтому достижение установившегося режима можно гарантировать за конечный промежуток времени. Для схемы с добавочными емкостями математическая модель схемы будет представлять собой систему нелинейных дифференциальных уравнений, нелинейность которых обусловлена односторонним протеканием тока в идеальном диоде.

При использовании метода установления в математическая модель схемы должна учитывать появление дополнительных емкостных токов. Так как характер изменения во времени искомых узловых потенциалов линейный, что объясняется зарядом или перезарядом идеальных емкостей от идеальных источников тока, то достаточно воспользоваться простейшей формулой численного интегрирования Эйлера для емкости:

$$\phi_{rk+1} = \phi_{rk} + \frac{h}{C_r} \Delta I_{rk}, \quad (2)$$

где ϕ_r – потенциал узла r или напряжение емкости, k – номер узла, с которым связано начало ветви C_r , h – шаг численного интегрирования, ΔI_r – невязка тока в узле r .

Теперь представим алгоритмическое описание поиска максимального потока по транспортной сети.

1. Составить список узлов и инициализировать поля его элементов. Для простоты в поле узла, предназначенное для хранения значений узлового потенциала, устанавливается нулевое начальное значение.
2. Составить список ветвей и заполнить поля его элементов.
3. Рассчитываются напряжения ветвей. Определяются токи всех ветвей системы, за исключением дополнительной ветви, по формуле (1).
4. Вычисляются токи невязки каждого узла путем алгебраического суммирования токов ветвей, связанных с рассматриваемым узлом. Если узел является началом ветви, то ток записывается со знаком плюс, в противном случае – со знаком минус.
5. Прежние значения узловых потенциалов замещаются вычисленными по формуле (2).
6. Если ток невязки каждого узла равен нулю, то кратчайший путь определен, в противном случае перейти к п. 3.
7. Окончание расчета. Максимальный поток по транспортной сети определяется по набору ветвей, в которых ток не равен нулю.

Работоспособность модели была подтверждена на различных графах, в частности на графе, включающем 16 узлов и 31 ребро.

Литература. 1. Хмельник, С.И. Электрические цепи постоянного тока для моделирования и управления. Алгоритмы и аппаратура / С.И. Хмельник: 2-я ред. – Израиль; Россия, 2006. – 177с. 2. Деннис, Дж.Б. Математическое программирование и электрические цепи / Дж.Б. Деннис: перев. с англ. – М.: Изд-во иностр. лит., 1961. – 214с. 3. Шутеев, Э.И. Моделирование нелинейных электрических цепей постоянного тока для решения задачи поиска кратчайшего пути / Э.И. Шутеев, Д.О. Белокопытов, Д.Ф. Димитров // Тр. Одес. политехн. ун-та. – 2009. – Вып 2(32). – с. 88–91. 4. Основы теории цепей: Учебник для вузов.: / Г.В. Зевеке, П.А. Ионкин, А.В. Нетушил, С.В. Страхов. / Под ред. Г.В. Зевеке. – М.: Энергоатомиздат, 1989. – 530 с. 5. Автоматизация схемотехнического проектирования: Учебное пособие для вузов / В.Н. Ильин, В.Т. Фролкин, А.И. Бутко и др. / Под ред. В.Н. Ильина. – М.: Радио и связь, 1987. – 368 с.

Городько Н.А.¹, Бояринова Ю.Е.^{1,2}

¹Институт проблем регистрации информации НАН Украины, Киев, Украина; ²КПИ им. Игоря Сикорского, Киев, Украина

Адаптивные системы организационного обеспечения критических инфраструктур

Система организационного управления объектами критических инфраструктур (СОУ ОКИ) – сложная информационная система совокупности связанных между собой элементов, связи между которыми строятся по определенным правилам [1]. Эти элементы и их взаимодействие можно охарактеризовать с помощью сложных сетей, которые в свою очередь являются объектом как теоретических, так и эмпирических исследований, в которых топология сетей играет ведущую роль. Главной предпосылкой анализа сети является то, что структура связей между узлами влияет на результаты деятельности каждого узла и СОУ в целом.

Интенсивное развитие теории сложных систем предусматривает переход от анализа и оценки надежности к анализу и оценке живучести системы. Живучесть информационной системы – это способность системы хранить и возобновлять выполнение основных функций в заданном объеме и в течение заданного времени в случае изменения структуры системы и/или алгоритмов и условий ее функционирования в результате деструктивных влияний (ДВ) [2].

Задание обеспечения адаптивности системы к ДВ можно рассматривать как одно из актуальных научных заданий повышения живучести в СОУ ОКИ. Для решения этого задания предлагается использовать нейронные сети – системы, которые состоят из многих простых вычислительных элементов (нейронов), функция которых определяется структурой сети, силой взаимосвязанных связей, а вычисления проводятся в самих элементах или узлах.

Работу искусственного нейрона можно описать таким образом: нейрон получает входные информационные потоки через множество входных каналов; каждый входной поток проходит через соединение, которое имеет определенную интенсивность (вес); вычисляется величина активации нейрона, которая равняется разнице между взвешенной суммой входов и пороговым значением нейрона; с помощью функции активации (передаточной функции) информационный поток активации превращается в исходящий поток нейрона (прогнозы или управляющие информационные потоки).

Обучение сети – это процесс минимизации ошибок прогнозов выдаваемых сетью, за счет адаптации модели системы к ДВ. Алгоритмы обучения нейронной сети позволяют по некоторому множеству входных потоков формировать необходимое множество исходящих потоков. Каждое множество информационных потоков при этом рассматривается как вектор. Обучение выполняется путем последовательной подачи входных векторов с одновременным редактированием весов и порогов. При этом в зависимости от выбранного алгоритма значения весов становятся такими, что каждый входящий вектор порождает необходимый исходящий вектор [3].

Благодаря тому, что в нейронно-сетевых моделях реализуется процесс автоматического обучения, их использование при моделировании механизмов повышения живучести позволит не только отслеживать и прогнозировать разные по тяжести состояния системы, но и разработать механизмы адаптации алгоритмов функционирования СОУ ОКИ в условиях постоянно действующих деструктивных влияний.

Литература. 1. 1. Додонов О.Г., Путятін В.Г., Куценко С.А., Ланде Д.В. Побудова узагальненої структури інформаційної системи організаційного управління // Математичні машини і системи – 2017. – № 3. – С. 3–22. 2. 2. Додонов А.Г. Живучесть информационных систем / А.Г. Додонов, Д.В. Ландэ.- К.: Наук. думка, 2011. – 256 с. 3. 3. Ландэ Д.В. Интернетика: Навигация в сложных сетях: модели и алгоритмы/ Д.В. Ландэ, А.А. Снарский, И.В. Безсуднов.- М.: Либроком (Editorial URSS), 2009. – 264 с. ISBN 978-5-397-00497-8.

Грушенко К.Ю., Петрашенко А.В.
КПІ ім. Ігоря Сікорського, ФПМ, Київ, Україна

Модифікований алгоритм виявлення руху з камер відеоспостереження

Вступ. Системи відеоспостереження стають все більш популярними в наші дні. Вони встановлюються в багатьох громадських місцях, складах та інших приміщеннях для попередження небезпечних ситуацій. Проте людський фактор не завжди дозволяє вчасно та правильно діяти в таких випадках. Саме тому використання даних, отриманих з системи відеоспостереження, для швидкого виявлення візуально визначених факторів певних небезпечних ситуацій завжди буде актуальним питанням.

Існує декілька підходів для виявлення руху в безперервному потоці відео, заснованих на порівнянні двох кадрів. Порівняння поточного кадру з фоном (еталонним кадром) має багато недоліків: виявлення нових статичних об'єктів як рухомих, чутливість до зміни освітленості та неважливих рухів фону. Порівняння поточного кадру з попереднім кадром вирішує недоліки попереднього методу, але не враховує незначні зміни за невеликий проміжок часу між кадрами, що можуть мати важливий характер [1].

Модифікований алгоритм виявлення руху з камер відеоспостереження. Для вирішення такого недоліку, запропоновано модифікувати алгоритм за допомогою додавання часового контексту. Часовий контекст – це збережені дані за певний проміжок часу, що використовуються для аналізу змін між двома кадрами. Блок-схему модифікованого алгоритму наведено на рисунку 1.

Запропонований алгоритм передбачає наступні кроки:

1. Порівнюємо кадри у поточний та попередній момент часу методом на основі режиму змішування, що полягає у виконанні операції попіксельного віднімання середнього арифметичного значення кольорних складових зображень. Враховуючи порогове значення, отримуємо чорно-білу маску руху, де білі пікселі – зміни в зображенні.
2. Зберігаємо отриману різницю порівняння кадрів в масиві даних часового контексту.
3. Якщо отримані на кроці 1 зміни перевищують порогове значення змін ПЗЗ, то нівелюємо ділянки з фіктивними змінами (відблисками, шумами) за допомогою даних часового контексту: для кожного пікселю зображення вираховуємо частоту появи змін в часовому контексті та прибираємо з маски ті, частота появи яких більша за порогове значення частот.
4. В результаті отримуємо фінальну маску руху. Якщо порогове значення змін ПЗЗ все одно буде перевищено, то фіксуємо наявність руху.



Рис. 1. Блок-схема алгоритму виявлення руху

Висновки. Запропонований алгоритм дозволяє нівелювати незначні зміни освітленості приміщення, але виявити малопомітні зміни, що мають важливе значення.

Література. 1. Nishu Singla. (2014). Motion Detection Based on Frame Difference Method. International Journal of Information & Computation Technology, 15, 1559–1565. https://www.ripublication.com/irph/ijict_spl/ijictv4n15spl_10.pdf.

Замятін Д.С., Книш П.К.

КПІ ім. Ігоря Сікорського, ФПМ, Київ, Україна

Визначення місцеположення джерела акустичного сигналу

Вступ. Задача локалізації джерела сигналу є базовою в багатьох спеціалізованих системах. До неї зводиться чимало інших задач, наприклад, при побудові охоронних систем, при визначенні місцеположення абоненту мобільного зв'язку та у військових цілях. Актуальність дослідження розподілених сенсорних систем є надзвичайно високою, оскільки вони повсякчас використовуються в перерахованих системах, завдяки низькій вартості сенсорів [1].

Постановка задачі. Задача полягає в оцінці максимальної кількості мікрофонів у моделі, які дозволяють локалізувати джерело акустичного сигналу в сенсорній мережі з заданою точністю. Головною частиною системи, що моделюється, є сервер на який в реальному часі надсилаються дані з мікрофонів за допомогою протоколу UDP та відбувається локалізація за допомогою методу TDOA. Необхідно визначити можливості такого серверу з перспективи максимальної кількості мікрофонів, дані з яких може обробляти сервер, що дозволить підвищити точність локалізації з більшою їх кількістю.

Опис методу. Time difference of arrival (TDOA) є методом, в якому час приходу конкретного сигналу, при фізично окремих приймальних станціях з точно синхронізованими тимчасовими посланнями, дозволяє визначити місцеположення джерела сигналу. У військовому контексті, вона є частиною як радіоелектронної боротьби та спостереження на полі бою [2]. За рахунок використання різниці в часі приходу звуку до мікрофонів, TDOA дозволяє визначити місцеположення джерела звуку.

Дослідження роботи моделі. Для дослідження роботи моделі було використано два комп'ютери, що мають наступні конфігурації: процесори Intel Core I7-4500U, Intel Core I7-7700HQ та 8 і 16 Гб ОЗП відповідно. Результати можна побачити в табл. 1.

Табл. 1. Час локалізації на різних комп'ютерах

К-ть мікрофонів/Час с.	I7-4500U/8 Gb	I7-7700HQ/16 Gb
8	64.17299985	29.26875996
16	141.5010001	58.51470804
32	281.463	123.0097718
64	558.3629999	247.4244539
128	1081.65	485.4320499
160	1345.707	591.5806970

Час виконання зменшився пропорційно до збільшення частоти процесору — приблизно 2,5 рази при такому же збільшенні частоти. Тестування відбувалося на однакових наборах даних.

Висновки. Створено модель, що дозволяє локалізувати джерело звуку за допомогою мікрофонів, кількість яких обмежена об'ємом оперативної пам'яті, використовуючи метод локалізації TDOA. Було проведено експерименти з локалізації джерела звуку на двох комп'ютерах, з використанням різної кількості мікрофонів. Слід зазначити, що збільшення об'єму оперативної пам'яті дозволить значно підвищити кількість мікрофонів, роботу яких можна моделювати. У подальшому було б доцільно зменшити потреби моделі в оперативній пам'яті та модифікувати алгоритм так чином, щоб модель могла працювати з більшою кількістю мікрофонів.

Література. 1. Design of a TDOA location engine and development of a location system based on chirp spread spectrum. Rui-Rong Wang, Xiao-Qing Yu, Shu-Wang Zheng, and Yang Ye. 2016. [Електронний ресурс]. – Режим доступу: <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC5108751/>. 2. Geometric Location Based on TDOA for Wireless Sensor Networks. Xin-long Luo, Wei Li, and Jia-ru Lin. 2012. [Електронний ресурс]. – Режим доступу: <https://bit.ly/2uugusB>.

Замятін Д.С., Перевертайло А.І.

КПІ ім. Ігоря Сікорського, ФПМ, Київ, Україна

Розподіленні обчислення координат місцезнаходження об'єктів за акустичним сигналом у мережі сенсорів

Вступ. Акустичні сенсорні мережі все частіше використовуються в багатьох прикладних областях людської діяльності. Також продуктивнішою стає апаратна складова обчислювального ресурсу таких мереж. Разом з цим збільшуються вимоги щодо швидкодії та граничної кількості об'єктів, які мають бути одночасно локалізовані у таких мережах, тому питання застосування нових алгоритмів для забезпечення розподілених обчислень на вузлах мережі є актуальним.

Постановка задачі. Задача полягає в оцінці можливості збільшення кількості об'єктів, які можуть бути одночасно локалізовані за акустичним сигналом в розподіленій мережі сенсорів шляхом розподілу обчислень координат їх місцезнаходження.

Опис алгоритму. З метою вирішення поставленої задачі побудовано модель, що складається з мережі сенсорів та серверів, які утворюють кластер, а також розроблено алгоритм, що дозволить використати метод TDOA (Time Difference of Arrival) [1, 2] для локалізації сигналів та розподілити обчислення між серверами кластера реалізуючи модель Apache Spark. Сенсори синхронізовані за часом і генерують безперервний потік даних. Сервери в режимі реального часу приймають дані сенсорів та обчислюють координати джерел сигналів. TDOA (Time difference of arrival) – це метод, що дозволяє знайти місцеположення джерела сигналів за рахунок використання різниці в часі надходження звуку з джерела до сенсорів мережі. Визначення координат джерела зводяться до розв'язку матричного рівняння, де елементами шуканої матриці є координати джерела, а елементами відомих матриць є коефіцієнти, що залежать від координат мікрофонів та часу надходження сигналу на мікрофони [1, 2].

Дослідження роботи моделі.. Для дослідження можливості збільшення кількості об'єктів, які можуть бути одночасно локалізовані за акустичним сигналом проведено запуск моделі з різною кількістю джерел на одному комп'ютері та на кластері з двох комп'ютерів. Результати запуску продемонстровані у табл. 1.

Табл. 1. Час обчислень координат на одному комп'ютері та на кластері з двох комп'ютерів

Кількість джерел	Один комп'ютер, с	Кластер з двох комп'ютерів, с
1000	74	100
3000	125	114
5000	180	120
6000	210	111
7000	320	168
8000	444	181
15000	–	304

Висновки. При більших навантаженнях розрахунок координат швидше виконується на кластері, ніж на одному комп'ютері. Це дає підстави стверджувати, що при збільшенні джерел сигналу до певної кількості або збільшенні кількості сенсорів доцільно організовувати кластер і динамічно масштабувати обчислювальні ресурси з метою збільшення кількості об'єктів, які можуть бути одночасно локалізовані за акустичним сигналом у розподіленій мережі сенсорів.

Література. 1. TDOA Acoustic Localization. Steven Li. July 5, 2011. [Електронний ресурс]. – Режим доступу: https://s3-us-west-1.amazonaws.com/stevenjl-bucket/tdoa_localization.pdf.
2. Solution and performance analysis of geolocation by TDOA. [Електронний ресурс]. – Режим доступу: <https://www.hindawi.com/journals/isrn/2012/710979/>.

Капшук О.О.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Використання біометричних технологій при дистанційному навчанні в мережевих інформаційно-освітніх середовищах

Розглядаються методи та засоби оцінки знань при дистанційному навчанні в мережевих інформаційно-освітніх середовищах (МІОС), проблеми верифікації студентів при роботі в онлайн-режимах, при тестуванні рівня набутих знань (заліки, екзамени), а також питання реалізації систем прокторингу з використанням біометричних технологій ідентифікації.

В роботі проводиться аналіз методів та засобів авторизації студентів і контролю їх знань в системах дистанційного навчання (СДН), які використовуються в Україні. Як показує аналіз, в переважній більшості систем дистанційного навчання (СДН) авторизація користувачів виконується з використанням систем парольного захисту при вході в СДН. Основні проблеми при реалізації СДН пов'язані з надійністю ідентифікації користувачів МІОС.

Враховуючи той факт, що в СДН легальний користувач може бути зацікавлений в несанкціонованому доступі сторонньої особи під своїми обліковими даними з метою проходження процедури авторизації. Таким чином, СДН повністю не захищена від загрози отримання доступу до МІОС сторонньою особою від імені легального користувача і результати оцінки знань можуть не відповідати знанням легального користувача.

В класичних кейсових і мережевих СДН, з метою підвищення рівня довіри до результатів навчання, набутих в процесі навчання, використовують підсумкове тестування на території навчального закладу під наглядом осіб, уповноважених навчальним закладом. Необхідність відвідувати навчальний заклад на заключному етапі значно зменшує привабливість СДН, а в деяких випадках взагалі унеможлиблює здобуття відповідної освіти в дистанційному форматі.

Таким чином, перед розробниками систем дистанційного навчання гостро стоїть проблема достовірного розпізнавання (верифікації) користувачів не тільки при вході в СДН, а також при виконанні всіх завдань, передбачених навчальним планом (лабораторні завдання, тести, заліки та іспити).

Проблеми верифікації користувачів (студентів) і їх поведінки, як правило, вирішуються шляхом використання системи прокторингу, яка представляє собою процедуру дистанційного супроводу поведінки і верифікації користувачів (проктор). В якості проктора може виступати уповноважена людина, яка дистанційно спостерігає за діями студента або може перевірити відеозаписи всіх його дій і оцінити ступінь довіри до отриманих студентом результатів.

Використання біометричних технологій розпізнавання обличчя, голосу і клавіатурного почерку для автоматизації процесів ідентифікації користувачів в СДН дозволяє підвищити ефективність роботи прокторів, зменшити їх кількість, або повністю здійснювати спостереження за студентами під час проходження онлайн-іспитів в автоматичному режимі. У цьому випадку функції прокторинга просто інтегруються в платформу онлайн-навчання, при цьому порядок та звичний сценарій проходження перевірки знань не змінюється [1].

Висновки.

1. Впровадження біометричних технологій ідентифікації користувачів в СДН надає можливість суттєво підвищити якість оцінки рівня підготовки студентів і дозволяє проводити підсумковий контроль в дистанційному режимі з використанням веб-технологій.
2. Комплексне використання технологій розпізнавання обличчя, голосу і клавіатурного почерку користувачів в автоматизованих підсистемах прокторингу СДН може суттєво підвищити надійність ідентифікації користувачів, зменшити ризики шахрайства і підвищити якість оцінки знань здобувачів освіти.

Література. 1. Повышаем доверие к результатам дистанционного обучения. Верификация личности и независимая оценка онлайн-экзаменов. [Електронний ресурс] // – Режим доступу: – <http://proctoredu.ru/>. – Дата доступу 1.04.2018. – Назва з екрану.

Киричек Г.Г., Курай В.І.

Запорізький національний технічний університет, Запоріжжя, Україна

Система розпізнавання об'єктів

В роботі запропоновано використання методу дерева квадрантів, для рекурсивної розбивки і порівняння зображень з метою розпізнавання об'єктів, як можливе рішення, коли компактне представлення зображення та високоефективний доступ до елементів, являються ключовими вимогами до системи.

Вступ. Задачі аналізу зображень, розпізнавання образів та порівняння зображень поширені в наш час. Існує багато інструментів, які дозволяють вести дослідження у цьому напрямку. Найбільш поширеними є: перцептивні хеш-алгоритми [1] і штучна нейронна мережа [2]. Найвідомішою бібліотекою для роботи з зображеннями наразі є OpenCV, яка надає засоби для обробки і аналізу вмісту зображень, у тому числі розпізнавання об'єктів [3]. Хоча наведені рішення це добре відточені алгоритми і бібліотеки для роботи з зображеннями, все ж таки інколи розгортання подібної системи є невиправданим з точки зору швидкодії, задіяння ресурсів обладнання та складності роботи з ним. Тому в даній роботі досліджено та запропоновано свій підхід до вирішення цієї проблеми, шляхом впровадження системи визначення об'єктів з використанням методу дерева квадрантів, для порівняння зображень, як можливе рішення, коли компактне представлення зображення і високоефективний доступ до елементів є ключовими вимогами.

