

РОССИЙСКАЯ НЕЙРОСЕТЕВАЯ АССОЦИАЦИЯ
РОССИЙСКАЯ АКАДЕМИЯ НАУК
МИНИСТЕРСТВО НАУКИ И ВЫСШЕГО ОБРАЗОВАНИЯ РОССИЙСКОЙ ФЕДЕРАЦИИ
МОСКОВСКИЙ ФИЗИКО-ТЕХНИЧЕСКИЙ ИНСТИТУТ
(НАЦИОНАЛЬНЫЙ ИССЛЕДОВАТЕЛЬСКИЙ УНИВЕРСИТЕТ)
ФЕДЕРАЛЬНЫЙ НАУЧНЫЙ ЦЕНТР НАУЧНО-ИССЛЕДОВАТЕЛЬСКИЙ ИНСТИТУТ СИСТЕМНЫХ
ИССЛЕДОВАНИЙ РОССИЙСКОЙ АКАДЕМИИ НАУК
ARTIFICIAL INTELLIGENCE RESEARCH INSTITUTE

**XXVI МЕЖДУНАРОДНАЯ
НАУЧНО-ТЕХНИЧЕСКАЯ
КОНФЕРЕНЦИЯ**

НЕЙРОИНФОРМАТИКА-2024

СБОРНИК НАУЧНЫХ ТРУДОВ

- НЕЙРОИНФОРМАТИКА И ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ
- ТЕОРИЯ НЕЙРОННЫХ СЕТЕЙ
- ВЫЧИСЛИТЕЛЬНЫЕ ИССЛЕДОВАНИЯ ПРИНЦИПОВ И МЕХАНИЗМОВ РАБОТЫ ЕСТЕСТВЕННЫХ НЕЙРОННЫХ СИСТЕМ
- НЕЙРОСЕТЕВЫЕ ПАРАДИГМЫ И АРХИТЕКТУРЫ
- НЕЙРОКОМПЬЮТЕРЫ, НЕЙРОМОРФНЫЕ ВЫЧИСЛЕНИЯ И ИМПУЛЬСНЫЕ НЕЙРОННЫЕ СЕТИ
- ГЛУБОКОЕ ОБУЧЕНИЕ
- ОБУЧЕНИЕ С ПОДКРЕПЛЕНИЕМ В НЕЙРОБИОЛОГИИ И В СИСТЕМАХ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА
- ПРИКЛАДНЫЕ НЕЙРОСЕТЕВЫЕ СИСТЕМЫ
- ИНТЕРФЕЙС "МОЗГ-КОМПЬЮТЕР". НЕЙРОТЕХНОЛОГИИ
- НЕЙРОННЫЕ СЕТИ И КОГНИТИВНЫЕ НАУКИ. АДАПТИВНОЕ ПОВЕДЕНИЕ И ЭВОЛЮЦИОННОЕ МОДЕЛИРОВАНИЕ

МОСКВА
МФТИ
2024

УДК 001(06)+004.032.26(06)

ББК 72я5+32.818я5

М 43

**XXVI Международная научно-техническая конференция
"НЕЙРОИНФОРМАТИКА-2024" : сборник научных трудов. –
Москва : МФТИ, Физтех. 2024. – 258 с. : ил.
ISBN 978-5-7417-0797-5**

Представлены доклады XXVI Международной научно-технической конференции «НЕЙРОИНФОРМАТИКА-2024», проходившей 21–25 октября 2024 года в Московском физико-техническом институте (национальном исследовательском университете) (МФТИ, Физтех).

Тематика конференции охватывает широкий круг вопросов, посвящённых теоретическим и прикладным исследованиям в следующих областях: теория нейронных сетей, глубокое обучение, нейросетевые парадигмы и архитектуры, нейроморфные вычисления и импульсные нейронные сети, нейронные сети в когнитивных науках, обучение с подкреплением, адаптивное поведение и эволюционное моделирование, вычислительные исследования принципов и механизмов работы естественных нейронных систем.

Конференция проводилась при поддержке программы «Приоритет 2030».

УДК 001(06)+004.032.26(06)

ББК 72я5+32.818я5

ISBN 978-5-7417-0797-5

© Федеральное государственное автономное образовательное учреждение высшего образования «Московский физико-технический институт (национальный исследовательский университет)», 2024

ОРГАНИЗАТОРЫ КОНФЕРЕНЦИИ

- Московский физико-технический институт (национальный исследовательский университет) (МФТИ, Физтех)
- Научно-исследовательский институт системных исследований РАН (НИИСИ РАН)
- Национальный исследовательский ядерный университет «МИФИ» (НИЯУ МИФИ)
- Министерство науки и высшего образования Российской Федерации
- НИЦ «Курчатовский институт»
- Государственная корпорация по атомной энергии «Росатом»
- Московский авиационный институт (национальный исследовательский университет) (МАИ)
- Российская нейросетевая ассоциация
- Институт искусственного интеллекта AIRI

ПРЕДСЕДАТЕЛЬ КОНФЕРЕНЦИИ

Чл.-корр. РАН Крыжановский Борис Владимирович – председатель координационного совета Российской нейросетевой ассоциации, НИИСИ РАН, Москва