Постановка завдання. До популярних методів сегментації зображення можна віднести метод дерева квадрантів (Q-дерево). Основна ідея квадро-дерева – комбінування однакових або схожих елементів даних і кодування великих однорідних сукупностей даних малою кількістю бітів. Причини вибору цього методу наступні: компактне представлення зображення та високоефективний доступ до елементів, які досягаються за рахунок деревовидної структури даних [4]. При розробці програмного забезпечення, необхідно вирішити наступні задачі: дослідити метод дерева квадрантів для рекурсивного розбиття двовірного простору; провести порівняння результатів, отриманих при розбитті зображень, на наявність відмінностей; на підставі отриманих відмінностей створити третє – результуюче зображення, що зберігає різницю перших двох вхідних. Програмне забезпечення розробляється як консольна утиліта, тому параметри для його роботи, такі як вхідні зображення та шлях до міста збереження результуючого зображення, потрібно вказувати як аргументи командного рядка.

Реалізація системи. Виходячи з основних завдань роботи наведено метод реалізації дерева квадрантів при розбитті зображення. Перший крок розбиття – отримання зображення у вигляді двовірного масиву пікселів, шляхом створення методу, що повертає двовірний масив, елементами якого є масив довжиною 3. Використовуючи доступ до елементів через індекс, отримано значення кольору, у відповідній точці зображення. Змінюючи мінімальний розмір вузла регулюємо якість вихідного зображення, за рахунок обмеження подальшого створення нащадків. Так як, великі області одного кольору представляються у вигляді одного прямокутника, то на їх подальше розбиття зміна мінімального розміру вузла не впливає. Натомість, області зображення, на яких знаходиться декілька кольорів, породжують більшу кількість вузлів. У цих випадках за допомогою зменшення мінімального розміру вузла можна завадити подальшому розбиттю зображення. Для визначення середнього значення кольору створено метод, який приймає на вхід вузол дерева. Далі у вкладеному циклі, через індекси проводиться доступ до пікселів зображення, що знаходяться в області вузла дерева. Проводиться підрахунок трьох складових кольору та кожна складова кольору ділиться на кількість пікселів у заданому вузлі, для знаходження середнього значення. Таким чином метод повертає три складових кольору. В тілі методу, що приймає в якості аргументу прямокутник (вузол) проводиться підрахунок різниці кольорів. Розрахунок проводиться три рази, для кожної складової кольору. Результати підсумовуються. Алгоритм порівняння зображень такий: провести розбиття зображень, задавши значення мінімального розміру вузла; зробити порівняння отриманих вузлів

дерева двох зображень; для знайдених квадрантів, повторювати порівняння, поки досягнемо заданого значення мінімального розміру квадранта; після отримання результатів, створити новий графічний файл, та на основі отриманого дерева квадрантів отримати зображення. В результаті проведення досліджень та визначення основних етапів роботи системи, при порівнянні зображень, розроблено консольну програму. Користувач (або система), при цьому, має можливість передати в програму наступні аргументи: шлях до зображень; мінімальний розмір вузла дерева; мінімальну різницю кольорів та шлях до міста збереження результуючого зображення. Представимо етапи роботи програми та її окремих функцій у вигляді блок-схеми, яка відображає процес розбиття зображення на квадранти (рис. 1).

Результатом розбиття зображення є масив квадрантів вузлів дерева. Кожний вузол містить в собі об'єкт типу прямокутника та значення середнього кольору і відстані кольорів. Маючи набір характеристик для кожного вузла дерева, можна провести порівняння двох зображень через порівняння їх дерев квадрантів [5]. Наукова новизна у розробці методу порівняння зображень на предмет їх відмінностей, шляхом порівняння їх дерев квадрантів, з можливістю визначення, з заданою точністю, ділянок зображень, які відрізняються та отримання частини зображення у вигляді квадрантів, що містять координати та колір частини зображення. Підхід, дозволяє здійснювати зручне маніпулювання даними та зменшує витрати пам'яті для їх представлення, забезпечуючи швидкий доступ до елементів зображення. Практична цінність результатів в тому, що розроблено систему розпізнавання об'єктів, шляхом фіксації зображень місцевості, яка дозволяє провести порівняння двох зображень та вивести результат порівняння у третє результуюче зображення. Дана система є механізмом для обробки і маніпуляції над графічною інформацією для її подальшої інтерпретації людиною, що є корисним у медицині (аналіз рентгенівських знімків), зйомках місцевості (аналіз появи та відсутності об'єктів), проведення аналізу супутникових знімків.

Висновки. В роботі проведено дослідження, отримано та наведено варіант програмної реалізації метода дерева квадрантів для порівняння двовірних зображень. Встановлена залежність між якістю порівняння зображень та такими характеристиками дерева квадрантів як мінімальний розмір вузла і відстань кольорів в середині вузла.

Література. 1. «Выглядит похоже». Как работает перцептивный хеш. – [Електронний ресурс] Режим доступу: habrahabr.ru/company/icl_services/blog/262173/. 2. Лукін В.Є. Аналіз використання технології штучних нейронних мереж в якості нового підходу до обробки сигналів // В.С. Лукін / Телекомунікаційні та інформаційні технології. – №3. – Київ, 2014. – С. 81–88. 3. Петин В. Микрокомпьютеры Raspberry Pi: Практическое руководство // В. Петин. – СПб: Питер, 2015. – 240 с. 4. Hunter G.M. Operations on Images Using Quad Trees // G.M. Hunter, K. Steiglitz / IEEE Transactions on Pattern Analysis and Machine Intelligence, 1979. – vol. PAMI-1 (issue 2). – pp. 145 – 153. 5. Kirichek G. Implementation quadtree method for comparison of images // G. Kirichek, V. Kurai / Proceedings of 14th International Conference on Advanced Trends in Radioelectronics, Telecommunications and Computer Engineering (TCSET), Lviv-Slavske, Ukraine, February 20 – 24, 2018. – p. 49.

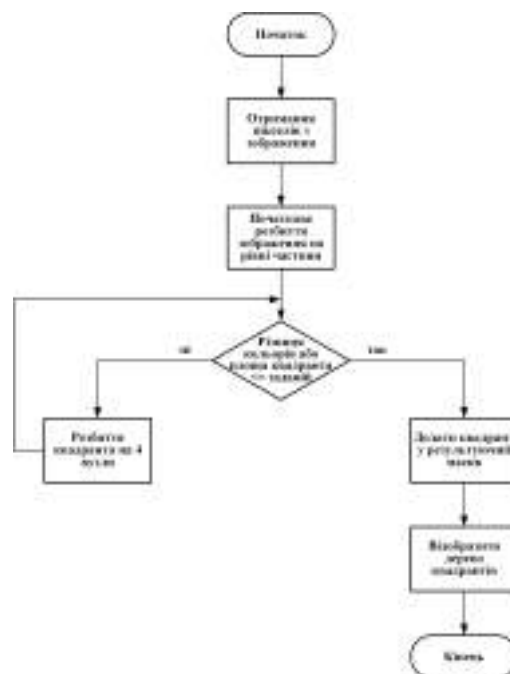


Рис. 1. Блок-схема розбиття зображення

Кінда В.В.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Модель оцінки страхового випадку із використанням розподілу Твіді

Вступ. Страхові компанії використовують багато різних моделей, особливо якщо вони шукають оптимальні рішення проблем ціноутворення і розробки продукту, повний профіль ризиків компанії, кількісну оцінку ризиків, достатність капіталу, адміністрування і управління капіталом в багатьох других областях. Крупні страхові компанії по всьому світу вкладають великі суми в розвиток моделей, орієнтованих на страхові ризики. Є доля ризиків, які на сьогоднішній день дуже складні і майже не страхуються, проте, скоріш за все, вони будуть в майбутньому впливати на ринок страхування та перестрахування. Це ризик глобальних кліматичних змін, енергетичні ризики і таке інше [1]. Значимість моделювання важлива в страхуванні, це пов'язано з його стохастичним характером. Випадковий характер подій і їхня непередбачуваність являються основними атрибутами страхування. Він полягає в виникненні невизначеності конкретної неблагополучної ситуації, а також її часовою невизначеністю. Для кожного суб'єкта проблематично оцінювати, якщо з'являється відповідний ризик і показник збитків, визваний випадковою подією.

Мета досліджень. Метою є дослідження існуючих методів аналізу і прогнозування настання страхових випадків, розробка програмного забезпечення для попередньої обробки даних і побудови моделі для оцінки настання страхового випадку за допомогою GLM (Generalised Linear Model) з використанням розподілу Твіді, та перевірка побудованої моделі на адекватність [2].

Математична основа. Узагальнені лінійні моделі (GLM) — це клас моделей, які являються узагальненням класичних лінійних моделей. В них розподіл елементів Y може бути будь-яким із розподілів експоненційної групи [3]. До них належать біноміальний, пуасонівський, нормальний та гама розподіл і знаходиться у вигляді

$$h(y; \theta, \phi) = \exp\left[\frac{y\theta - b(\theta)}{a(\phi)} + c(y, \phi)\right] \quad (1)$$

для деяких функцій a, b, c і θ , які відомі як канонічні параметри і зв'язані із початковими моментами для конкретного розподілу. Для класичних лінійних моделей функція зв'язку є тотожна, в той час, коли в страхуванні широко використовується експоненційна. Одержимо модель для ризикової премії застрахованого.

$$f(x) = \exp\left(\beta_0 + \sum_{i=1}^p \beta_i x_i\right). \quad (2)$$

Ця модель має властивість завжди бути додатною — $f(x) > 0$. Відповідно ризикова премія ніколи не буде від'ємною. Більше того є можливість комбінувати різні фактори ризику мультиплікативно, що є адекватною комбінацією моделі в страхових задачах. Окрім того, до переваг GLM із експоненційною функцією зв'язку є швидке знаходження параметрів моделі, параметри легко перевіряються на рівень статистичної зачимості і важливість різних незалежних змінних моделі легко аналізується. Також, додавання додаткових змінних в GLM не змінює суттєво час конвергенції для алгоритма оцінки параметрів [3].

Одержані результати. Виконавши препроцесінг даних та розбивши на навчальну та тестову вибірки запустили для порівняння 4 моделі: Multilayer Perceptron, Tweedie, Poisson-Gamma, Intercept — оцінювали результати за значенням середньоквадратичної похибки.

Model	Prediction MSE (Mean Square Error)
MLP	$1.63 * 10^8$
Tweedie	$1.57 * 10^8$
Poisson-Gamma	$1.65 * 10^8$
Intercept	$1.68 * 10^8$

Висновки. Однією з найважливіших проблем страхового бізнесу є встановлення премії для клієнтів. На конкурентному ринку страховикам вигідно стягувати справедливую премію відповідно до очікуваної втрати застрахованого. Необхідність прогнозних моделей впливає з того, що очікувані втрати значною мірою залежать від характеристик індивідуальної політики. Завдяки своїй здатності одночасно моделювати нулі та безперервні позитивні результати, Tweedie GLM є широко використовуваним методом в актуарних дослідженнях. Дана робота показує нам метод використання моделі Твіді в страховому бізнесі. Модель Твіді за середньоквадратичною похибкою дала кращі результати в порівнянні із іншими моделями.

Література. 1. Лемер Ж. Автомобильное страхование. Актуарные модели. / Ж. Лемер – М.: Янус-К, 1998. — 316 с. 2. Anderson, D. et al. “A Practitioner’s Guide to Generalized Linear Models”. In: 2004 Discussion Paper Program – Applying and Evaluating Generalized Linear Models Including Research Papers on the Valuation of PC Insurance Companies. Casualty Actuarial Society, 2004, pp. 1–116. 3. Bent Jørgensen and Marta C. Paes de Souza. Fitting tweedie’s compound poisson model to insurance claims data. Scandinavian Actuarial Journal, pages 69–93, 1994.

Кісільчук Б.Я., Тесленко О.К.

КПІ ім. Ігоря Сікорського, ФПМ, Київ, Україна

Ключі алгоритму шифрування на базі підстановок довільної розрядності

Вступ. В роботах [1, 2] торетично обґрунтована можливість реалізації прямих та обернених підстановок довільної розрядності на базі апаратних або програмних структур лінійної складності від кількості розрядів – одновимірних каскадів конструктивних модулів (ОККМ). Практичний інтерес мають насамперед регулярні структури, які складаються, наприклад, з однакових конструктивних модулів (КМ). При програмній реалізації КМ можна подати у вигляді чотирьох таблиць підстановок $P[h, X, r]$, та чотирьох таблиць вибору підстановок – на молодші розряди – $C_h[h, X]$ і на старші розряди – $C_r[r, X]$, $h, r \in \{0, 1\}$, $X \in \{0, \dots, 2^m - 1\}$. В таблицях підстановок всі значення в таблиці різні, та визначаються випадковим чином із множини $\{0, \dots, 2^m - 1\}$. В таблицях вибору підстановок значення належать множині $\{0, 1\}$. Згідно із [1] таблиці вибору підстановок можна визначити по будь яких таблицях підстановок. Згідно з [2] підстановки КМ для розшифрування є оберненими підстановкам для шифрування. Для шифрування та розшифрування довільних файлів використовується один і той же алгоритм [3], змінюються лише таблиці підстановок та таблиці вибору підстановок. Для виконання процедури шифрування необхідно задати значення h_0 та r_0 для крайніх КМ. Зауважимо, що в алгоритмах шифрування та розшифрування використовуються лише операції пересилання. Теоретично, завдяки наявності таблиць вибору, всі біти результату шифрування залежать від всіх бітів початкового повідомлення.

Виникає наступна проблема – які дані алгоритму віднести до секретних (тобто вважати їх секретним ключем), та визначити попередню оцінку атак на такий ключ.

Аналіз ключів та атак на ключ. Перший варіант. В комп'ютерній інженерії прийнято оперувати байтами, тому приймемо $m = 8$. В такому разі об'єм таблиці підстановок – 2^{11} біт, на всі 4 таблиці – 2^{13} біт. Оскільки таблиці вибору формуються на базі таблиць підстановок, то їх можна не задавати. Необхідно додати ще два біти для початкових значень вибору першого та останнього КМ, але для таблиць підстановок достатньо визначити лише для 255 значень X , 256 – те значення визначається за відсутністю в 255 значеннях таблиці. Цього достатньо, щоб два біти вибору не враховувати. Таким чином ключ може мати розмір в 1 Кбайт.

Розглянемо класифікацію атак на ключ. До першої групи атак віднесемо атаки, коли криптоаналітику відомі лише криптограми (ciphertext-only attack). До другої групи віднесемо атаки, коли криптоаналітику відомі пари – початкове повідомлення та криптограма (known plaintext attack). До третьої групи віднесемо атаки зі спеціально підібраними початковими повідомленнями (chosen plaintext attack), або атаки зі спеціально підібраними криптограмами (chosen ciphertext attack). Такі атаки пов'язані з доступом криптоаналітика до шифруючих (або розшифровуючих) пристроїв, або з можливістю використання (безпосередньо чи дистанційно) сервіси криптографічних перетворень комп'ютерних систем та мереж, в яких таємні ключі задані і не доступні.

Для атак першої групи для блочних криптографічних перетворень існує метод Шеннона по визначенню ненадійності ключів, який ґрунтується на наявності правил створення початкових повідомлень. Метод дозволяє з кожним наступним блоком криптограми відкидати частину ключів. В нашому випадку цей метод не можна застосувати, так як при добавленні кожного наступного байту початкового повідомлення змінюються всі попередні байти криптограми. Зауважимо, що в режимах шифрів Калина та AES, зміна блоку приводить до змін шифrogram лише наступних блоків, а не попередніх.

У випадку атак другої групи криптоаналітик в змозі визначити декілька значень підстановок

ключа, але не підстановки в цілому. А при методі силової атаки необхідно проаналізувати $4^{(w+1)}$ підстановок, де $w = (2^m)!$, які реалізуються всіма можливими каскадами при будь-яких h_0 та r_0 , що вже при $m = 3$ практично не можливо.

У випадку атак третьої групи криптоаналітик в змозі задати всі 16-розрядні початкові дані (або криптограми) та визначити всі чотири прямі (або обернені) підстановки, і тим самим повністю визначити ключ. Отже, алгоритм шифрування, який ґрунтується на одному каскаді можна вважати стійким лише для атак першої та другої групи. Але для запобігання атак із третьої групи достатньо в сервісі (пристрої) криптографічних перетворень аналізувати появу численних вхідних даних, які відрізняються лише в декількох байтах, та блокувати, в цьому випадку, роботу сервісу.

Розглянемо наступний, другий, варіант визначення ключів. Створимо множину $\mathbf{Q}$ КМ, яка не є секретною [3]. Наприклад, при $|\mathbf{Q}|=64$ така множина матиме розмір в 64 Кбайт, і може розглядатись як частина алгоритму криптографічного перетворення. Шифрування (розшифрування) виконується за допомогою підстановки, яка є суперпозицією підстановок із випадково вибраних КМ (операція множення в групі підстановок). Секретними даними (ключем) в даному випадку є послідовний перелік вибраних КМ, включаючи їх кількість – q , та початкові значення h_0 та r_0 для кожного каскаду. Розмір K ключа в бітах обчислюється наступним чином:

$$K = \log_2 |\mathbf{Q}| + 2q + q \log_2 |\mathbf{Q}|.$$

Розглянемо наступний приклад. Для формування множини $\mathbf{Q}$ КМ використаємо 32 загальновідомі 4-х розрядні ($m = 4$) підстановки (S-box) алгоритму DES, властивості яких для протидії лінійному та диференціальному криптоаналізу в значній мірі вивчені. Всього теоретично може бути створено $32^4 = 2^{20}$ КМ, а у випадку коли всі підстановки КМ різні (для забезпечення залежності від всіх розрядів початкового повідомлення) – 35960 КМ. Із них виберемо, наприклад, 64 КМ та $q = 16$, тоді $K = 134$, що відповідає розмірам ключів у сучасних криптографічних перетвореннях. В даному варіанті шифрування результати аналізу атаки на ключ в загальному відповідають результатам попереднього варіанту, але необхідно відмітити наступні особливості. Силова атака на ключ відповідає перебору всіх 2^K ключів як і в стандартних алгоритмах, але при цьому збільшується обсяг обчислень пов'язаний з визначенням $2q$ таблиць вибору підстановок. В сервісах шифрування та розшифрування таблиці вибору підстановок можуть бути визначені заздалегідь.

Висновки. Перевага першого із розглянутих варіантів полягає в більшій швидкості виконання криптографічних перетворень, при значному розмірі ключа. Другий варіант може мати широкий діапазон розміру ключа, який може відповідати діапазону розмірів ключів сучасних стандартних криптографічних перетворень. Силкові атаки на ключ в обох варіантах вимагають значних обсягів обчислень. Щодо атак лінійного криптоаналізу, то відповідний вибір підстановок КМ їх унеможливорює. Подальші дослідження полягають у визначенні КМ, використання яких унеможливають алгебраїчні та диференційні атаки на ключ, при аналізі криптограм довільної розрядності.

Література. 1. Тарасенко В.П., Тесленко О.К., Яновська О.Ю. Реалізація повних підстановок на простому двох модульному каскаді конструктивних модулів // Інформаційні технології та комп'ютерна інженерія, №1 (11), 2008, с. 88–97. 2. Тарасенко В.П., Тесленко О.К., Яновська О.Ю. Реалізація обернених підстановок на простому двомодульному каскаді конструктивних модулів // Інформаційні технології та комп'ютерна інженерія 2013, №3 (28) с. 16–25. 3. Невмержицький П.А., Тесленко О.К. “Алгоритм шифрування, який ґрунтується на суперпозиції підстановок із найпростіших регулярних ОККМ” Свідectво про реєстрацію авторського права на твір. № 66618 Дата реєстрації 13.07.2016.

Краштан Т.Є.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Аналіз існуючих алгоритмів колоризації зображень

Колоризація – процес додавання кольору до чорно-білих зображень чи відео таким чином, щоб результат виглядав візуально прийнятним і правдоподібним. Методи колоризації умовно можна розділити на три групи залежно від того, яку участь бере людина у роботі алгоритму:

- Колоризація з використанням кольорових маркерів.
- Колоризація з використанням семантично подібних зображень.
- Автоматизована колоризація.

Колоризація з використанням кольорових маркерів. Користувач подає на вхід чорно-біле зображення та додає до кожної значущої області маркер того кольору, який має бути наданий цій області. Типовими алгоритмами цього типу є Colorization using Optimization [1] та Interactive Image Colorization and Recoloring based on Coupled Map Lattices [2]. Перший метод є простішим у реалізації, однак використання другого дозволяє швидше переобчислювати кольори при внесенні в зображення додаткових маркерів. Обидва алгоритми добре колоризують розмиті зображення, а також можуть якісно працювати на детальних зображеннях (за умови маркування кожної значущої області).

Обидва розглянуті маркерні алгоритми можуть залишати помітні артефакти при поганому підборі параметрів маркерів (розташування, яскравість, товщина), зокрема при наявності на зображенні об'єктів з нечіткими межами. Перевагою щодо алгоритмів, які розглядатимуться нижче, є можливість використання цих методів не лише для колоризації, а і для заміни кольорів.

Колоризація з використанням семантично подібних зображень. Для іншої групи алгоритмів замість нанесення маркерів користувач разом з чорно-білим зображенням, яке необхідно розфарбувати, подає на вхід алгоритму зображення-приклад, яке й стає джерелом кольорів. Тут можна розглянути алгоритми, порівняння яких подано у таблиці 1.

Табл. 1. Порівняння алгоритмів колоризації, які використовують подібність зображень

	Colorization Using Similar Images [3]	Transferring color to Greyscale Images [4]	Colorization by Example [5]
Детальні зображення	+	+	-
Зображення з нечіткими межами між об'єктами	-	-	+
Розмиті зображення	+	+	+
Колоризація відео	-	-	+
Можливість використання підказок від користувача у вигляді позначок на зображеннях	-	+	+

Внаслідок того, що алгоритми з цієї групи зазвичай включають в себе використання алгоритмів з першої групи, вони вимагають більшої кількості часу для роботи і мають деякі

такі ж проблеми, як і попередні алгоритми: погане розпізнавання розмитих меж між об'єктами та складність обробки областей з різнокольоровою текстурою. Натомість, такі алгоритми не усувають необхідність людського втручання в роботу алгоритму.

Автоматизовані алгоритми колоризації. Такі методи не вимагають від користувача жодного підбору кольорів, але являють собою згорткові нейронні мережі, які необхідно тренувати перед використанням. Тут було розглянуто алгоритми Colorful Image Colorization [6], Learning Representations for Automatic Colorization [7] та Let there be Color! [8]. Останній алгоритм, на відміну від перших двох, працює з усіма типами зображень та дозволяє поєднувати колоризацію одного зображення зі стилем іншого. Усі три алгоритми коректно обробляють детальні зображення (за умови їх чіткості), однак, навіть вони часто виводять кольорові результати, що є тьмянішими за оригінали чи подекуди містять розмитий колір на межах об'єктів.

Висновки. За результатами огляду алгоритмів можна зробити висновок, що найкраще задачу колоризації нерухомих зображень розв'язують автоматизовані алгоритми, зокрема Let there be Color! Натомість маркерні алгоритми, як то Colorization using Optimization, працюють найшвидше і дозволяють користувачу самому обирати кольори, що будуть присутні на кольоровому результаті.

Література. 1. Levin A. Colorization using optimization / A. Levin, D. Lischinski, Y. Weiss. // Proceedings of ACM SIGGRAPH. – 2004. – №23. – С. 689–694. 2. Konushin V. Interactive image colorization and recoloring based on coupled map lattices / V. Konushin, V. Vezhnevets. // Graphicon. – 2006. – С. 231–234. 3. Image colorization using similar images / [R. Gupta, A. Chia, D. Rajan та ін.]. // Proceedings of the 20th ACM international conference on Multimedia. – 2012. – С. 369–378. 4. Welsh T. Transferring color to greyscale images / T. Welsh, M. Ashikhmin, K. Mueller. // ACM Transactions on Graphics. – 2002. – №21. – С. 277–280. 5. Irony R. Colorization by example / R. Irony, D. Cohen-Or, D. Lischinski. // Proceedings of the Sixteenth Eurographics conference on Rendering Techniques. – 2005. – С. 201–210. 6. Levin A. Colorization using optimization / A. Levin, D. Lischinski, Y. Weiss. // Proceedings of ACM SIGGRAPH. – 2004. – №23. – С. 689–694. 7. Larsson G. Learning Representations for Automatic Colorization / G. Larsson, M. Maire, G. Shakhnarovich. // European Conference on Computer Vision. – 2016. – С. 577–593. 8. Iizuka S. Let there be Color!: Joint End-to-end Learning of Global and Local Image Priors for Automatic Image Colorization with Simultaneous Classification / S. Iizuka, E. Simo-Serra, H. Ishikawa. // ACM Transactions on Graphics (Proc. of SIGGRAPH 2016). – 2016. – №35. – С. 110:1–110:11.

Кривоніс А.В., Волонтир О.О.

КПІ ім. Ігоря Сікорського, ФПМ, Київ, Україна

Поєднання Machine Learning та VR

Віртуальна реальність або VR (Virtual Reality) раніше сприймалась як технологія, яка здатна лише до симуляції людини, про що людство змінило свою думку після останніх хвиль інновацій, які принесла нам Oculus RIFT [1].

Сьогодні візуалізація графіки з високою частотою кадрів у віртуальній реальності відповідає швидкості та точності функціонування роботів та камер. Віртуальна реальність вже є таким інструментом, яка здатна моделюючи фізичне середовище, рух та взаємодію предметів один з одним, навчати різноманітних роботів, безпілотних апаратів та інших машин до того, як вони почнуть функціонувати в реальному світі.

В робототехніці ця технологія зробила маленький крок вперед, але передбачається, що цей крок буде значно більшим для штучного інтелекту або ML (Machine Learning). Останні досягнення вказують на потенційно руйнівну комбінацію віртуальної реальності та штучного інтелекту, яка відкриває майбутнє з безпечними та компетентними інтелектуальними машинами, здатними до самонавчання.

Не так давно компанія NVIDIA оголосила про створення моделі віртуальної реальності, розробленій в хмарі, яка використовує точне моделювання фізики для моделювання реальних середовищ. Ця система “гіперреалізації” добре підходить для навчання роботів та штучних систем в модельованих середовищах. Раніше NVIDIA [2] продемонструвала використання вхідних даних віртуальної реальності для підготовки безпілотників, використовуючи імітований візуальний вхід та перевіряючи точність навігації.

Стереоскопічні імітаційні візуальні матеріали дозволили користувачеві використовувати алгоритми візуального 3D-положення для підтримання точного положення та навігації. Цей тест був раннім свідченням того, що безпілотники та самокеровані машини можуть скоро навчитися розвиненій навігації з поєднанням реальних середовищ і візуальних віртуальних реальностей.

Аналітичний центр OpenAI оголосив у серпні, що команда розробила нейронну мережу та навчила її грати в реальному часі в популярну стратегічну гру «DOTA II Valve» [3]. Цей агент був навчений, використовуючи в якості вхідних параметрів візуальний вид зверху, тобто так як справжній гравець буде взаємодіяти з грою.

Тим не менш, зламавши гру для роботи в хмарі та надаючи системі візуального контролю агента для навчання машинам, команда розробників змогла тренувати агента шляхом самостійної гри – граючи сам з собою знову і знову, швидше, ніж в режимі реального часу, і в хмарі.

Машинний гравець перетворився з дійсно хорошого та навіть “неперевершеного” гравця, протягом тижня, перемігши кращих гравців у світі. Без користі багаторічного контексту в тому, як грають реальні люди, або ж будь-яке поняття стратегії чи тактики, нейронна мережа дізналася лише з власних успіхів та невдач у структурованому середовищі інтерактивної гри.

Отже, віртуальні середовища разом з використанням машинного навчання можуть бути цілеспрямованими та значно складнішими, здатним навчатися в процесі функціонування та змінювати свої параметри під час використання. Це дає змогу для створення нових систем прискоринного навчання нейронних мереж, та інших алгоритмів ML в будь-якій галузі діагностування та прогнозування, а також виявлення помилок на ранніх етапах.

Література. 1. Oculus Rift. (2016). The Oculus Rift – virtual headset Retrieved from https://en.wikipedia.org/wiki/Oculus_Rift. 2. Blogs Nvidia. (2016). How VR Training Keeps Fighter Pilots on Top of Their Game. Retrieved from blogs.nvidia.com/blog/2016/11/26/virtual-reality-flight-simulation/. 3. Blog OpenAI. (2018). Bot which beats the world's top professionals at 1v1 matches of Dota 2 under standard tournament rules. Retrieved from <https://blog.openai.com/dota-2/>.

Кузнецова Ю.А.¹, Цешевський В.В.^{1,2}

¹Національний аерокосмічний університет ім. М.Є. Жуковського “Харківський авіаційний інститут”, Харків, Україна; ²Sigma Software, Харків, Україна

Аналіз застосування метапрограмування для проекту Mikz Market

Кодогенерація або метапрограмування – вид програмування, пов’язаний зі створенням програм, які породжують інші програми як результат своєї роботи (зокрема, на стадії компіляції їх вихідного коду), або програм, які змінюють себе під час виконання (код, що самомодифікується). Перше дозволяє отримувати програми при менших витратах часу і зусиль на кодування, порівняно з тим, коли програміст пише їх вручну. Зменшення витрат часу – це зменшення витрат грошей. Отже, це завжди актуально для будь-якої організації, що займається програмуванням як бізнесом. Але метапрограмування не набуло широкого розповсюдження через велику складність написання кодогенераторів. Відповідь на те, як зменшити складність метапрограмування та зробити його застосування більш розповсюдженим, і намагається дати це дослідження [1].

Об’єктом дослідження є аналіз параметрів різних компонентів програмного забезпечення, що роблять їх придатними для кодогенерації. *Предметом* – методи знаходження компонентів, найбільш придатних для кодогенерації, методами експертного оцінювання й експериментальними засобами. Розрахунок вигоди застосування кодогенерації в залежності від кількості компонентів, які можна згенерувати у системі, й складності написання генератора. *Мета роботи* – прискорення процесів розробки та реінжинірингу програмного забезпечення за допомогою автоматизації написання однотипних компонентів – автоматизованої генерації коду.

Застосування. Створення програм, що генерують код, має численні застосування.

По-перше, є можливість писати програми, які будуть заздалегідь генерувати таблиці даних для використання під час виконання програми. Наприклад, якщо при створенні гри є задача виконувати швидкий пошук у таблиці синусів для всіх 8-бітових цілих чисел, то можна або обчислювати кожен синус і закодувати це таким чином, щоб програма будувала таблицю при запуску, або написати програму для створення спеціалізованого коду для таблиці перед компіляцією. Хоча для такого маленького набору чисел може мати сенс створення таблиці під час виконання, інші подібні завдання можуть надмірно уповільнити час запуску програми. У таких випадках написання програми побудови статичних даних зазвичай є найкращим варіантом [2]. *По-друге*, якщо є великий додаток, в якому безліч функцій включає довгий стереотипний код, можна створити «міні-мову», яка буде створювати стереотипний код замість програміста і надасть йому можливість кодувати тільки важливі частини програми. В такому випадку найкраще абстрагувати стереотипні фрагменти у функцію. Але часто ці фрагменти є не настільки простими. Можливо є список змінних, які потрібно оголосити в кожному примірнику, можливо є необхідність зареєструвати обробники помилок, можливо існує кілька стереотипних фрагментів, які повинні включати код за певних обставин. Все це робить створення простої функції неможливим. Часто у таких ситуаціях гарною ідеєю є створення доменної мови, що дозволить працювати з таким кодом більш простим способом. Ця доменна мова потім конвертується у початковий код на звичайній мові програмування перед компіляцією [2]. *Нарешті*, багато мов програмування змушують писати багатослівні оператори для виконання простих завдань. Програми, які генерують код, дозволяють скоротити ці оператори і зберегти багато часу, який витрачається на набір коду, що також зменшує шанс на виникнення помилок. У міру набуття мовою додаткових можливостей, програми, що генерують код, стають менш «привабливими». Те, що є доступним як стандартна можливість в одній мові, може бути доступним в іншій тільки при використанні програми, яка генерує код. Однак, неадекватний дизайн мови є не єдиною причиною необхідності в програмах, що генерують код. Важливу роль відіграє більш легке обслуговування. Програми, що генерують код, дають можливість розробляти і використовувати маленькі, специфічні для конкретної області мови, на яких легше писати і підтримувати програми, ніж писати їх одразу на цільовій мові.

Отже, огляд і аналіз сучасного стану кодогенерації та метапрограмування підтверджують доцільність використання метапрограмування на великих проєктах, а також доводять актуальність теми дослідження.

Планування. В рамках планування експериментальних досліджень доцільності використання кодогенерації у великому проєкті наведено схему експериментального дослідження, яка складається з наступних етапів: експеримент з написання універсального кодогенератору; аналіз результатів експерименту з написання універсального кодогенератору; розробка вимог до експерименту з написання вузькоспеціалізованого кодогенератору; експеримент з написання вузькоспеціалізованого кодогенератору; аналіз результатів експерименту з написання вузькоспеціалізованого кодогенератору; розробка рекомендацій з використання кодогенерації на великих проєктах. Також створено прототип універсального кодогенератору. Розроблено Web API, що генерує метадані про сутності, що знаходяться в базі даних, а потім дозволяє працювати з даними, що знаходяться в цій базі на основі раніше згенерованих метаданих. Розроблено Web-клієнт, який використовує метадані, що віддаються Web API, який вміє генерувати інтерфейс користувача на підставі цих метаданих. Наведено архітектури Web-клієнту та серверного додатку.

На жаль, як виявилось, універсальний кодогенератор підходить лише для швидкого прототипування додатків, й абсолютно не підходить для реальних бізнес-завдань. З однієї простої причини – у ньому немає ніякої можливості описати спеціалізовану бізнес-логіку для різних компонентів системи. Тож були розроблені вимоги до наступного експерименту. Визначено відгуки та фактори експерименту. Факторами, що впливають на результати, є: складність компонентів, що будуть генеруватись, та продуктивність праці програміста, що буде писати кодогенератор і використовувати його для кодогенерації. Відгуками експерименту є підходи до розробки без кодогенерації, з частковою та з повною кодогенерацією, а також кількість компонентів, що мають бути створені.

Результати. До основних результатів проведеного експериментального дослідження доцільності використання кодогенерації у великому проєкті належать такі: оцінено з точки зору доцільності генерування п'ять типів компонентів, що найчастіше зустрічаються на проєкті Mikz Market; вибраний засіб кодогенерації та необхідний інструментарій; розроблено предметно-орієнтовну мову, специфічну саме для таблиць проєкту Mikz Market; розроблено кодогенератор для часткової кодогенерації, що на основі предметно-орієнтованої мови генерує усі необхідні файли для роботи кодогенератору; підраховано кількість годин, яку було витрачено на розробку таблиць без використання кодогенерації; проведено експеримент з написання та застосування генератору для часткової кодогенерації та екстрапольовано його результати на загальну кількість таблиць у проєкті Mikz Market; підраховано очікувану кількість годин, що буде потрібна на розробку генератору для повної кодогенерації.

Таким чином, було проведено експериментальне дослідження доцільності використання кодогенерації на великому проєкті та розроблено методику аналізу доцільності використання кодогенерації на великих проєктах. Розроблено два прототипи кодогенераторів:

- універсальний кодогенератор, що генерує увесь додаток, маючи лише доменні класи;
- кодогенератор для часткової генерації лише одного типу компонентів з великого проєкту.

У ході експерименту доведено, що використання часткової кодогенерації є найшвидшим способом виконати необхідний обсяг робіт, який дозволив би заощадити до 134 годин, якби він використовувався з самого початку проєкту.

Подальші дослідження доцільно проводити в напрямку збору статистики застосування кодогенерації та збільшення числа параметрів, за якими можна визначити доцільність застосування кодогенерації.

Література. 1. Joshi Prateek. What Is Metaprogramming? Retrieved from <https://prateekvjoshi.com/2014/04/05/what-is-metaprogramming-part-22/>. 2. Jonathan Bartlett. The Art Of Metaprogramming. Retrieved from <https://www.ibm.com/developerworks/library/l-metaprogl/>.

Лимар Б.О., Олефір О.С.

КПІ ім. Ігоря Сікорського, ФПМ, Київ, Україна

Метод виявлення паркувальних місць в інтелектуальній системі паркування

Пошук місця для стоянки – звична діяльність для багатьох людей та призводить до спалення близько мільйона барелів світової нафти щодня. Водії продовжують кружляти по автостоянці, доки не буде знайдене вільне місце для паркування. Ця проблема, як правило, виникає в міських районах, де кількість транспортних засобів вище, ніж наявність місць для паркування.

Для вирішення задач виявлення пошуку вільних місць розробляються інтелектуальні системи паркування. У даній роботі пропонується використовувати камеру як датчик розпізнавання вільного місця на автостоянці з відеозображення шляхом комбінації методу регіональної сегментації зображення та сегментації базованої на кластеризації та розпізнавання круглого зображення для однозначної ідентифікації паркувального місця.

Задачами даного дослідження є створення архітектури інтелектуальної системи паркування та ефективного методу виявлення вільних паркувальних місць.

На рис. 1 зображена архітектура інтелектуальної системи паркування.



Рис. 1. Архітектура інтелектуальної системи паркування

Регіональна сегментація або порогова сегментація – загальний алгоритм сегментації, що опрацьовує зображення сірого кольору на основі сірого значення різних об'єктів та буває локальною і глобальною. Метод глобального порогу розділює зображення на дві області – об'єкти зображення та фон, як окремий поріг. Метод локального порогу розділяє зображення на декілька градацій та до кожної градації ділить зображення на кілька цільових регіонів та фонів за допомогою декількох порогів.

Перевагою методу є швидка обробка зображення за рахунок розпізнавання лише сірих об'єктів (мішеней). Метод є чутливим до шуму та напівтоновій нерівності, так як враховує лише сіру інформацію зображення [1].

Для усунення недоліків методу порогового значення було запропоновано використати метод Оцу, що зводить зображення сірого кольору до бінарного зображення. Алгоритм передбачає, що зображення містить два класи пікселів бі-модальної гистограми: пікселі переднього плану і пікселі фону, обчислюється оптимальний поріг, що розділяє два класи, так, що їх комбінований діапазон (дисперсія кластера) є мінімальним або рівноцінним (тому що сума попарних квадратичних відстаней постійна), так, що їх міжкластерна дисперсія є максимальною [2].

Література. 1. Nobuyuki Otsu (1979). A threshold selection method from gray-level histograms. IEEE Trans. Sys., Man., Cyber. 9 (1). с. 62–66. 2. Ping-Sung Liao and Tse-Sheng Chen and Pau-Choo Chung (2001). «A Fast Algorithm for Multilevel Thresholding». J. Inf. Sci. Eng.17: 713–727.

Лисецкий Ю.М.

ДП «ЭС ЭНД ТИ УКРАИНА»

Защита информации: системы резервного копирования

С ростом любой организации неизбежно растет и объем данных, находящихся на жестких дисках персональных компьютеров, серверов и систем хранения данных. Поэтому задача надежности их хранения для обеспечения целостности, конфиденциальности и доступности информации [1] является весьма важной и актуальной, особенно в условиях участвующих кибератак, вирусной активности, различного рода технических сбоев и негативного влияния «человеческого фактора». Одним из инструментов защиты информации является резервное копирование и создание архивов из исторических данных для быстрого доступа к ним при необходимости [2]. В таких случаях необходимо иметь механизмы для быстрого восстановления данных из резервных копий. Традиционные методы архивирования и восстановления данных (ручное резервное копирование файлов, архивирование или использование средств операционной системы для переноса данных на съемный носитель или в сетевую папку) малоэффективны и не способны реализовать требования по надежной сохранности и доступности, критических для организации приложений. Решить эту задачу могут системы резервного копирования (СРК), способные работать в динамично развивающейся информационно-технологической инфраструктуре (ИТИ) организации и интегрироваться с новыми приложениями и сервисами.

На сегодняшний день существует достаточно большое количество промышленных СРК с различным функционалом. Из них можно выделить основных производителей, СРК которых максимально соответствуют требованиям современных организаций. Это компании IBM с Tivoli Storage Manager, EMC с NewWorker и Avamar, Veritas с Backup Exec и NetBackup, которая является лидером в этом сегменте.

Backup Exec и NetBackup – это комплексные решения, позволяющие выполнять всевозможные виды резервного копирования, начиная от отдельных физических или виртуальных систем, заканчивая восстановлением приложений и баз данных в масштабах всей организации. Эти СРК работают на основе запатентованной технологии V-Ray, способной контролировать виртуальные файловые системы и приложения, с дальнейшим выполнением прозрачного резервного копирования, применяя гибкие варианты технологии дедупликации файлов, позволяющей в свою очередь существенно сократить объем резервной копии на системе хранения. Технология V-Ray существенно экономит время в распределенной ИТИ и упрощает управление процессами резервного копирования и восстановления. Благодаря этой технологии стало возможным уменьшить нагрузку на сеть и хранилище и восстанавливать данные быстрее без снижения производительности виртуальной среды, так как для ее работы не требуются агенты. Для пользователя эта технология позволяет восстановить нужные данные за несколько минут. Предоставляется такая возможность за счет технологии гранулярного восстановления, которая исключает необходимость долгого поиска нужного архива в системе хранения, чтобы найти нужный файл. Пользователь напрямую обращается к backup-образам систем и выбирает для восстановления нужную информацию. V-Ray автоматически защищает новые виртуальные машины, которые перешли в состояние online, при этом политика резервного копирования остается неизменной.

Также в Backup Exec и NetBackup реализована дедупликация данных – технология, позволяющая устранить дублирование данных и сократить при этом объемы физических носителей для хранения. До ее появления каждая резервная копия могла иметь одинаковый набор неизменившихся файлов со времени последней резервной копии, что приводило к увеличению объемов хранения, а с ростом изменяемых данных увеличивало и время резервного копирования. Работает она на уровне блоков данных, записанных на диск, где используются хеш-функции. Когда система дедупликации находит совпадающие хеш-функции для разных блоков, она предлагает сохранить блоки в виде единственного экземпляра и набора ссылок на него. Эта технология также позволяет сравнивать блоки данных с разных компьютеров

(глобальная дедупликация), это еще больше увеличивает эффективность дедупликации, так как на дисках разных компьютеров с одной и той же операционной системой может храниться много повторяющихся данных. В СРК компании Veritas реализованы три способа применения дедупликации: дедупликация на стороне клиента/источника, дедупликация на сервере, дедупликация на устройстве. Особенно эффективной технология дедупликации показала себя в виртуальных инфраструктурах VMware и Hyper-V, где большинство систем таких, как Windows и Linux развернуты с эталонного дистрибутива и содержат одинаковый набор программ, компонент и библиотек на своих виртуальных дисках, могут отличаться между собой только файлами и папками персональных настроек. Каждая такая виртуальная машина занимает десятки гигабайт в системе хранения. В этом случае технология дедупликации помогает снизить объем резервных копий и уменьшить время их создания. Использование технологии дедупликации позволяет ускорить резервное копирование в 10 раз, а объем резервных копий сократить примерно на 95 процентов.

Наиболее распространенным СРК является Backup Exec, который способен защитить не только физические среды, но и виртуальные. Многофункциональность Backup Exec позволяет реализовать всевозможные сценарии работы с такими объектами, как виртуальные машины, физические серверы, приложения, их компоненты, файлы, папки. Его основные преимущества: резервное копирование на любые устройства хранения и быстрое создание моментальных копий виртуальных машин; защита виртуальных машин; интегрированная функция восстановления после аварии (технология Bare Metal Restore).

NetBackup – кроссплатформенное решение резервного копирования, способное работать с очень большими объемами данных для комплексной защиты физических и виртуальных систем, приложений и баз данных в корпоративных центрах обработки данных. Эта СРК обладает высоким уровнем производительности, автоматизации и масштабируемости, ориентирована на корпоративные организации в ИТИ которых могут находиться тысячи как физических, так и виртуальных серверов. Также предлагает большой выбор компонентов для оптимизации резервного копирования и восстановления в гетерогенных средах, и поддерживает такие платформы, как HP-UX, HP Tru64, IBM AIX, Linux, Microsoft Windows, Novell NetWare, SGI IRIX и Sun Solaris. Для оперативного восстановления критически важных баз данных и приложений NetBackup использует агентов для продуктов IBM, Microsoft, Oracle, SAP и Sybase. NetBackup в своей типовой конфигурации имеет трехуровневую архитектуру: мастер-сервер, медиа-сервер, агент. Основные преимущества NetBackup: централизованное управление, автоматизированное аварийное восстановление, управление удаленными ленточными накопителями (NetBackup Vault Option). Также имеются дополнительные функции, расширяющие функциональность NetBackup: NetBackup Accelerator for VMware, NetBackup Accelerator for Microsoft Hyper-V, NetBackup for vCloud Director.

Таким образом, Backup Exec подойдет организациям, в ИТИ которых основные сервисы работают под управлением операционной системой Windows, резервное копирование которых выполняется как правило на несколько устройств хранения, в отведенное для этого ночное время, и не требует для это дополнительных ресурсов. NetBackup будет лучше применим в организациях корпоративного уровня с гетерогенной инфраструктурой, в которых сервисы работают под различными операционными системами Unix, Linux, Windows и с большими объемами данных, для резервного копирования которых понадобятся более гибкие инструменты и настройки, такие как скорость копирования, время копирования, количество копий на разные устройства и т.д.

Рассмотренные СРК являются необходимым элементом безопасности любых рабочих мест как личных, так и корпоративных, а также эффективным инструментом защиты информации.

Литература. 1. Домарев В.В. Безопасность информационных технологий. Методология создания систем защиты / Домарев В.В. – К.: ООО “ТИД” ДС 2001. – 688 с. 2. Бережной А.Н. Сохранение данных: теория и практика / Бережной А.Н. – М.: ДМК Пресс, 2016. – 317 с.

Литвинов В.А.¹, Майстренко С.Я.¹, Хурцилава К.В.¹, Костенко С.В.²

¹Институт проблем математических машин и систем НАН Украины, Киев, Украина;

²Национальный университет пищевых технологий, Киев, Украина

Оценка контролирующих и корректирующих свойств референтного словаря в системе обнаружения и исправления типовых ошибок пользователя

В настоящее время функция проверки орфографии является обязательной компонентой функционала текстовых редакторов, поисковых систем, почтовых клиентов и т.п. Центральным элементом таких систем является референтный орфографический словарь (РОС), содержащий “правильные” слова некоторой предметной области, с которыми сравниваются проверяемые слова. В [1] предлагается возможный подход к оценке и улучшению контролирующих свойств РОС, выраженных через относительное количество необнаруживаемых типовых ошибок.

Наряду с обнаружением орфографических ошибок, многие текстовые редакторы предлагают функцию автоматического исправления. Одним из типовых решений при выборе алгоритмов реализации таких функций является использование различных модификаций фонетических алгоритмов [2,3], обеспечивающих для ошибочного слова быстрый поиск “близких” (в смысле расстояния Дамерау-Левенштейна) слов РОС. Вопросам же оценки потенциальных корректирующих свойств самих РОС практически не уделяется внимания.

С целью частичного заполнения отмеченного пробела проведено моделирование корректирующих свойств словарей русского и украинского языков на основе имитационной модели [4]. В процессе моделирования для каждого слова A_j фиксировались возможные типовые ошибки $A_j \rightarrow \bar{A}_{jk}$ и определялись вероятности следующих частных исходов:

- ошибка $A_j \rightarrow \bar{A}_{jk}$ не обнаружена (вероятность Q_{jno});
- ошибка $A_j \rightarrow \bar{A}_{jk}$ обнаружена, корректировка выполнена правильно (вероятность Q_{ojnk});
- ошибка $A_j \rightarrow \bar{A}_{jk}$ обнаружена, корректировка выполнена ошибочно (вероятность Q_{ojnk});
- ошибка $A_j \rightarrow \bar{A}_{jk}$ обнаружена, корректировку выполнить не удалось (вероятность Q_{ojlk}).

Словарь	$Q_{но}$	$Q_{опк}$	$Q_{олк}$
Словарь Русской литературы $N = 161730$ слов	0.0184	0.7549	0.1443
Словарь Лопатина $N = 150213$ слов	0.0060	0.8282	0.0709
Словарь Зализняка $N = 92555$ слов	0.0054	0.8281	0.0710
Усеченный словарь Лопатина $N = 84575$ слов	0.0038	0.8518	0.0474
Украинская версия усеченного словаря Лопатина $N = 84575$ слов	0.0028	0.8610	0.0382

Ключевые результаты моделирования, обобщенные для словарей в целом, приведены в таблице. Данные таблицы с одной стороны определяют в принятых терминах показатели корректирующих (вероятности $Q_{опк}$, $Q_{олк}$) и контролирующих ($Q_{но}$) свойств исследованных словарей, и подобные показатели могут служить основой сравнения и выбора РОС.

С другой стороны, полученные данные иллюстрируют явную корреляцию между значениями $Q_{олк}$ и $Q_{но}$. Это дает основания полагать, что РОС, оптимизированные для целей улучшения контролирующих свойств [1], обладают и лучшими корректирующими свойствами.

Литература. 1. В.А. Литвинов, С.Я. Майстренко, К.В. Хурцилава, Дисфункция референтного словаря системы проверки орфографии и подход к ее снижению // Математичні машини і системи. – 2017. – №2. – С. 39–48. 2. Phonetic Algorithms [Электронный ресурс] – Режим доступа к ресурсу: <https://deparkes.co.uk/2017/12/01/phonetic-algorithms/>. 3. Фонетические алгоритмы [Электронный ресурс] – Режим доступа к ресурсу: <https://habrahabr.ru/post/114947/>. 4. В.А. Литвинов, С.Я. Майстренко, К.В. Хурцилава, Оценка контролирующих свойств базового словаря допустимых слов в системе автоматического обнаружения ошибок пользователя // Математичні машини і системи. – 2014. – №2. – С. 65–70.

Литвинюк А.А., Левчук С.Б.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Система автоматичної оптимізації асортименту в ритейлі

Вступ. Криза 2015–2016 років в Україні найбільше торкнулась споживчих галузей. За 2 роки оборот роздрібної торгівлі скоротився на 15% опустившись до показників 2010–2011 років. Зі зменшенням платеспроможності покупців і з нестабільним ринком здатність ритейлера швидко реагувати на зміни є ключовою. А це можливо тільки використовуючи новітні засоби ІТ, які дозволяють оперативного корегувати процеси ритейлера. ІТ системи використовуються для досягнення ключових задач ритейлу:

1. Забезпечення доступності товарів на полицях;
2. Розширення аудиторії покупців;
3. Зменшення списань та товарних залишків;
4. Встановлення оптимальних цін.

На перші три пункти має значний вплив правильна викладка товару на полиці, визначення максимально ефективного розміру асортименту товару в товарних групах та його розміщення в торговій точці (ТТ). На сьогодні рішення цієї задачі залежить виключно від суб'єктивної експертизи менеджера ТТ і приймаються без належного обґрунтування [1,2]. Для автоматизації та оптимізації задачі побудови асортименту було розроблено алгоритм, що базується на використанні регресійного аналізу та накопиченій історії продаж товарів.

Попередня підготовка даних. Для реалізації алгоритму необхідні наступні дані за 2 останні роки:

1. Історія продаж по товарних групах;
2. Історія середньої ціни та знижки по товарних групах;
3. Історія асортименту по товарних групах;
4. Середня ширина однієї позиції в товарній групі;
5. Загальна довжина простору полиць в торговій точці.

Використовуючи методи машинного навчання, визначається вплив сезонності, тренду, ціни та знижки на продажі. Для визначення впливу асортименту на продажі історія продажів очищається від впливу сезонності, тренду, ціни та знижки. Після очищення історії продаж визначається функція залежності продажів від асортименту.

Визначення оптимального асортименту. Алгоритм являється ітераційним з основним параметром – кількістю ітерацій. Суть полягає у оптимальному розподілі простору магазину на асортимент товарних груп. Для розуміння роботи алгоритму необхідно визначити наступні терміни:

1. Раунд — одна ітерація;
2. Гравець — окрема товарна група;
3. Сила — середня ширина однієї позиції в товарній групі;
4. Ставка — нормований градієнт функції залежності продажів від ширини асортименту при наявному асортименті;
5. Бали — виграна ширина простору в торговій точці;
6. Приз — оптимальна ширина асортименту розрахована на основі балів і сили гравця;
7. Кон — бали які розігруються в раунді.

Додаткові параметри алгоритму:

1. Кількість перемог для того, щоб гравець заснув (необхідно для того, щоб кожен гравець мав шанс на перемогу);
2. Кількість раундів сну;
3. Радіус переможців (необхідно, щоб визначити найбільш сильних гравців, для використання двох попередніх параметрів).

Алгоритм розділений на наступні етапи:

Логін В.В., Відюк П.І.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Моделювання та прогнозування характеристик трафіку цифрової Інтернет-реклами

Інтернет-реклама – реклама, що розміщується в мережі Інтернет; представлення товарів, послуг або підприємства в мережі Інтернет, адресована масовому клієнту і має характер переконання.

Ключовою відмінністю Інтернет-реклами від будь-якої іншої є можливість відстеження рекламних контактів. За рахунок можливості відслідковування реакції і дій користувача в мережі Інтернет, рекламодавець може швидко вносити зміни до чинної рекламної кампанії. Бажані дії користувача називаються «конверсією».

Найважливішими характеристиками трафіку цифрової Інтернет-реклами є показники заповнюваності («fill rate» – співвідношення кількості показів до загальної кількості запитів реклами) та доходу. Тому у даній роботі досліджувались саме вони. Тобто задача була наступною – на основі статистичних даних, що збираються рекламною компанією, потрібно зробити довгостроковий прогноз параметрів заповнюваності та доходу.

Оскільки потрібно зробити прогноз на основі статистичних даних, які збирались із однако-вим інтервалом, то маємо справу із часовими рядами. Для того, щоб побудувати оптимальну модель прогнозу, будемо будувати нашу модель методом поступового ускладнення. Тобто почнемо із найпростішої моделі та будемо її ускладнювати крок за кроком.

Перед тим, як приступити до прогнозування, потрібно визначити чи є всі наші параметри статистично значущими. Для здійснення даної задачі було вирішено використовувати кореляційну матрицю, а також метод градієнтного бустингу. Варто зазначити, що у наявному наборі статистичних даних є 19 параметрів. Із них перші 2 – нечислові (дата та назва рекламного тегу) та 17 числових.

Спочатку було досліджено показник доходу. Після кореляційного аналізу було відібрано 7 параметрів із рівнем кореляції, більшим за 0,3. Для заповнюваності кількість значущих параметрів виявилась аналогічною, але самі значущі параметри для цих двох показників відрізняються. Відібрані параметри можна побачити у табл. 1 та 2.

Табл. 1. Значення кореляції параметра «fill rate» з іншими параметрами

Параметр	«fill rate»
matched requests	0.328889
impressions	0.334531
cpm	-0.322844
skip ratio	-0.333545
eligible count	0.334614
measurable count	0.333803
viewable count	0.335137

Далі було здійснено побудову моделей для прогнозування. Спочатку проводилось прогнозування доходу. Для початку потрібно було побудувати модель множинної регресії. Після цього – модель авторегресії та модель авторегресії з ковзним середнім. Потім було вирішено до вже побудованих моделей додати декілька найважливіших регресорів, які були отримані після виконання першої частини дослідження, коли проводився аналіз значущості параметрів. Далі було вирішено побудувати модель авторегресії з інтегрованим ковзним середнім, та спробувати також додати декілька важливих регресорів [1, 2].

Табл. 2. Значення кореляції параметра «revenue» з іншими параметрами

Параметр	«revenue»
requests	0.668877
matched requests	0.459112
impressions	0.413456
skippable views	0.426618
eligible count	0.413536
measurable count	0.414247
viewable count	0.409011

У даному випадку $y(k)$ розглядається як дохід, і здійснюється побудова моделей прогнозу доходу.

Наступним етапом було прогнозування показника заповнюваності. Для виконання даного етапу було вирішено скористатись аналогічною схемою, що була використана для побудови моделей прогнозу доходу. А саме – було побудовано такі моделі:

$$\text{Множинна регресія: } y(k) = a_0 + a_1x_1(k) + a_2x_2(k) \dots + a_px_p(k) + \varepsilon(k), \quad (1)$$

$$\text{АР(p): } y(k) = a_0 + \sum_{i=0}^p a_i y(k-i) + \varepsilon(k), \quad (2)$$

$$\text{АРКС(p, q): } y(k) = a_0 + \sum_{i=0}^p a_i y(k-i) + \sum_{j=0}^q b_j \varepsilon(k-j) + \varepsilon(k), b_0 = 0. \quad (3)$$

Після цього було вирішено побудувати модель АРІКС(p, d, q), а також спробувати додати до вже побудованих моделей декілька найважливіших регресорів, отриманих на етапі визначення значущих параметрів.

У цьому випадку $y(k)$ – це показник заповнюваності або «fill rate», тобто відношення кількості показів до загальної кількості запитів реклами. Тому на даному етапі здійснювалось прогнозування показника «fill rate».

Насамкінець потрібно було порівняти точність отриманих моделей. Для того щоб виконати порівняння результатів прогнозування, було використано критерій найменшої середньоквадратичної похибки:

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 \rightarrow \min. \quad (4)$$

В результаті порівняння середньоквадратичних похибок побудованих прогнозів виявилось, що найкращою моделлю прогнозу, як для показника доходу, так і для параметра заповнюваності, є модель АРІКС із декількома важливими регресорами, оскільки саме для цієї моделі значення похибки MSE було найменшим. Це можна побачити у табл. 3.

Табл. 3. Значення похибки MSE для кінцевих моделей

Параметр	АР	АРКС	АРІКС	АРІКС+регресори
Дохід	85	64	56	43
Fill rate	0.0086	0.0072	0.0061	0.0054

Література. 1. Бідюк П.І. Аналіз часових рядів (навчальний посібник) / П.І. Бідюк, В.Д. Романенко, О.Л. Тимошук. – Київ: Політехніка, 2010. – 317 с. 2. Hamilton J.D. Time Series Analysis / James Douglas Hamilton. – Princeton: Princeton University Press, 1994. – 820 p.

Лютенко І.В., Курасов О.І., Соколов Д.В.

Національний технічний університет "Харківський політехнічний інститут", Харків, Україна

Використання методів послідовного ранжування альтернатив для оцінювання веб-сайтів

Завдяки бурхливому розвитку мережі Інтернет в усьому світі клієнт може придбати необхідні товари або послуги за допомогою веб-сайту. За таких умов якість веб-сайту є критичним фактором успішності ведення бізнесу.

Необхідність забезпечення високої якості веб-сайту зумовлює виникнення задачі оцінювання його якості, що є складовою частиною задачі підтримки прийняття рішень. Згідно неї, дається множина варіантів $A_1, \dots, A_p$, що можуть представляти веб-сайти, що оцінюються. Кожен з них характеризується певними критеріями $K_1, \dots, K_m$. Кожен критерій K_i має шкалу $X_i = \{x_i^1, \dots, x_i^{g_i}\}$, $i = 1, \dots, m$, дискретні числові або вербальні градації в якій у більшості випадків є впорядкованими.

У разі, якщо виконується порівняння з врахуванням великої кількості критеріїв (більше 20), виникає проблема, яка полягає у тому, що формальне порівняння лише за одними значеннями критеріїв стає неможливим, а спроби зменшити їх кількість призводить до зниження якості остаточного результату внаслідок його віддалення від дійсності. Тому необхідно знайти метод, який забезпечить вирішення задачі багатокритеріального вибору у просторі великої розмірності шляхом скорочення кількості вимірів, опираючись на специфіку предметної області.

Для вирішення задачі оцінювання веб-сайтів може бути використаний метод ПАКС («Последовательное Агрегирование Классифицируемых Состояний»), який заснований на використанні методів вербального аналізу для зменшення розмірності задачі [2].

Процедура оцінювання в такому разі включає декілька основних етапів:

1. Побудова ієрархічної системи агрегованих критеріїв веб-сайту. Процес побудови полягає у створенні інтегральних показників, які агрегують початкові критерії, які були відібрані особою, що приймає рішення, як ключові для оцінювання якості веб-сайту.
2. На другому етапі виконується послідовна побудова шкали кожного складеного критерію, що є найбільш показовою комбінацією оцінок вихідних критеріїв.
3. На цьому етапі виконується кінцевий відбір кращого сайту в отриманому просторі комплексних критеріїв за допомогою методу АРАМИС («Агрегирование и Ранжирование Альтернатив Многопризнаковых Идеальных Ситуаций») [2] за яким сайти упорядковуються за наближенням до еталонного зразка A_+ , за відстанню до найгіршого зразка A_- або за значенням показника відносної близькості до найкращого зразка

$$l(A_q) = d(A_+, A_q) / [d(A_+, A_q) + d(A_-, A_q)], \quad (1)$$

$d(A_+, A_q)$ – відстань до найкращого зразка A_+

$d(A_-, A_q)$ – відстань до найгіршого зразка A_- .

Такий методологічний підхід до зниження розмірності простору якісних ознак має певну універсальність, тому що в загальному випадку дозволяє оперувати як з символіною (якісною), так і з числовою (кількісною) інформацією, представляючи кожен градацію шкали складеного (агрегованого) критерію у вигляді комбінації градацій оцінок вихідних показників, що дозволяє знизити суб'єктивність оцінки під час оцінювання веб-сайтів. Для побудови шкал складених критеріїв можна використати більшість методів ранжування альтернатив, що дозволяє обрати для практичної задачі найкращий набір та способи побудови складених критеріїв.

Література. 1. Петровский А.Б. Многокритериальный выбор с уменьшением размерности пространства признаков: многоэтапная технология ПАКС / А.Б. Петровский, Г.В. Ройзензон // Искусственный интеллект и принятие решений. – 2012. – № 4. – С. 88 – 103. 2. Петровский А.Б. Теория принятия решений / А.Б. Петровский Москва: Издательский центр «Академия». – 2009. – 402 с.

Максим К.Є.

КПІ ім. Ігоря Сікорського, ФПМ, Київ, Україна

Дослідження вразливостей спекулятивного виконання інструкцій на процесорах ARM

На початку січня стало відомо про вразливість, до якої схильні практично всі сучасні комп'ютери та телефони. Головний корінь проблеми – так зване спекулятивне виконання команд. Розглянемо детальніше проблеми використання цієї технології.

Спекулятивне, або випереджаюче виконання дозволяє підвищити продуктивність, оптимізуючи завантаження операційних блоків CPU. Спрощено, це технологія в сучасних мікропроцесорах, яка дозволяє «передбачати», які команди буде потрібно виконати в майбутньому. Випереджаюче виконання поширюється не тільки на лінійні ділянки програми, а й на умовні переходи та непрямі форми передачі управління, а саме перехід за адресою, що знаходиться в регістрі або комірці пам'яті [1]. Щоб точки розгалуження програмного коду не зупиняли його випереджальне виконання, використовується логіка передбачення розгалужень. Збираючи статистику про раніше виконані розгалуження, процесор визначає найбільш ймовірні цільові адреси, віддаючи їм пріоритет при випереджальному виконанні інструкцій, з розрахунку на те, що після умовного або непрямого переходу, виконання програми продовжиться по передбаченій адресі. У цьому випадку результати завчасно виконаних дій будуть затребувані. Але при будь-якому прогнозі можливі помилки. Якщо умовний або непрямий перехід відбувся за адресою, відмінною від передбаченого і завчасно виконана гілка програми не отримала управління, то результати спекулятивного виконання повинні бути знищені [1]. Все це відбувається у внутрішній пам'яті процесора на кількох рівнях кеша. Уразливість спекулятивного виконання полягає в тому, що зломисники можуть отримати доступ до даних, які знаходяться в цій загальній пам'яті [3]. Група дослідників з безпеки з Google Project Zero досліджували чи можна передати вміст незаконно прочитаного байта зі спекулятивної гілки виконання в нормальну. З'ясувалося, що це можливо. Для цього спекулятивний код може виконати дії, що впливають на таймінги виконання програми. Причому рівень або тип цього впливу повинен стійко залежати від вмісту забороненого байта, так, щоб нормальна гілка виконання змогла методом вимірювання таймінгів розшифрувати своєрідне повідомлення, сформоване спекулятивної гілкою [2]. Один з варіантів: створити в загальнодоступному просторі користувача пам'яті вимірювальний масив. Виконавши в спекулятивному коді звернення до одного елементу масиву, адреса якого залежить від вмісту забороненого байта, автор шкідливого коду забезпечить кешування цього елемента. Далі, в нормальній гілці слід порівняти час читання елементів масиву, скориставшись таймером, наприклад лічильником процесорних тактів (TSC). Попередньо кешований елемент прочитає швидше за інших. Знаючи його адресу, обчислюємо вміст забороненого байта.

Для запобігання проблеми значення забороненого байта, перед використанням в якості індексу для адресації вимірювального масиву, розумно помножити на розмір сторінки пам'яті (4 кілобайти), що рівнозначно зрушенню на 12 біт вліво. Такий прийом дозволяє задіяти в вимірювальному механізмі кеш трансляції сторінок, що, в порівнянні з кеш-пам'яттю даних забезпечить більш стабільний результат, оскільки кожній відміченій комірці буде відповідати власна сторінка пам'яті. Отже, можна підсумувати, що наявність уразливості і особливості її уникнення будуть залежати від глибини спекулятивних механізмів а також розміру, асоціативності і алгоритмів роботи кеш-пам'яті, а розглянутий алгоритм дозволить уникнути проблеми значення забороненого байта і частково зменшить вразливість системи.

Література. 1. FireEye Threat Intelligence (9 July 2015). "HAMMERTOSS: Stealthy Tactics Define a Russian Cyber Threat Group". Retrieved 7 August 2015. 2. Конахович Г.Ф. Микропроцесоры / Г.Ф. Конахович, А.Ю. Пузыренко. – К.: «МК-Пресс», 2006. – 288 с. 3. Holub, V. Low Complexity Features for JPEG Steganalysis Using Undecimated DCT / V. Holub, J. Fridrich // IEEE Transactions on Information Forensics and Security. – 2015. – vol. 10 (2). – P. 219–228.

Пациента].pdf (для создания pdf файла используется некоторые функции из дополнительной библиотеки reportlab [7]). Работа анализатора построена непосредственно по стандартной методике анализа полученных результатов исследования. В начале создается словарь слов для заключения, и по ходу анализа программа в зависимости от полученных цифр выбирает необходимые слова, компоуя их в предложения. После этого происходит общая компоновка файла, который включает в себя шапку национального аллергоцентра, основную информацию о пациенте, дату проведения анализа, массив результатов анализов с таблицей для проверки, а также само заключение по каждому препарату в отдельности. Внизу добавляется подпись иммунолога, который проводил данный анализ и врача-аллерголога. Функция CreatePDF() возвращает ссылку на готовый pdf файл, которая принимается функцией Sender(). Эта функция с помощью библиотеки email [8] подключается к SMTP-серверу почты ukr.net, и формирует письмо, состоящее из комментария к результатам анализа и прикрепленного файла с результатами анализа. Данное письмо отправляется по указанной пациентом почте. После отправки всех писем программа завершает работу и выводит отчет по отправке результатов.

Выводы. Данный программный продукт позволяет повысить эффективность выполнения процедуры алергодиагностики за счет увеличения скорости обработки результатов анализа, уменьшения нагрузки на врача-иммунолога, что позволит ему оказывать помощь большему количеству пациентов. Помимо того, исключена ошибка при формировании и отправке отчета, так как это проводится автоматически данной программой.

Литература. 1. Бабаджан В.Д. Лабораторна діагностика медикаментозної алергії // Астма та Алергія. 2013. № 3. С. 21–27. Медикаментозна алергія, включаючи анафілаксію. Адаптована клінічна настанова, заснована на доказах // Державний експертний центр міністерства охорони здоров'я України, Українське товариство спеціалістів з імунології, алергології та імунореабілітації. 2015. URL : http://www.lorlife.kiev.ua/rhinology/2015/2015_3_4_29.pdf (дата обращения: 26.03.2018). 2. Кайдашев И.П. Гиперчувствительность к лекарственным препаратам :руководство для врачей // Киев: Медкнига, 2016. 288 с. (дата обращения: 15.03.2018). 3. Калинина Н.М., Дрыгина Л.Б., Алхутова Н.А. Современные технологии in vitro алергодиагностики // Лабораторная диагностика ИНВИТРО. 2011. URL: <http://www.labdiagnostic.ru/docs/specialists/allergy.html> (дата обращения: 15.03.2018). 4. Анисимов М.В., Аллергические реакции на местные анестетики в стоматологии// Газета «Новости медицины и фармации» Аллергология и пульмонология (380) 2011 URL: <http://www.mif-ua.com/archive/article/21219> (дата обращения: 10.03.2018). 5. A Python library to read/write Excel 2010 xlsx/xlsm files (2014). Retrieved from <https://openpyxl.readthedocs.io/en/stable/>. 6. NumPy User Guide (2018). Retrieved from <https://docs.scipy.org/doc/numpy/user/index.html>(дата обращения: 01.03.2018). 7. ReportLab PDF Library (2016). Retrieved from <https://www.reportlab.com/docs/reportlab-userguide.pdf> (дата обращения: 01.03.2018). 8. 18.1.11. email: Examples (2014). Retrieved from <https://docs.python.org/2/library/email-examples.html> (дата обращения: 01.03.2018).

Огурцов М.І.

Інститут кібернетики ім. В.М.Глушкова НАН України, Київ, Україна

Розробка протоколу захищеного обміну даними для спеціальних мереж

Вступ. Через стрімке зростання масштабу, складності задач і розширення сфер практичних застосувань мереж спеціального призначення підвищується актуальність питання захищеного обміну даних у таких мережах. Зокрема, велику актуальність набувають питання розробки алгоритмів захищеної передачі даних та їх маршрутизації для спеціалізованих мереж, призначених для функціонування системи керування рухомими технічними об'єктами з мультимедійною інформацією і кодованих команд у зашифрованому виді. Існуючі стандарти для тимчасових безпроводних мереж і мобільні пристрої, які доступні на ринку, не передбачають роботу в умовах використання засобів радіоелектронної боротьби, високого рівня активних завад та хакерських атак [1–6], а також не відповідають ДСТУ.

Якщо не використовувати технології захисту інформації, то потік даних, що циркулює в мережі, буде доступний будь-кому, хто має відповідні технічні засоби [7]. Тому функціонування засобів захисту даних є невід'ємною частиною інформаційних процесів у цих мережах. На сьогоднішній день в Україні у більшості випадків засоби захисту інформації в спеціальних мережах не застосовуються [8]. А в тій незначній частині, де вони застосовуються, звичайно використовують лише стандартні засоби захисту інформації від виробника обладнання, що застосовується для побудови спеціальних мереж. В поточних геополітичних умовах для нашої держави такий підхід є неприйнятним, тому розробка загального протоколу захищеного обміну даними для спеціальних мереж є актуальною науковою задачею.

Мета досліджень. Метою досліджень стала розробка протоколу, математичного апарату і відповідного програмного забезпечення з використанням синхронної і асинхронної передачі зашифрованих пакетів даних та кодованих команд в спеціалізованій безпроводній мережі.

Передача зашифрованих повідомлень. Завдання пересилання зашифрованих повідомлень є актуальним як для передачі сигналів управління, так і для передачі фото- та відеоданих в мережі [8]. Для шифрування після проведеного аналізу існуючих алгоритмів був обраний для використання симетричний алгоритм AES [1, 7, 9, 10].

Протокол захищеного обміну даними в спеціальній мережі. Для розробки протоколу захищеного обміну даними в спеціальній мережі, був виконаний аналіз існуючих методів та алгоритмів захисту інформації в подібних системах. На основі отриманих результатів проведеного аналізу був розроблений новий протокол захищеного обміну даними в мережі.

Розглянемо *алгоритм його роботи*.

1. Для випадку зв'язку типу точка-точка (наприклад, за наявності лише двох вузлів) алгоритм роботи захищеної передачі даних за цим протоколом є очевидним. Дані передаються з підтвердженням отримання задля збереження синхронізації. Цей алгоритм повторюється для кожного нового сеансу зв'язку.
2. Алгоритм роботи протоколу захищеного обміну даними між вузлами для випадку однорангової mesh-мережі, що містить більше одного учасника:
 - На першому етапі роботи, виконується ініціалізація мережі, на кожному її елементі незалежно будується таблиця маршрутизації. Таблицю маршрутизації кожен вузол мережі будує незалежно – це зручніше, ніж її весь час розсилати, та й це б суперечило логіці однорангової mesh-мережі на вимогу. Визначається, з якими з елементів мережі є прямий зв'язок, з якими – зв'язок лише через ретрансляцію. Далі, коли є необхідність передати дані за якоюсь адресою, визначається, чи є прямий зв'язок з отримувачем, чи потрібна ретрансляція. В будь-якому випадку визначається, кому потрібно передати пакет даних.
 - Пакет даних це: 128 біт зашифровані дані + 4 біта адреса + (за необхідності) додаткові службові дані. У більшості випадків маємо невелику групу вузлів мережі

(зазвичай до 4), тому логічно зробити маленький адресний пул з прив'язкою адрес – для 8 вузлів це буде 2 в 3 ступені – 3 біта. Отже, 4 біта нададуть необхідний рівень надлишковості.

- Якщо з цим вузлом мережі обмін даними попередньо не виконувався, відбуваються ті ж дії, що і у випадку 1 – на основі фізичних випадкових даних генерується початковий вектор IV. Він шифрується довготерміновим ключем, передається цьому вузлу і далі зберігається обома вузлами у таблиці маршрутизації для відповідного вузла.
- В подальшому цей вектор (що змінюється після кожного переданого пакета) починає використовуватись для передачі даних між ними з підтвердженням отримання кожного пакету. Таким чином, для кожної пари вузлів буде використовуватись індивідуальний вектор IV. Це призведе до зростання обсягу пам'яті, що потрібно мати на кожному вузлі зв'язку (додаткові 128 біт на кожен вузол мережі), але ці вимоги не є значними, а при цьому надається високий рівень захисту даних, що циркулюватимуть цією мережею.
- Алгоритм передачі пакету по поточній таблиці маршрутизації: вузол мережі визначає, чи безпосередньо він має посилати пакет до отримувача, чи через ретранслятор. Адреса отримувача передається в незашифрованому вигляді, тому одержувач отримує цю адресу без необхідності розшифрування пакету. Якщо адреса не його – він визначає по власній таблиці маршрутизації, як і відправник, куди слати цей пакет – і повторює процедуру з першого підпункту.

Висновки. В результаті наукових досліджень створено і апробовано протокол захисту передачі даних для мереж спеціального призначення з урахуванням особливостей спеціальних мережі, зокрема, у складних ситуаціях, що відповідають українським та міжнародним стандартам. Розроблено алгоритми криптографічного захисту інформації, яка циркулює в таких мережах. Розроблені алгоритми дозволяють виконувати захист інформаційних потоків для децентралізованих чарункових та ad hoc мереж у польових умовах. Розроблене програмно-алгоритмічне забезпечення апробовано шляхом створення модельних зразків мереж спеціального призначення та проведення тестування мережевої взаємодії їх вузлів. Проведені натурні експерименти з апробації розробленого програмно-алгоритмічного забезпечення у лабораторних умовах підтвердили його застосовність і працездатність та довели потенційну можливість його впровадження. Проведене тестування на практиці показало, що затримки шифрування для реалізації розробленого протоколу склали декілька десятків мс, що дозволяє без проблем передавати сигнали, текст, службові команди, відео, зображення та звук. Досягнуті результати дозволять підвищити надійність, захищеність та керованість спеціальних мереж у польових умовах.

Література. 1. Хорошко В.А., Чекатков А.А. Методы и средства защиты информации. – К.: “Юниор”, 2003. – 504 с. 2. Домарев В.В. Безопасность информационных технологий. Методология создания систем защиты. – К.: ООО «ТИД ДС», 2001. – 688 с. 3. Петраков А.В. Основы практической защиты информации. – М.: Радио и связь, 2005. – 384 с. 4. Корольов В.Ю. Персоналізація віртуальних обчислювальних ресурсів і інформаційних джерел в сервісо-орієнтованих архітектурах // Вісті академії інженерних наук України. – № 4 (34). – 2007. – С. 13–20. 5. Степанов Е.А. Корнеев И.К. Информационная безопасность и защита информации. – М.: ИНФРА-М, 2001. – 304 с. 6. Олейников Е. Экономическая и национальная безопасность. – М.: Экзамен, 2005. – 768 с. 7. Bhaiji Yu. Network Security Technologies and Solutions. – Cisco Press, 2008. – 840 p. 8. Корольов В.Ю., Поліновський В.В., Огурцов М.І. Моделювання мереж зв'язку рухомих дистанційно керованих систем на базі HLA // Вісник Хмельницького національного університету. Технічні науки. – 2017. – № 1 (245). – с. 160–165. 9. Фергюсон Н., Шнайер Б. Практическая криптография.: Пер. с англ. – М.: Издательский дом “Вильямс”, 2005. – 424 с. 10. Schneier B. Applied cryptography: protocols, algorithms, and source code in C. – John WileyAndSons, 2007. – 816 p.

Орлова М.М., Шитий Д.В.

КПІ ім. Ігоря Сікорського, ФПМ, Київ, Україна

Способи та засоби підтримки багатопотокої передачі в комп'ютерних мережах

Від сучасних систем передачі даних вимагають широкосмуговий та багатопотоковий доступ. Актуальним є питання підвищення стійкості з'єднань при перериванні з'єднання в експериментальних системах, що направлені на використання в корпоративних умовах (сумісні з IPv6). Тому на сьогодні собою представляє інтерес виявлення впливу розриву з'єднання на фізичному рівні для додатків користувача у комунікаційних системах, що дозволить досягти зазначених цілей: стійкість при перериванні з'єднання на рівні інтерфейсу на основі протоколу IPv6 [1].

На сьогодні актуальним є питання забезпечення стійкості з'єднань при перериванні з'єднання та збільшення пропускної спроможності в системах, орієнтованих на використання у корпоративних мережах з протоколом IPv6 та множинну адресацію.

Аналіз сучасного стану предметної області. Протокол Multipath Transmission Control Protocol [2] є розширенням протоколу TCP [3], яке орієнтоване на використання робочими станціями, і забезпечує надійне багатопотокове з'єднання між кінцевими вузлами мережі в разі розриву одного з TCP потоків [4]. За рахунок передачі додаткових потоків суттєво збільшується функціональна надійність взаємодіючих модулів та загальна пропускна спроможність мережі без заміни обладнання та комунікаційних каналів. При декількох потоках передачі даних з'являється необхідність у планувальнику пакетів, що повинен враховувати всі можливі сценарії розподілення пропускної спроможності та відповідно реагувати на переривання встановленого з'єднання.

Пропозиції автора, направлені на досягнення поставленої мети. Для вирішення поставленої задачі запропоновано спосіб забезпечення швидкодії функціонування планувальника MRTCP у разі переривання зв'язку або шляхів з однаковою пропускною спроможністю. Спосіб полягає в реалізації кроків, представлених нижче. Спочатку ініціюється $P = p_1, \dots, p_n$, де P – набір шляхів. Кожному шляху відповідає інтервал часу Round Trip Time – r_i , і на кожен шлях пов'язаний з відмінним вікном перенавантаження.

Планувальник генерує набір $Q = q_1, \dots, q_n$, де елемент q_1 – пара (q_1, p_1) , тобто пакет₁ буде відправлено шляхом₁. Представлений планувальник модифікований для випадку розриву з'єднання та для випадку черг з однаковою пропускною спроможністю. Особливість запропонованого способу полягає у тому, що він враховує випадок, коли канали мають однакову пропускну спроможність та надає інструменти для реагування на розрив з'єднання без втрати загальної пропускної спроможності.

У разі коли канали мають пропусну спроможність, то планувальник пакетів розподіляє їх на одноаківні частини, беручи до уваги також вікно перентавантаження. Якщо відбувається переривання каналу, то пакети, що вже були поставлені у відповідність деякому шляху необхідно поставити у відповідність тому шляху, де є активне з'єднання. Тобто, якщо буде перерване з'єднання, в результаті якого не функціонує швидкий шлях, необхідно всі пакети з черги на швидкий шлях перемістити на повільний шлях.

Та ж сама операція відбудеться і для повільного каналу. Якщо повільний канал буде перервано, всі пакети, що вже поставлені у відповідність повільному каналу будуть переставлені у відповідність швидкому каналу.

Висновок. Запропоновано спосіб планування пакетів для протоколу MRTCP, що відрізняється від відомих підтримкою однакових або різних за пропусною спроможністю шляхів, а також плануванням пакетів у разі переривання з'єднання.

Література. 1. [RFC 2460], <https://www.ietf.org/rfc/rfc2460.txt>. 2. [RFC 6824], <https://tools.ietf.org/html/rfc6824>. 3. [RFC 793], <https://tools.ietf.org/html/rfc793>. 4. [RFC 6897], <https://tools.ietf.org/html/rfc6897>.

Пашаева М.М.

Научно-исследовательский институт аэрокосмической информатики НАКА, Баку, Азербайджан

Классификация различных типов растительного покрова на основе космических снимков и технологии ГИС

Инвентаризация сельскохозяйственных культур в республике и анализ состояния сельскохозяйственных земель является актуальной проблемой. Выбраны для исследования экономический регион Губа – Хачмазского района является ведущим в республике производителем сельскохозяйственной продукции. В статье использована технология ГИС для классификации различных типов посевов в выбранном для исследований тестовом экономическом районе, на основе обработки космических изображений охватывающих разные периоды времени. Оценка состояния посевов с использованием вегетационных индексов растительности NDVI в исследуемом регионе и их распределение по территории Республики является одной из наиболее актуальных и практических проблем. В работе использовались космические снимки Landsat 8 за период времени 06.04.2016, 15.05.2016 и 31.05.2016, охватывая различные периоды вегетации сельскохозяйственных культур (рис. 1).

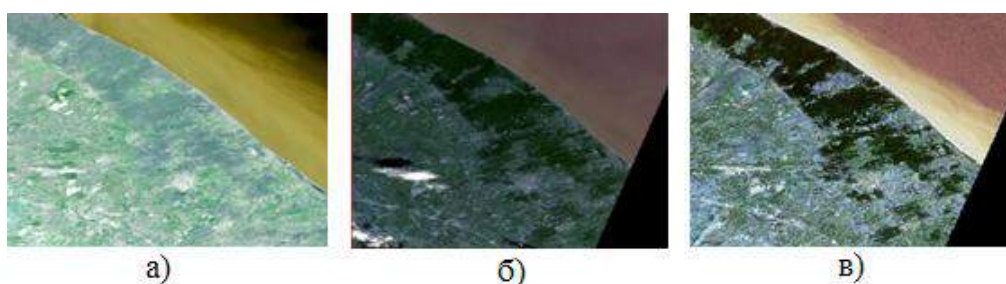


Рис. 1. Космические изображения Губа – Хачмазского района: а) Landsat 8 – 06.04.2016, б) Landsat 8 – 15.05.2016, в) Landsat 8 – 31.05.2016 г.

Объектами классификации в выбранном тестовом районе являются:

- Фруктовые сады;
- Виноградники;
- Лесные посадки;
- Зерновые посадки;
- Пастбищные угодья.

Для мониторинга динамики созревания сельскохозяйственных культур использовались вегетационные индексы растительности [2]. NDVI, рассчитанные на основе спутниковых изображений Landsat 8, охватывающих 3 разных периода времени созревания культур (рис. 2).

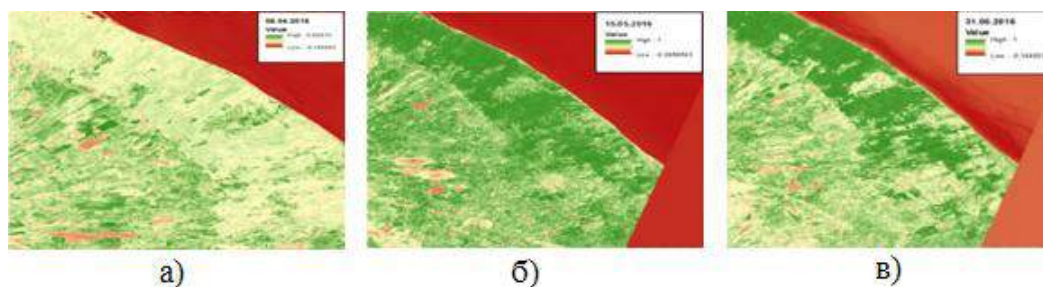


Рис. 2. Космические снимки значений тестового района со спутника Landsat 8 а) 06.04.2016, б) 15.05.2016, в) 31.05.2016 г. и индексы NDVI, рассчитанные на их основе

Согласно результатам классификации, сельскохозяйственных культур, были оценены значения пикселей растительности на основе созданных моделей NDVI (табл. 1). Модели NDVI, разработанные на основе космических снимков за разные периоды времени анализировались на основе значений пикселей, а значения пикселей были представлены графически для каждого периода времени (рис. 3).

Таблица 1. Исследование динамики развития сельскохозяйственных культур в соответствии с индексными моделями NDVI рассчитанными на основе космических снимков Landsat 8 за разные периоды времени

Сельхоз. посадки	06.04.2016	15.05.2016	31.05.2016
Виноградники	0,40	0,56	0,57
Зерновые посадки	0,45	0,43	0,24
Пастбищные угодья	0,38	0,52	0,39
Фруктовые сады	0,42	0,54	0,51
Лесные посадки	0,27	0,53	0,56

Процесс классификации выполнен на основе значений пикселей, учитывая, что методы дистанционной классификации основаны на различных спектральных представлениях объектов [3]. Существование подобных или схожих классов растительности в классификации приводит к определенным ошибкам в результатах классификации. Тем не менее, можно частично устранить неточности с классификацией на основе использования космических снимков за разные периоды времени. Первоначальные результаты классификации по космическим снимкам от 06.04.2016, в основном, относили пиксели к классам объектов “зерновые посадки” и “садовые культуры”. Однако добавление космических снимков от 31.05.2016 и 15.05.2016 позволило значительно улучшить полученные результаты.

Как мы уже упоминали, классификация была выполнена для выбранного тестового участка в экономическом регионе Губа – Хачмаз площадью в 7630 гектаров, с координатами 41047', 48034'. На рис. 5 показан результат классификации с определением площадей классов растительности расположенных в исследуемой тестовой зоне. Полученные результаты были представлены в виде таблицы (таблица 2) (зеленый цвет леса, светло-коричневый цвет – зерновые посадки, светло-голубой – трава, темно-коричневый – виноградник, фиолетовый цвет – фруктовые сады и красный цвет – другие виды, как показано на карте).

Таблица 2. Показатели классификации растительных объектов в области исследования по соответствующим классам

N	Класс	Площадь, га
1	Лесные посадки	3040,83
2	Виноградники	1119,51
3	Фруктовые сады	1084,14
4	Зерновые посадки	971,19
5	Другие культуры	819,9
6	Пастбищные угодья	595,08
	Итого:	7630,65

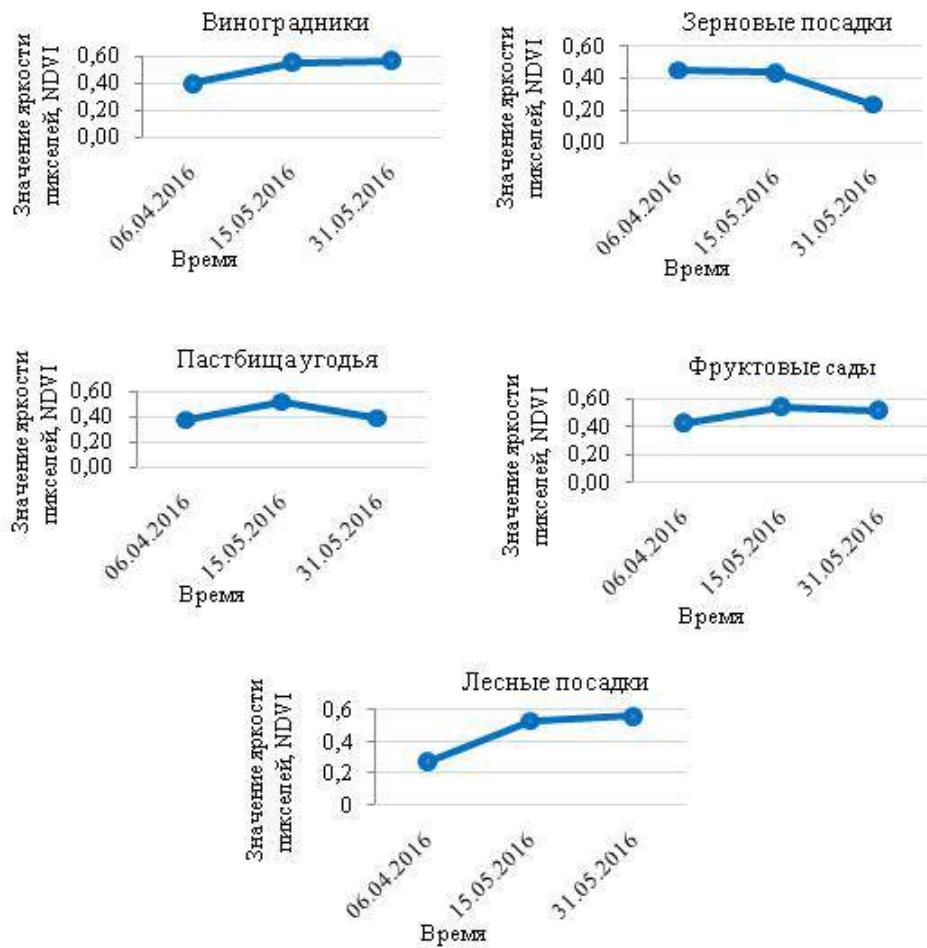


Рис. 3. Графическое представление значений пикселей классов растительности рассчитанных на основе индексов NDVI и космических снимков Landsat 8, за разные периоды времени

Как видно из таблицы, в составленной электронной карте тестового региона значительная часть тестовой зоны (40%) лесистая. Очевидно, это связано с низким спектральным разрешением снимков со спутника Landsat 8. Однако метод с точки зрения точности является оптимальным. Так, учитывая геометрическую характеристику спутника, нет сомнения, что этот метод способен дать верный результат.

Однако предложенный метод попиксельной классификации при использовании снимков с высоким пространственным разрешением способен значительно улучшить точность классификации растительных объектов с близкими спектральными характеристиками.

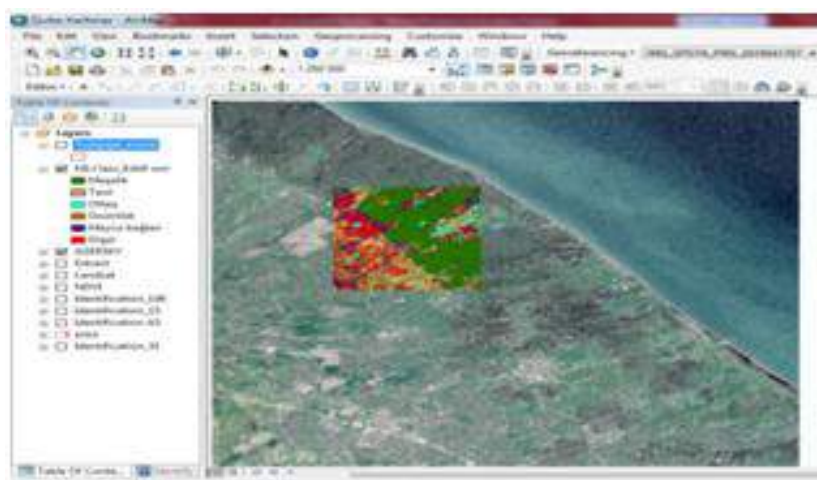


Рис. 4. Классифицированное изображение классов растительности исследуемой зоны экономического района Губа – Хачмаз, на основе космического снимка Landsat 8

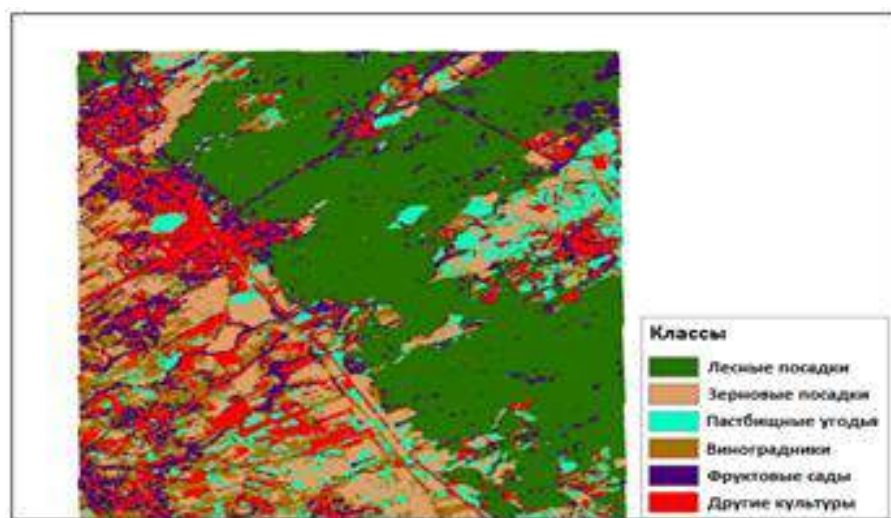


Рис. 5. Классифицированное изображение исследуемой зоны экономического района Губа – Хачмаз на основе космических снимков от 06.04.2016, 15.05.2016, 31.05.2016

Литература. 1. http://azerbaijan.az/_Economy/_Agriculture/_agriculture_a.html. 2. С.Duran Uzaktan algılama teknikleri ile bitki örtüsü analizi// Doğu Akdeniz Ormancılık Araştırma Enstitüsü, 2017, sah. 16–19. 3. E.Ayhan, F.Karsli, E.Tunç “Uzaktan algılanmış görüntülerde sınıflandırma ve analiz”, 2017, sah. 32–36.

Петрішенко С.О.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Сучасна концепція Інтернету речей

Вступ. Інтернет речей – концепція взаємопов’язаних фізичних пристроїв, які мають вбудовані датчики або сенсори, а також програмне забезпечення, яке дозволяє здійснювати передачу і обмін даними в режимі реального часу з комп’ютерними системами і фізичним світом [1]. Концепція IP (Інтернет речей) стає все більш популярною в сучасному світі, адже вона має на меті замінити багато процесів і подій, які б відбувалися без участі людини на основі аналізу взаємодії “розумних пристроїв”, наприклад Apple Watch – це пристрій IP, автомобіль, який отримує оновлення програмного забезпечення, як-от автомобілі компанії Tesla – теж частина цієї концепції. Але більша частина IP прихована в промисловому обладнанні, спокійно роблячи виробництво більш ефективним.

Об’єкти даних. Додатки IP повинні мати гнучке та масштабоване управління даними для пристроїв та керування подіями. Платформа IP повинна задовольняти ці ключові вимоги, щоб забезпечити гарну основу для цих додатків. З точки зору управління даними, пристрої та події є двома різними поняттями. Тому можна виділити два центральних об’єкта даних [2]:

- *Пристрій.* Більшість програм IP побудовані на певному фізичному пристрої. Типи пристроїв можуть сильно відрізнятися – від ручних електроінструментів до великих вантажних автомобілів.
- *Події.* Керування подій, створених пристроями, є критичним для більшості додатків IP. Ці події можуть варіюватися від технічних подій, таких як відмова пристрою до вираження інтересів покупця в конкретному торговому центрі.

Ієрархії управління. Ефективне управління пристроями та подіями залежить від створення незалежних ієрархій [2]:

- *Невеликі автономні елементи апаратури,* такі як датчики, контролери тощо.
- *Пристрій,* які можуть підключатися до кількох елементів обладнання через локальну мережу, використовуючи спеціалізовані драйвери для різних типів обладнання. Пристрій зазвичай підключається до платформи додатків через мережу мобільного оператора (GSM, EDGE) або через пряме підключення до Інтернету.
- *Системи систем* – це кілька пристроїв, які можуть бути згруповані разом, щоб сформувати так звану систему систем, наприклад електромережа, що складається з декількох електростанцій тощо.

Управління даними. Виділяють наступні типи даних [2]:

- *Елементи пристрою:* кожен пристрій стає окремим елементом у базі даних центрального пристрою, включаючи інформацію про його розташування, а також додаткові атрибути та властивості.
- *Події, пов’язані з окремими пристроями:* кожен пристрій містить повну інформацію про усі події, пов’язані із ним, включаючи технічні події (наприклад, сповіщення про відмову тощо).

Автоматизоване керування учасниками. Існує декілька технологій, які використовуються в контексті IP, такі як:

- *iBeacon* – технологія, яка використовує можливості смартфона (BLE, NFC, GPS) для надання місцезнаходження, надсилання сповіщень тощо. Окрім контролю місцезнаходження, програмний додаток, який встановлено на мобільний пристрій, має інформацію про те, як далеко перебуває датчик iBeacon. Замість того, щоб використовувати довготу і широту для визначення положення, iBeacon використовує сигнал BLE, джерелом якого являється фізичний пристрій (смартфон, “розумний годинник” тощо) [3].
- *RFID* – технологія, в якій розпізнавання здійснюється за допомогою закріплених за об’єктом спеціальних міток, що містить інформацію про нього [4]. Вона дозволяє автоматично

реєструвати користувачів, взаємодіяти один з одним.

Технології iBeacon використовуються в торговельних центрах, залізничних вокзалах, аеропортах тощо для своєчасного надання корисної інформації. RFID використовують у логістиці, бібліотеках, автоматичних системах оплати, парках розваг тощо. Дані технології допомагають відстежувати та збирати потрібні дані, скорочують час очікування та трудові затрати для управління подіями.

Карта тепла і рухів. Використовуючи датчики BLE і WiFi, можна створювати карту тепла і рухів, на яку наноситься, наприклад, кількість відвідувачів торговельного центру відповідно до їх місцезнаходження. На основі цих даних менеджери можуть створювати так звані “гарячі зони”, в яких знаходиться акційні пропозиції тощо. Це створює нові шляхи розвитку для маркетингу [5].

Розумне харчування. IP – це мережа повсякденних об’єктів, включаючи навіть ті пристрої, які використовуються для приготування їжі. IP реєструє всі події використання харчових продуктів, створює кількісні дані про них. Автоматичні системи копіювання, що використовують технологію RFID, можуть містити інформацію про продукти, які було придбано. На основі цієї інформації формується припущення, які продукти потрібно буде купити в майбутньому. Це дозволяє мінімізувати фінансову частину, яка витрачається в процесі [5].

Інтелектуальне освітлення. Система освітлення, яка працює з врахуванням погодних умов: система управління самостійно визначатиме час увімкнення освітлення ввечері чи під час появи туману. Проект “розумного освітлення” передбачає оснащення ламп датчиками, які реагуватимуть на рух людей та автомобілів. Також може мати можливість контролювати всіма системами освітлення “розумного будинку” бездротовим шляхом через власний мобільний телефон.

Розумні будинки. Дозволяють керувати бездротовим способом будь-якими системами будинку, які з’єднані в одну мережу. Однією із функцій “розумного будинку” є контроль параметрів навколишнього середовища, залежно від чого здійснюється регулювання температури в приміщеннях. Це означає, що є можливість контролю ефективного споживання енергоресурсів.

Інтерактивні плакати. Являють собою пристрої, що використовують технологію NFC для надсилання сповіщень певним користувачам, які знаходяться біля них.

Висновки. Таким чином, сучасна концепція Інтернету речей містить в собі безліч пристроїв, які можуть самостійно виконувати обробку зібраних даних, проводити аналіз, здійснювати обмін накопиченою інформацією між собою і залежно від результатів приймати рішення та виконувати певні дії, які в свою чергу можуть використовуватися не тільки в бізнесі або маркетингу для підвищення ефективності та якості роботи, зменшення фінансових та матеріальних затрат, а й для полегшення повсякденного життя.

Література. 1. Internet of things [Online]. Available: https://en.wikipedia.org/wiki/Internet_of_things. 2. Steffen Schenkluhn. (2014, September 8). Device and event management as foundation of IoT applications [Online]. Available: <https://blog.bosch-si.com/internetofthings/>. 3. iBeacon [Online]. Available: <https://en.wikipedia.org/wiki/iBeacon>. 4. Radio-frequency identification [Online]. Available: https://en.wikipedia.org/wiki/Radio-frequency_identification. 5. Eugene Yang. (2015, December 13). Integrating the IoT with Event Management [Online]. Available: <https://www.gevmc.com/blog/integrating-the-iot-with-event-management/>.



Рис. 1. Взаємодія пристроїв з датчиком iBeacon

Прогонов Д.О.

КПІ ім. Ігоря Сікорського, ФІОТ, Київ, Україна

Теоретико-інформаційні оцінки стійкості методів UNIWARD до стегоаналізу

Вступ. Важливою складовою захисту інформаційних ресурсів даних організацій є протидія витоку конфіденційної інформації, зокрема з використанням прихованих (стеганографічних) систем. Обмеженість апіорних даних щодо особливостей застосовуваних стеганографічних методів призводить до вагомого зниження ефективності відомих методів виявлення прихованих повідомлень (стеганограм). Тому актуальною задачею є розробка методів сліпого стегоаналізу, що дозволять забезпечити високу імовірність виявлення стеганограм навіть у випадку відсутності апіорних даних щодо використаних методів приховання повідомлень.

Огляд сучасних методів стегоаналізу. Одним з поширених підходів до виявлення стеганограм є побудова статистичної моделі файлу-контейнеру, зокрема цифрового зображення, та подальшого аналізу відмінностей у значеннях параметрів моделі для вихідного (незаповненого) контейнеру та стеганограми [1, 2]. Для підсилення слабких змін параметрів зображення-контейнеру (ЗК) при формуванні стеганограм згідно адаптивних стеганографічних методів застосовується метод калібрування контейнеру – оцінки статистичних параметрів незаповненого зображення-контейнеру [3]. При цьому виявлення стеганограм проводиться шляхом порівняння “віддаленості” розподілів значень яскравості пікселів вихідного (каліброваного) P_C та заповненого P_S зображень-контейнерів.

Теоретико-інформаційні оцінки стійкості стеганографічних методів. Для оцінки “віддаленості” розподілів P_C та P_S широко використовується відстань Кульбака-Лейблера D_{KL} [1]. Для підсилення слабких відмінностей у розподілах P_C та P_S при використанні новітніх UNIWARD-подібних методів [4], зокрема J-UNIWARD та SI-UNIWARD, в роботі запропоновано використовувати також спеціалізовані оцінки відстаней між імовірнісними розподілами, зокрема відстані Бхаттачарая D_B , Хеллінгера D_H , χ^2 -відстань та спектр відстаней Реньї D_R .

Отримані результати. Дослідження ефективності застосування запропонованих оцінок відстаней між розподілами P_C та P_S проводилося з використанням вибірки 10,000 зображень зі стандартного тестового пакету MIRFlickr-25k [5]. Тестові зображення-контейнери були масштабовані до однакового розміру $640 \cdot 480$ (пікселів) та представлені в градаціях сірого кольору. За результатами аналізу отриманих даних були побудовані залежності значень відстаней Кульбака-Лейблера D_{KL} , Бхаттачарая D_B , Хеллінгера D_H , χ^2 -відстані та спектру відстаней Реньї D_R від ступеня заповнення зображення-контейнеру стегоданими.

Висновки. За результатами проведеного дослідження виявлено, що застосування спектру відстаней Реньї D_R дозволяє підсилити слабкі відмінності в розподілі значень яскравості пікселів незаповненого зображення-контейнеру P_C і стеганограми P_S . Це дає можливість підвищити точність виявлення стеганограм сформованих згідно UNIWARD-подібних методів приховання повідомлень навіть у випадку слабого заповнення зображення-контейнеру стегоданими (менше 10%).

Література. 1. Fridrich J. Steganography in Digital Media: Principles, Algorithms, and Applications / Fridrich J. – 1st Edition. – New York: Cambridge University Press, 2009 – p. 437. 2. Progonov D. Multiclass detector for modern steganographic methods / Progonov D. // International Journal “Information Theories and Applications”. – Vol. 24, No. 3, 2017. – pp. 55–71. 3. Kodovsky J., Fridrich J. Calibration revisited / Kodovsky J., Fridrich J. // Proceeding of the 11th ACM workshop on multimedia and security. – Princeton, 2009. p. 63–74. 4. Holub V., Fridrich J. Universal distortion function for steganography in an arbitrary domain / Holub V., Fridrich J. // EURASIP Journal on Information Security. – 2014. 5. Huiskes M.J., Lew M.S. The MIR Flickr Retrieval Evaluation / Huiskes M.J., Lew M.S. // ACM International Conference on Multimedia Information Retrieval.

Продан А.О.

КПІ ім. Ігоря Сікорського, ФПМ, Київ, Україна

Оптимізація пошуку та індексації документів формату XML

На сьогодні спостерігається безперервне зростання кількості даних в електронній формі. Для вирішення проблеми систематизації та індексації даних було розроблено велику кількість структурованих форматів файлів. Одним з найбільш вживаних форматів на сьогодні є формат XML. В зв'язку з цим постає проблема покращення ефективності пошуку по документам формату XML великого обсягу.

У цій роботі розглядаються проблеми індексації наборів даних XML і знаходження ефективних методів пошуку для індексованих даних. Також визначаються структури даних та алгоритми, які використовують паралельний підхід до проблеми пошуку. Реалізація представляє користь для багатоядерних процесорів.

Декомпозиція роботи залежить від вибраної моделі потоків. Найкращим рішенням є створення одного потоку для кожного ядра логічного процесора, за рахунок якого повністю використовується апаратний потенціал.

Робота розподіляється на невеликі фрагменти, що називаються завданнями. Кожне завдання насправді є функторним об'єктом, що має посилання на код, який буде виконуватися, і дані, що обробляються кодом. Завдання є неподільними з логічної точки зору, однак вони можуть викликати інші завдання. Кожен потік містить свій власний пул, де зберігаються завдання, що готові до запуску. Коли останнє посилання, що вказує на завдання, видаляється, завдання позначається як готове до виконання та додається в готовий пул потоку, який створив завдання. Готові завдання підбираються один за одним і виконуються відповідними потоками.

Підхід до обробки запитів передбачає, що документ XML повністю завантажується в основну пам'ять для прискорення виконання операції пошуку за заданим шляхом XPath. Коли XML-документ завантажується в пам'ять, він повинен бути представлений належним чином. WWW Consortium розробив специфікацію DOM, яка описує об'єктну модель для даних XML. Ця модель насправді є деревом, де вузли представляють елементи XML, атрибути та текстові дані. Так як модель DOM є досить складною, в роботі використовується аналогічна структура у вигляді дерева, яка є більш простою, але не повністю сумісною з DOM.

Структура DOM може бути проіндексована багатьма способами. Деякі представлені індекси є досить складними і більш придатні для систем управління базами даних. В роботі використовуються лише ті, які досить прості для реалізації у автономних програмах і які мають суттєвий вплив на паралельне виконання XPath.

Оскільки ми розглядаємо обробку великих документів (понад 10^5 вузлів) пріоритетом, нам потрібен алгоритм, який дозволяє виконувати запити в просторі $O(kN)$, де k – розмір запиту. Тому наша реалізація базується на рекурсивному підході з двома покращеннями.

Мінімальні контекстні правила застосовуються для кожного підзапиту. Якщо підзапит не потребує повної інформації контексту для виконання, можлива оптимізація підзапиту. У деяких випадках можуть бути використані ітераційні методи замість рекурсивних, в той час як лінійна просторова складність залишається сталою. Таким чином створюється швидше рішення, яке можна застосовувати навіть для великих документів.

Друге покращення полягає в оптимізації використання пам'яті. Якщо кількість доступної пам'яті обмежена, то експоненціальна часова складність для деяких запитів не може бути уникнута. Однак, її можна значно зменшити, використовуючи кешування для підзапитів.

Використання розробленого рекурсивного алгоритму з наведеними покращеннями дозволяє виконувати декомпозицію задач на під задачі, що можуть виконуватись паралельно, що дозволяє збільшити швидкість виконання запитів XPath.

Література. 1. Reinders. (2007). Intel threading building blocks. O'Reilly & Associates, Inc., Sebastopol, CA, USA. 2. Cooper B.F., Sample N., Franklin M.J., Hjaltason G.R., Shadmon M. (2001). A fast index for semistructured data. Morgan Kaufmann.

Романкевич В.О., Липка Т.Б.

КПІ ім. Ігоря Сікорського, ФПМ, Київ, Україна

Покращення якісних характеристик стеганосистем з використанням зображень в якості контейнера

Вступ. Стеганографія – давня і міждисциплінарна наука, що активно розвивається. Зокрема, починаючи приблизно з 2008 року нею стали цікавитися не тільки математики-криптографи, а й лінгвісти, філологи і навіть хіміки [1]. На даний момент в США зареєстровано 119 патентів по стеганографії [2], в Росії – 63 [3], в Україні – 12 [4].

Класичною ціллю стеганографії є прихована передача даних. Вона полягає в передачі інформації таким чином, щоб противник не здогадався про сам факт існування прихованого повідомлення. Окрім цього, її цілями також є: цифрові підписи (Digital Fingerprint), стеганографічні водяні знаки (Stego Watermarking) [2].

Актуальні проблеми. Аналіз останніх досліджень і публікацій [5–7] показав, що найбільшу популярність в комп'ютерній стеганографії здобули методи, що використовують у ролі контейнера зображення.

Проведений аналіз також показав, що однією з ключових проблем стеганографії є питання оптимального використання стего-контейнера, оскільки потреба у передачі даних великого обсягу виникає частіше, ніж потреба у передачі даних невеликого обсягу [6].

При роботі з даними великого обсягу важливим завданням є прискорення процесів шифрування, дешифрування інформації.

Часто сучасні стеганографічні методи використовують великі ключі, що змушує користувачів зберігати їх у вигляді файлів. В такому випадку існує небезпека у викраденні цих файлів. Тому важливим завданням є приведення ключа до такого вигляду, щоб його можна було просто пам'ятати.

Критеріями оцінки стеганосистем є їхні якісні (пропускна здатність, стійкість, невидимість, захищеність, складність вбудовування і вилучення секретного повідомлення) та кількісні (співвідношення «сигнал/шум» (SNR), якість зображення (IF)) показники [8]. Отже, актуальним є питання покращення цих показників.

Вищевказані критерії є взаємно конкуруючими і не можуть бути оптимальними одночасно. Наприклад, якщо необхідно приховати велике повідомлення, то неможливо вимагати абсолютної невидимості і високої стійкості. Завжди необхідний оптимальний компроміс [8].

Мета. Метою даної роботи є покращення деяких якісних показників методу стеганографії з використанням матриці sudoku. Показниками, що покращуються є:

1. **Пропускна здатність** – кількість бітів секретного повідомлення, які можуть бути передані за допомогою даного методу в стего-контейнері фіксованого розміру.
2. **Складність вбудовування і вилучення** – кількість стандартних операцій, які необхідно виконати для вбудовування і виявлення секретного повідомлення.

Метод стеганографії з використанням матриці sudoku. Класична реалізація даного методу передбачає, що контейнер та секретне повідомлення – це зображення у форматі bmp, а ключ – матриця sudoku. Метод полягає у модифікації виділених пар пікселів стего-контейнера базуючись на цій матриці [9].

Даний метод дозволяє закодувати 1 піксель (3 байти) секретного зображення за допомогою 6-ти пікселів (18 байт) контейнера.

Матриця sudoku. Квадратна матриця $N \times N$ є sudoku матрицею, якщо виконуються наступні умови:

- Для заповнення матриці використовується лише заданий набір з N цифр,
- Цифри в межах кожного з рядків унікальні,
- Цифри в межах кожного з стовпчиків унікальні,
- Цифри в межах кожного з районів унікальні (Матриця послідовно ділиться на райони – квадрати зі сторонами $\sqrt{N}$)

Завдяки останньому правилу побудови sudoku, така стеганосистема характеризується високою невидимістю (характеристика, що відповідає за неспроможність людського зору виявити стеганографічне повідомлення без використання спеціальних засобів).

Модифікація методу стеганографії з використанням матриці sudoku. Для покращення якісних показників стеганосистем у даній роботі пропонується наступне:

1. Збільшення розміру ключа – замість стандартної матриці sudoku 9×9 використовувати 256×256 та не виконувати її дублювання.
2. Відмова від використання дев'яткової системи числення. Завдяки цьому зникає необхідність виконувати перетворення чисел з десятикової в дев'яткову систему числення та навпаки.
3. Перехід до оперування повними значеннями складових кольору (значення R, G, або B), а не по цифрах, як в існуючих методах.

Завдяки цьому модифікований метод дозволяє закодувати 1 піксель (3 байти) секретного зображення за допомогою 2-х пікселів (6 байт) контейнера.

Висновки. Приховування інформації – актуальна тема на протязі всього часу існування людства. Поширення мережі Інтернет призводить до збільшення обсягів інформації, що передається, обробляється та зберігається. А це створює підґрунтя для активного використання стеганографії.

Аналіз досліджень і публікацій показав, що найпопулярнішими методами стеганографії є ті, що використовують у ролі контейнера зображення. Тому було вивчено і запропоновано модифікацію одного з них – методу стеганографії з використанням матриці sudoku.

Модифікований метод відрізняється від існуючих покращенням якісних показників стеганосистеми – пропускну здатності та складності вбудовування і вилучення секретного повідомлення.

Місткість стего-контейнера збільшується у 3 рази. А зникнення необхідності у виконанні перетворень між десятиковою та дев'ятковою системами числення і оперування повними складовими кольорів дозволяє прискорити процеси шифрування і дешифрування даних.

Література. 1. <https://habrahabr.ru/post/253045/>. 2. http://patents.com/search?top_keyword=steganography&keyword=steganography. 3. <http://www.freepatent.ru/>. 4. <http://uapatents.com/>. 5. Рябко Б.Я., Фионов А.Н. Основы современной криптографии и стеганографии. 2-е изд. / Рябко Б.Я., Фионов А.Н. // М.: Горячая линия – Телеком, 2013. 232 с. 6. В.Г. Бабенко, В.М. Зажома, О.Б. Нестеренко. Метод вбудовування стегоповідомлення на основі ключового елемента / В.Г. Бабенко, В.М. Зажома, О.Б. Нестеренко. // Захист інформації. 2014. – С. 53–58. 7. Кузнецов О.О. Стеганография: навчальний посібник / О.О.Кузнецов, С.П. Євсєєв, О.Г. Король. // – Х. : Вид. ХНЕУ, 2011. – 232с. 8. Вовк О.О. Методи підвищення стійкості та пропускну здатності систем прихованої передачі інформації http://nure.ua/wp-content/uploads/dis_Vovk.pdf. 9. Sanmitra I., Shivananda P., Shrikant B., Usha B, Image Steganography using Sudoku Puzzle for Secured Data Transmission // International Journal of Computer Applications (0975 – 888) Volume 48– No.17, June 2012 <https://pdfs.semanticscholar.org/17bf/1c8c04a5ae0da1cb644908fe4126d6d2f17f.pdf>.

Романов В.В., Шахун О.В.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Інформаційні технології в освіті

В останні роки швидкий і динамічний розвиток мережі Інтернет і технологій побудови повноцінних інформаційних систем, послужив поштовхом до формування нової предметної області «Інформаційні технології в освіті». До неї можна віднести проблеми відкритої освіти та дистанційного навчання, проблеми інтелектуальних навчальних систем, машинного навчання, а також інформаційних освітніх середовищ. При цьому відбувається посилення інтелектуальних можливостей учнів в інформаційному суспільстві, а також інтенсифікація процесу навчання та підвищення якості освіти.

Щоб розробити нові та нестандартні підходи, ефективність яких була б доведена експериментальним шляхом, перш за все слід розглянути класифікацію навчальних програм в цілому і основні креативні напрямки в навчанні. Найчастіше виділяють чотири типи навчальних програм: наставницькі, тренувальні та контролюючі, імітаційні і моделюючі та цілий ряд нових інформаційних технологій. Компанія Gartner використовує фази для опису життєвого циклу нових технологій, як показано на рис. 1.

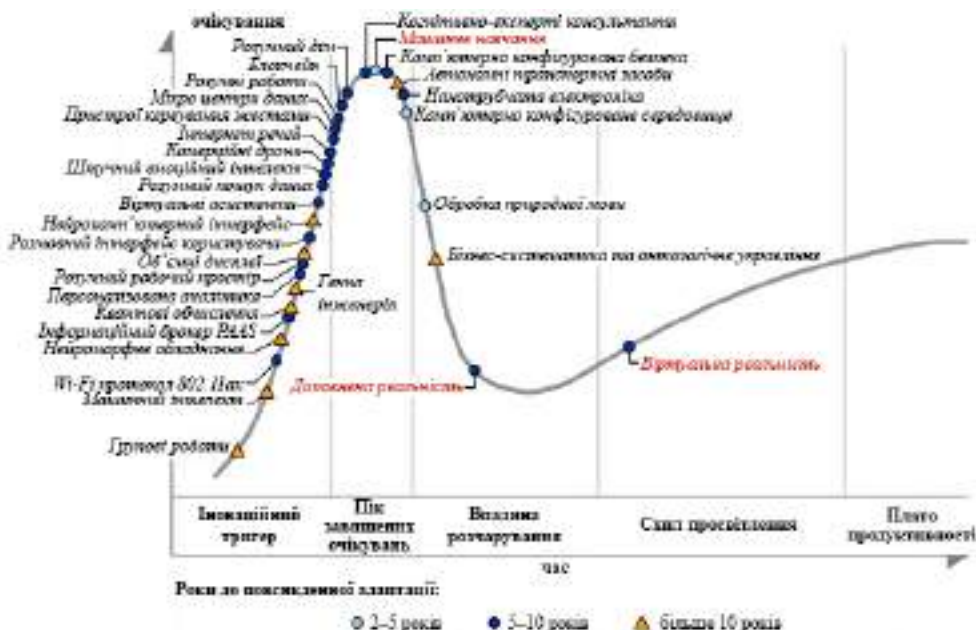


Рис. 1. Життєвий цикл інформаційних технологій

До абсолютно нових напрямків в навчанні можна віднести три напрямки: машинне навчання, як підрозділ штучного інтелекту, освітні ігри з використанням доповненої реальності (англ. Augmented Reality, AR) і віртуальної реальності (англ. Virtual Reality, VR), які вимагають великої роботи психологів, фахівців в області вивчаемого предмета, програмістів, педагогів-методистів та віртуальних середовищ. Сьогодні вже існують організації, що займаються розробкою подібних, орієнтованих на віртуальну реальність інформаційних систем. Зупинимось на цих напрямках.

Машинне навчання. Математична дисципліна, яка використовує: розділи математичної статистики, чисельні методи оптимізації, теорію ймовірності, дискретний аналіз та виділяє знання з даних. За оцінкою McKinsey Global Institute, в 2018 році в одних тільки США попит на експертів по машинному навчанню буде перевищувати пропозицію на 170–190 тисяч чоловік. Крім того, виникне потреба більш ніж в півтора мільйонах управлінців з навичками

аналізу даних [1]. Розрізняють дедуктивний і індуктивний методи навчання в комп'ютерних системах. Дедуктивне, або аналітичне, навчання (експертні системи) — це наявні знання, сформульовані експертом і якимось формалізовані. Програма виводить з цих правил конкретні факти і нові правила. Індуктивне навчання (статистичне навчання) — це програма, яка на основі емпіричних даних будує загальне правило. Емпіричні дані можуть бути отримані самою програмою в попередні сеанси її роботи або просто завантажені. Етапи вирішення завдань машинного навчання: аналіз завдання і даних, препроцесинг даних і пошук ознак, побудова моделі, зведення навчання до оптимізації, рішення проблем оптимізації і перенавчання, оцінювання якості, впровадження та експлуатація [1]. Основна схема навчання: Множина X — об'єкти, приклади, зразки (samples). Множина Y — відповіді, відгуки, «мітки», класи (responses). Існує деяка залежність $g : X \rightarrow Y$, що дозволяє по $x \in X$ передбачити (або оцінити ймовірність появи) $y \in Y$. Залежність відома тільки на об'єктах з навчальної вибірки: $T = (x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$. Пара $(x_i, y_i) \in X \times Y$ — прецедент. Задача навчання по прецедентах: навчитися по нових об'єктах $x \in X$ передбачати відповіді $y \in Y$. Вирішити завдання машинного навчання означає розробити алгоритм або модель алгоритму, який би залежав від параметрів і дозволяв визначити значення мітки класу (Y) для нового об'єкта (X).

Використання віртуальної реальності. Віртуальна реальність — один з головних символів сучасності, яка поступово поширюється за межі розважальної індустрії, щоб зробити крок до чогось більшого. Інформаційні аналітики прогнозують, що до 2019 року ринок VR виросте до 15,9 млрд. дол. [2]. Віртуальна реальність створюється технічними засобами — це візуальний світ, який передається людині через його відчуття: зір, слух, нюх, дотик. VR імітує вплив і реакцію на нього. Вона намагається бути максимально схожою на фізичний світ. Сприйняття об'єктів в ній наближене до звичайної реальності настільки, наскільки це можливо з нинішнім рівнем розвитку технологій. Користувач може своїми діями впливати на об'єкти в відповідності до реальних законів фізики. Особливість VR в тому, що вона дозволяє порушувати і змінювати ці закони та робити все, на що вистачає фантазії. Доповнена реальність (англ. Augmented Reality, AR) змінює поточне сприйняття реального середовища, тоді як віртуальна реальність замінює навколишнє середовище його імітацією.

Плюси використання VR в освіті. Використання віртуальної реальності відкриває багато нових можливостей в навчанні та освіті. Можна виділити п'ять основних переваг застосування VR технологій. Наглядність — використовуючи 3D-графіку, можна деталізовано показати фізичні процеси аж до атомного рівня. Віртуальна реальність здатна не лише надати відомості про саме явище, а й продемонструвати його з будь-яким рівнем деталізації. Безпека — можливість занурити глядача в будь-яку “екстремальну ситуацію”. Залучення — віртуальна реальність дозволяє впливати на хід експерименту або вирішувати математичні задачі в ігровій та доступній формі. Зосередженість — дозволяє повністю зосередитись на матеріалі. Віртуальні уроки — вид від першої особи і відчуття власної присутності в віртуальному світі.

Проникнення інформаційних технологій в сферу освіти сприяє вирішенню педагогічних і психологічних проблем, оскільки дозволяє педагогам гармонійно змінювати зміст, методи і організаційні форми навчання і тим самим усувати слабку вмотивованість студентів у їхньому навчанні. Сьогодні покоління 2000-них відчуває себе дуже комфортно отримуючи онлайн освіту, знаходячи інформацію в інтернеті, навчаючись дистанційно за допомогою відео курсів. Безсумнівно, машинне навчання і віртуальна реальність — це наступний крок, оскільки використання нових інформаційних технологій, які інтегрують новітні і традиційні форми навчання, є найбільш перспективними.

Література. 1. Домингос Е. Верховный алгоритм: как машинное обучение изменит наш мир / Домингос Е – Москва: Манн, Иванов и Фербер, 2016. – 336 с. 2. Фещенко А.В. Технологии виртуальной и дополненной реальности в образовательной среде вуза / Фещенко А.В., Бахарева В.А., Захарова У.С., Сербин В.А. // Открытое и дистанционное образование. – 2015. – №4(60). – С. 12–20.

Северін С.І., Тесленко О.К.

КПІ ім. Ігоря Сікорського, ФПМ, Київ, Україна

Вибір алгоритму ущільнення для багаторозрядних підстановок

Вступ. В умовах роботи великих архівів, обробки даних міського управління та бібліотечних фондів, зростає необхідність у збереженні великих обсягів даних. Тому підвищення ефективності ущільнення даних є досить актуальною задачею. В [1] запропоновано метод ущільнення файлів, який поряд з використанням існуючих засобів ущільнення ґрунтується на проміжному перетворенні файлів за допомогою підстановок довільної розрядності. Для такого перетворення використані результати реалізації підстановок на регулярних одновимірних каскадах конструктивних модулів (ОККМ), де кожен із конструктивних модулів (КМ) є однорозрядним. В [2] експериментальним шляхом із 850 однорозрядних КМ визначені 96 пар прямих й обернених КМ, а також наведений приклад алгоритму перетворення файлів за допомогою регулярних ОККМ.

Згідно з [1], використання багаторозрядних підстановок на базі однорозрядних КМ дозволяє підвищити показники ущільнення файлів. Можна прогнозувати підвищення ефективності застосування КМ більшої розрядності, зокрема на 2, 4 та 8 розрядів. Визначення КМ довільної розрядності наведено в [4, 5].

Постановка задачі. В [3] запропоновано базовий алгоритм підвищення ефективності існуючих засобів ущільнення даних за допомогою прямих та обернених підстановок. Даний алгоритм передбачає перебір підстановок, що застосовуються до початкового файлу та заархівованого файлу, після чого знову виконується ущільнення. Найефективніша підстановка вибирається та записується у заголовок архіву. Однією з головних частин даного алгоритму є ущільнювач даних. Теоретично даний алгоритм можна застосувати з будь-яким ущільнювачем. При цьому існує проблема доцільності використання того чи іншого ущільнювача.

Основна частина. Важливими критеріями оцінки існуючих ущільнювачів є:

- швидкість ущільнення;
- швидкість розущільнення;
- коефіцієнт ущільнення;
- використані ресурси для ущільнення.

Враховуючи те, що запропонований в [1] алгоритм потребує багато часу на виконання пошуку ефективної підстановки, критерій швидкості ущільнення втрачає актуальність. Таким чином при виборі архіватора для даного алгоритму на перший план виходить коефіцієнт ущільнення та швидкість розущільнення. Отже, для алгоритму підвищення ефективності саме ці критерії стають найбільш значимими.

Для даної задачі пропонується використання алгоритму архівації LZMA. LZMA використовується в архіваторі 7zip. Даний алгоритм має наступні переваги:

- високий коефіцієнт ущільнення;
- невеликі вимоги до пам'яті для розущільнення даних;
- бібліотека з відкритим кодом (що дає можливість використання цієї бібліотеки для тестування).

Недоліком даного алгоритму є мала швидкість ущільнення.

Висновки. У даній роботі описані основні критерії вибору архіватора на якому має базуватися алгоритм підвищення ефективності існуючих засобів ущільнення даних. Головними критеріями є високий коефіцієнт ущільнення та швидкість розущільнення.

Відповідно до даних критеріїв запропоноване використання архіваторів, що базуються на алгоритмі ущільнення LZMA.

Література. 1. Тарасенко В.П., Тесленко О.К., Яновська О.Ю. Використання прямих та обернених підстановок довільної розрядності для підвищення ефективності існуючих засобів ущільнення даних // Тези доповідей Четвертої Міжнародної науково-практичної конференції «Методи та засоби кодування, захисту й ущільнення інформації». Україна, Вінниця, 23–25 квітня 2013 р. – С. 223–226. 2. Тесленко О.К., Невмержицький П.А. Загальна оцінка алгоритму шифрування, який ґрунтується на суперпозиції підстановок із найпростіших регулярних ОККМ // Збірник тез доповідей восьмої наукової конференції магістрантів та аспірантів «Прикладна математика та комп'ютинг» (ПМК-2016), Київ, НТУУ «КПІ», 20–22 квітня 2016 р. – С. 159–164. 3. Дробязко І.П., Северін С.І., Тесленко О.К. Способи використання генетичних алгоритмів і багаторозрядних підстановок для підвищення ефективності існуючих засобів ущільнення даних. // Тези доп. XVII міжнародної наукової конференції імені Т.А. Таран «Інтелектуальний аналіз інформації». Україна, м. Київ 17–19 травня 2017 р. – С. 71–74. 4. Тарасенко В.П., Тесленко О.К., Яновська О.Ю. Реалізація повних підстановок на простому двомодульному каскаді конструктивних модулів. Інформаційні технології та комп'ютерна інженерія. – К. : НТУУ «КПІ», 2008. – №1(11). – С. 88–97. 5. Тарасенко В.П., Тесленко О.К., Яновська О.Ю. Реалізація обернених підстановок на простому двомодульному каскаді конструктивних модулів // Інформаційні технології та комп'ютерна інженерія. – К. : НТУУ «КПІ», 2013. – №3 (28). – С. 16–25.

Сопілков М.Р.

КПІ ім. Ігоря Сікорського, ІПСА, Київ, Україна

Прогнозування виникнення збройних конфліктів з використанням ймовірісно-статистичних методів

Задача моделювання і прогнозування військових конфліктів у світі є надзвичайно актуальною. Як говориться у дослідженні Peace Research Institute Oslo (PRIO), виконаному спільно з Університетом м. Уппсала (Швеція), починаючи з 1946 року мали місце 256 збройних конфліктів [1]. Тобто за початкову точку для отримання та ведення статистики стосовно виникнення збройних конфліктів взято завершення Другої світової війни. Дослідники відзначають два головних тренди у процесах, пов'язаних з військовими конфліктами. Перший полягає у тому, що на початку формування статистичних даних найбільшу кількість становили міждержавні і колоніальні конфлікти. Згодом такі конфлікти зникли і на сьогодні домінуючою формою конфліктів є внутрішньодержавний, але часто із залученням учасників з інших держав. Другий тренд – збільшення кількості внутрішньодержавних громадянських конфліктів. Пік припадає на 1991 рік – 52 збройних конфлікти. Потім почалося зниження кількості конфліктів: так, у 2003 році налічувалося тільки 32 конфлікти. З 2003 року кількість збройних конфліктів знову почала зростати, а потім спадати – їх налічувалося від 30 до 40 на рік. Тобто очевидно, що світ змінюється. Він рухається шляхом демократизації та глобалізації (крім Росії). Однак, збройні конфлікти все ж виникають і їх причини в більшості випадків можна розділити на дві основні групи:

1. Ресурсні війни: якщо деяка країна має достатньо цінних ресурсів на своїй території, то агресивні країни-сусіди можуть застосувати різні методи для вилучення цих ресурсів на свою користь – у більшості випадків це збройний конфлікт.
2. Політичні війни: якщо країні вигідно підтримувати ведення збройних конфліктів, то вона їх реалізує. Наприклад це може мати місце, коли наявна влада хоче продовжити свій строк перебування на посадах. Або хоче відвести увагу власного народу від накопичених цією владою проблем у державі.

Внаслідок таких конфліктів деякі країни дійсно тимчасово перебувають у виграшній ситуації, але за це гинуть сотні тисяч людей: солдати на фронті і мирне населення. Крім того, внаслідок окупації території руйнується економіка, погіршується екологія та соціальні умови проживання (приклад – сучасна ситуація на окупованих територіях Донбасу). Аналіз поточної ситуації свідчить, що на сьогоднішній день загроза виникнення збройних конфліктів залишається великою.

В результаті огляду існуючих методів моделювання можливих військових конфліктів встановлено, що для цієї мети застосовують моделі у формі диференціальних рівнянь, теорію ігор, нечітку логіку та нейронечіткі моделі та сучасні методи імітаційного моделювання, які поєднують у собі практично всі вказані аналітичні моделі. У дослідженні поставлена задача опрацювання технологій та моделей для прогнозування ймовірності виникнення збройних конфліктів між країнами світу. Однією з таких технологій моделювання і прогнозування обрані байєсівські мережі. Застосування Байєсівських мереж має такі переваги: модель може містити велику кількість різномірних змінних (статистичні дані, експертні оцінки, категорійні дані та комбінації вказаних величин), структура моделі може бути оптимізована за допомогою альтернативних методів, для оцінювання ймовірнісного висновку існує множина методів з яких можна вибрати прийнятну процедуру для конкретного використання. Другою технологією обрані нейронні мережі, перевагами яких є: обробка великої кількості даних, можливості застосування альтернативних оптимізаційних методів навчання, можливості адаптації побудованих структур до нових даних, можна застосовувати альтернативні функції активації. В якості вхідних даних використовується база даних за 2013–2016 роки країн світу за їх основними економічними та соціальними показниками, які можуть впливати на прогнозне значення основної змінної (ВВП на душу населення, наявність ядерної зброї, рівень безробіття,

агресивність країни, індекс людського розвитку, чисельність армії, витрати на оборону на душу населення тощо). Перед використанням для побудови моделі деякі дані були кластеризовані, пронормовані та доповнені.

За допомогою отриманих даних та аналізу можливих структур великої кількості створених моделей-кандидатів спочатку побудована байєсівська мережа, яка давала найкращі результати прогнозування ймовірності виникнення конфлікту. Країною для перевірки прогнозу була обрана Україна. За даними створеної моделі, ймовірність виникнення збройного конфлікту на території країни у 2013 році була 0,38, а у 2014 вже 0,51.

Далі на основі тих самих вхідних даних, була побудована та навчена нейронна мережа. Вона також призначена для прогнозування можливості виникнення збройного конфлікту на території країн світу. Таким чином, буде отримано 2 оцінки: ймовірнісна оцінка за байєсівською мережею і оцінка можливості виникнення конфлікту в інтервалі $[0; 1]$ за нейромережею. Отримані оцінки комбінуються з однаковими або різними ваговими коефіцієнтами з метою обчислення комбінованої оцінки. На основі попередніх досліджень встановлено, що якість комбінованих оцінок, як правило, перевищує якість оцінок отриманих за кожною моделлю окремо [2]. Для виконання обчислювальних експериментів побудований веб-додаток на мові програмування Java з використанням фреймворку Spring, а також Front-end фреймворку Angular JS, які забезпечують швидке виконання обчислювальних експериментів та зручне відображення даних клієнту відповідно.

Також в якості системи управління базами даних була обрана MS SQL Server 2012, для комфортної роботи з даними за межами веб-додатку (заповнення початкових даних), а також для автоматичного перерахунку агрегованих предикторів та оновлення даних в цілому.

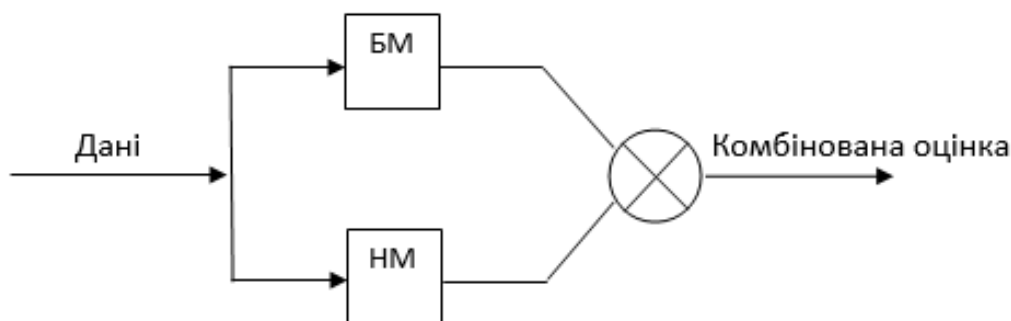


Рис. 1. Приклад загального розрахунку комбінованої оцінки

Висновки. В роботі виконано аналіз зібраних статистичних даних стосовно сучасного економічного та оборонного стану вибраних країн світу з Україною включно. Розроблена інтелектуальна система підтримки прийняття рішень, яка призначена для прогнозування виникнення військового конфлікту у конкретній країні. Отриманий прогноз стосовно України може інтерпретуватись таким чином: з кожним роком в Україні погіршувались економічні та соціальні показники її розвитку, а тому при прогнозуванні військового конфлікту на 2014 рік модель дала високу ймовірність його виникнення. У подальшому можна доповнити отриманий метод статистичною та регресійною обробкою вхідних даних, використанням нейронних мереж та нечіткої логіки з метою підвищення якості прогнозу. Комбіноване використання кількох методів моделювання і прогнозування, як правило, дає можливість підвищити якість оцінок прогнозів.

Література. 1. Scott Gates, Håvard Mogleiv Nygård, Håvard Strand, Henrik Urdal. Trends in Armed Conflict, 1946–2014 // Peace Research Institute Oslo (PRIO) – Oslo, Norway, 2016 – Р. 2. 2. Бідюк П.І., Меньяйленко О.С., Половцев О.В. Методи прогнозування. – Луганськ: Луганський Університет ім. Т.Г. Шевченка, 2008. – 600 с.

Authors · Авторы · Автори

- Aghaei Agh Ghamish Ovi Nafas, 18
 Aghaei Agh Ghamish Yaghoub, 18
 Akhmedova Daryna N., 26
 Baranovska Galyna G., 22
 Baranovska Lesia V., 20, 22
 Basyuk Taras M., 192
 Bobyr Anastasiia O., 107
 Boyarinova Yuliya E., 194
 Chernenko Vladyslav M., 204
 Chertov Oleg R., 106–108
 Donets Georgiy A., 18
 Dychko Alina O., 23
 Galib Hamidov, 120
 Globa Larisa S., 202
 Hasan Muhammad Jamal, 165
 Honcharevskyi Dmytro O., 108
 Kalinovskiy Yakov A., 194
 Kasyanchuk Igor V., 196
 Kharchenko Kostiantyn V., 118
 Khitsko Yana V., 194
 Kondratova Ljudmyla P., 24, 164
 Kopp Andrii M., 198
 Kurylenko Oleksandr M., 200
 Kyselov Gennadiy D., 30
 Kyselova Anna G., 30
 Lazuka Vitalii T., 109
 Makohon Roman O., 208
 Malysheva Milena O., 26
 Moroz Anastasiia M., 202
 Mustafa Rekar Qasim, 205
 Naderan Maryam, 111
 Nikolaiev Sergii S., 113
 Olefir Oleksandr S., 204
 Onanko Anastasiia M., 115
 Orlovskiy Dmytro L., 198
 Osaulenko Viacheslav L., 28
 Pechurin Nickolay K., 164
 Pechurin Sergey N., 164
 Prymak Ivan K., 26
 Quadir Safwan Husein, 167
 Rodionov Andrii M., 109
 Romanenko Yan V., 33
 Rudnyk Taras P., 106
 Sergiyenko Anatoliy M., 165, 167, 205
 Sergiyenko Pavlo A., 165
 Shanin Vladyslav O., 107
 Shaptala Roman V., 30
 Siryk Sergii V., 31
 Soloviev Vladimir N., 33
 Sukalo Alina S., 194
 Tarasenko-Klyatchenko Oxana V., 207
 Terentiev Oleksandr M., 208
 Tkachenko Dmytro A., 116, 118
 Tuparieva Valentyna A., 207
 Tymoshenko Yurii O., 113
 Ved Olena V., 35
 Vinogradov Juri N., 167
 Yermeev Igor S., 23
 Zaychenko Yuri P., 111, 120
 Акимов Вадим Сергійович, 121
 Арчвадзе Натела Нодариевна, 210
 Бакун Сабіна Антонівна, 211
 Бакурова Анна Володимирівна, 123
 Бахрушин Володимир Євгенович, 36
 Беседовский Михаил Олегович, 37
 Бідюк Петро Іванович, 56, 136, 152, 243
 Білушак Юрій Ігорович, 98
 Болдак Андрій Олександрович, 38, 41, 101
 Бояринова Юлія Євгенівна, 221
 Буценко Юрій Павлович, 124
 Васильєв Владимир Иванович, 39
 Васильєв Володимир Иванович, 213
 Видолоб Антон Валерійович, 215
 Виклюк Ярослав Ігорович, 216
 Вишталъ Дмитрий Михайлович, 39
 Вішталъ Дмитро Михайлович, 213
 Власов Максим Дмитрович, 41
 Волкова Виолетта Николаевна, 43

- Волонтир Олександр Олегович, 234
Волошин Олексій Федорович, 46
Галета Павло Олександрович, 169
Галушко Марія Олегівна, 125
Гіоргізова-Гай Вікторія Шавлівна, 217
Гнатенко Валерій Юрьевич, 219
Годна Анастасія Вікторівна, 47
Горелова Галина Викторовна, 49
Городецкий Виктор Георгиевич, 51
Городько Наталья Алексеевна, 221
Горох Олександр Сергійович, 38
Грачова Олександра Андріївна, 156
Гришин Игорь Юрьевич, 52
Грушенко Катерина Юріївна, 222
Гусак Олександр Вадимович, 126
Гусак Олена Михайлівна, 216
Данилов Валерій Яковлевич, 87
Денисенко Олександр Іванович, 53
Дмитренко Олег Олександрович, 127
Довганюк Анна Олеговна, 186
Доманецька Ірина Миколаївна, 129
Дорундяк Ксенія Миколаївна, 155
Духновская Ксения Константиновна, 54
Ерохин Глеб Вадимович, 130
Євграфова Ксенія Леонідівна, 56
Єфремов Костянтин Вікторович, 38
Жданова Олена Григорівна, 47, 58, 131
Забелин Станислав Игоревич, 134
Заводник Вячеслав Владленович, 61
Зайченко Анастасія Євгеніївна, 63
Зайченко Елена Юрьевна, 135
Зайченко Юрій Петрович, 143, 157
Замятін Денис Станіславович, 223, 224
Згуровський Михайло Захарович, 10
Зинченко Анастасія Александровна, 247
Івченко Дмитро Анатолійович, 171
Капшук Олег Олексійович, 225
Караюз Ірина Валентинівна, 136
Касумова Тахмина Агаджафар, 65
Киричек Галина Григорівна, 226
Кирюша Богдан Анатолійович, 175
Кінда Віталій Васильович, 228
Кірік Олена Євстафіївна, 67
Кісельова Олена Михайлівна, 68
Кісільчук Богдан Ярославович, 230
Клімова Олена Віталіївна, 138
Книш Петро Костянтинів, 223
Коваленко Сергей Владимирович, 75
Ковтун Оксана Івановна, 54
Колінько Анжела Михайлівна, 175
Комлев Олександр Олександрович, 69
Кондратенко Наталія Романівна, 139
Коробова Марина Віталіївна, 74
Косташ Татяна Владимировна, 61
Костенко Святослав Владимирович, 240
Кравчук Євгеній Сергійович, 70
Крамов Артем Андрійович, 72
Краснюк Роман Петрович, 93
Краштан Таміла Євгенівна, 232
Кривий Сергій Лук'янович, 11
Кривоніс Артем В'ячеславович, 234
Круш Ігор Володимирович, 177
Кудін Григорій Іванович, 74
Кузнецова Юлія Анатоліївна, 235
Кузнецова Наталія Володимирівна, 141
Кулик Володимир Васильович, 74
Курай Вячеслав Іванович, 226
Курасов Олексій Ігорович, 245
Куценко александр Сергеевич, 75
Лабжинський Володимир Анатолійович, 124
Лабуткина Татяна Викторовна, 76
Ланде Дмитро Володимирович, 127
Левчук Святослав Богданович, 241
Леонова Алла Євгенівна, 43
Лимар Борис Олегович, 237
Липка Тетяна Богданівна, 260
Лисецкий Юрий Михайлович, 238
Литвинов Валерій Андроникович, 240
Литвинюк Антон Андрійович, 241
Логвіненко Владислава Олексіївна, 36
Логинова Олександра Викторовна, 43
Логін Вадим Вікторович, 243
Лозінський Анатоль Павлович, 179
Луданов Денис Костянтинів, 143
Любашенко Наталія Дмитрівна, 213
Любашенко Наталья Дмитриевна, 39
Лютенко Ірина Вікторівна, 245
Майстренко Светлана Яковлевна, 240
Максим Катерина Євгенівна, 246
Маленко Анастасія Олексіївна, 58
Маляр Микола Миколайович, 46
Манілевич Дмитро Федорович, 86
Маняк Юрій Вікторович, 83
Матвиенко Юлія Александровна, 247
Михалько Віталій Геннадійович, 177
Мілявський Юрій Леонідович, 78

- Нагулов Сергій Васильович, 123
Назарага Інна Михайлівна, 79
Назаренко Сергій Юрійович, 150
Назарова Ірина Акопівна, 181
Нестеренко Владислав Ігоревич, 80
Ови Нафас Агаи аг Гамиш, 135
Огурцов Максим Ігорович, 249
Олефір Олександр Степанович, 237
Орехов Олександр Арсенійович, 183
Орехова Наталія Артурівна, 183
Орлова Марія Миколаївна, 251
Осадчук Николай Павлович, 51
Остапчук Ян Михайлович, 184
Павлов Дмитро Геннадійович, 144
Падалко Вадим Геннадійович, 68
Панкратов Владимир Андреевич, 81
Панкратова Наталія Дмитрівна, 83, 85
Пашаева Малахат Мухтар, 252
Пашинська Наталія Миколаївна, 150
Перевертайло Андрій Ігорович, 224
Перестюк Марія Миколаївна, 10
Петрашенко Андрій Васильович, 222
Петрішенко Сергій Олександрович, 256
Пишнограєв Іван Олександрович, 38
Поворознюк Анатолій Іванович, 146
Погорілий Сергій Дем'янович, 11, 159
Подковаліхіна Олена Олександрівна, 36
Пономаренко Роман Николаевич, 148
Попенко Володимир Дмитрович, 131
Притоманова Ольга Михайлівна, 68
Прогонов Дмитро Олександрович, 258
Продан Анастасія Олегівна, 259
Путренко Віктор Валентинович, 150
Пховелишвили Мераб Гайозович, 210
Роговий Андрій Владиславович, 152
Роїк Павло Дмитрович, 96
Романкевич Віталій Олексійович, 260
Романкевич Олексій Михайлович, 86
Романов Валерій Володимирович, 262
Савченко Данил Олександрович, 154
Савченко Ілля Олександрович, 63, 100
Сапсай Тат'яна Григорьевна, 186
Северін Сергій Іванович, 264
Селін Юрій Миколайович, 160
Семендяк Євген Сергійович, 126
Сергеев-Горчинський Олексій
Олександрович, 70
Ситников Валерій Степанович, 219
Слинько Максим Сергійович, 187
Сльота Максим Русланович, 85
Слюсар Андрій Вячеславович, 87
Снігур Ольга Олексіївна, 139
Соколов Дмитро Віталійович, 245
Сопілков Максим Романович, 266
Сперкач Майя Олегівна, 47, 58
Ступень Павел Вячеславович, 219
Терентьев Олександр Миколайович, 211
Терещенко Еліна Валентинівна, 123
Тесленко Олександр Кирилович, 230, 264
Тимиргалеева Рена Ринатовна, 52
Тимофієва Надія Костянтинівна, 89
Тимошук Оксана Леонідівна, 155
Товажнянский Владимир Ігоревич, 75
Фатенко Владислав Васильович, 91
Филатова Анна Евгеньевна, 146
Фіногенов Олексій Дмитрович, 156
Фомін Олександр Володимирович, 141
Хроленко Ярослав Олексійович, 129
Хурцилава Константин Викторович, 240
Цегелик Григорій Григорович, 93
Цешевський Валерій Валентинович, 235
Циганок Віталій Володимирович, 96
Чапалюк Богдан Володимирович, 157
Чернуха Ольга Юрївна, 98
Черный Юрий Юрьевич, 43
Чертов Олег Романович, 144
Чечельницкая Александра Борисовна, 158
Чечула Микита Володимирович, 159
Чикрий Аркадій Алексеевич, 13
Чучвара Анастасія Євгенівна, 98
Шакун Олександр Васильович, 262
Шеренковський Артем Олегович, 217
Шехна Халед, 146
Шибирин Анастасія Романівна, 100
Широков Владимир Анатолиевич, 14
Шитий Дмитро Вадимович, 251
Шудренко Євгеній Олегович, 101
Шулькевич Тетяна Вікторівна, 160