ОРГАНИЗАЦИОННЫЙ КОМИТЕТ КОНФЕРЕНЦИИ

Председатель – Райгородский Андрей Михайлович – МФТИ, Физтех, Москва

Зам. председателя – Юдин Дмитрий Александрович, МФТИ, Физтех, AIRI, Москва

Зам. председателя – Климов Валентин Вячеславович, НИЯУ МИФИ, Москва

Акопов Эдмунд Иванович – НИИСИ РАН, Москва

Карандашев Яков Михайлович – НИИСИ РАН, Москва

Киселев Михаил Витальевич – ЧГУ им. И.Н. Ульянова

Смирнитская Ирина Аркадьевна – НИИСИ РАН, Москва

Сохова Зарема Борисовна – НИИСИ РАН, Москва

Трофимов Александр Геннадьевич – НИЯУ МИФИ, Москва

Тюменцев Юрий Владимирович – МАИ, Москва

Ученый секретарь – Бесхлебнова Галина Александровна, НИИСИ РАН, Москва

ПРОГРАММНЫЙ КОМИТЕТ КОНФЕРЕНЦИИ

Председатель – Оселедец Иван Валерьевич – AIRI, Москва

Сопредседатель – Горбань Александр Николаевич, University of Leicester, Great Britain

Зам. председателя – Дунин-Барковский Виталий Львович – МФТИ, Физтех, Москва

Зам. председателя – Редько Владимир Георгиевич, Федеральный научный центр Научно-исследовательский институт системных исследований РАН, Москва

Зам. председателя – Тюменцев Юрий Владимирович, Московский авиационный институт (национальный исследовательский университет)

Члены Программного комитета

Академик РАН Анохин Константин Владимирович – Институт перспективных исследований мозга МГУ им. М.В. Ломоносова

Академик РАН Балабан Павел Милославович – Институт высшей нервной деятельности и нейрофизиологии РАН, Москва

Бурцев Михаил Сергеевич – London Institute for Mathematical Sciences, United Kingdom

Введенский Виктор Львович – Национальный исследовательский центр «Курчатовский институт», Москва

Доленко Сергей Анатольевич – НИИ ядерной физики им. Д.В. Скобельцына МГУ им. М.В. Ломоносова, Москва

Ежов Александр Александрович – Государственный научный центр Российской Федерации Троицкий институт инновационных и термоядерных исследований (ГНЦ РФ ТРИНИТИ), Москва

Каганов Юрий Тихонович – Московский государственный технический университет им. Н.Э. Баумана

Кирой Валерий Николаевич – Научно-исследовательский технологический центр нейротехнологий Южного федерального университета, Ростов-на-Дону

Литинский Леонид Борисович – University of Haifa, Israel

Макаренко Николай Григорьевич – Главная (Пулковская) астрономическая обсерватория Российской академии наук, Санкт-Петербург

Мишулина Ольга Александровна – Национальный исследовательский ядерный университет «МИФИ», Москва

Панов Александр Игоревич – Федеральный исследовательский центр "Информатика и управление" РАН, AIRI, МФТИ, Физтех, Москва

Самсонович Алексей Владимирович – Национальный исследовательский ядерный университет «МИФИ», Москва

Терехов Сергей Александрович – РАНИ, РАИИ

Трофимов Александр Геннадьевич – Национальный исследовательский ядерный университет «МИФИ», Москва

Ушаков Вадим Леонидович – Институт перспективных исследований мозга МГУ им. М.В. Ломоносова

Чижов Антон Вадимович – Физико-технический институт им. А.Ф. Иоффе РАН, Санкт-Петербург

Шапошников Дмитрий Григорьевич – Научно-исследовательский технологический центр нейротехнологий Южного федерального университета, Ростов-на-Дону

Шепелев Игорь Евгеньевич – Научно-исследовательский технологический центр нейротехнологий Южного федерального университета, Ростов-на-Дону

Шумский Сергей Александрович – МФТИ, Физтех, Москва

Юдин Дмитрий Александрович – МФТИ, Физтех, AIRI, Москва

Юров Артем Валерианович – Балтийский федеральный университет им. И.Канта, Калининград

Яхно Владимир Григорьевич – Институт прикладной физики РАН, Нижний Новгород

Abraham Ajith – Machine Intelligence Research Labs (MIR Labs), Scientific Network for Innovation and Research Excellence, Washington, USA

Baidyk Tatiana – The National Autonomous University of Mexico, Mexico

Borisyuk Roman – Plymouth University, United Kingdom

Cangelosi Angelo – Plymouth University, United Kingdom

Dosovitskiy Alexey – Albert-Ludwigs-Universität, Freiburg, Germany

Golovko Vladimir Adamovich – Brest State Technical University, Belarus

Hayashi Yoichi – Meiji University, Kawasaki, Japan

Husek Dusan – The Institute of Computer Science of Academy of Sciences of the Czech Republic, Prague

Izhikevich Eugene – Braincorporation, San Diego, USA

Jankowski Stanislaw – Warsaw University of Technology, Poland

Kecman Vojislav – Virginia Commonwealth University, USA

Kernbach Serge – Cybertronica Research, Research Center of Advanced Robotics and Environmental Science, Stuttgart, Germany

Koprinkova-Hristova Petia – Institute of Information and Communication Technologies, Sofia, Bulgaria

Narynov Sergazy – Alem Research, Almaty, Kazakhstan

Pareja-Flores Cristobal – Universidad Complutense de Madrid, Spain

Prokhorov Danil – Toyota Research Institute of North America, USA

Rutkowski Leszek – Czestochowa University of Technology, Poland

Sandamirskaya Yulia – Institute of Neuroinformatics, University of Zurich, Switzerland

Sirota Anton – Ludwig-Maximilians-Universität, München, Germany

Snasel Vaclav – Technical University Ostrava, Czech Republic

Tikidji-Hamburyan Ruben – Louisiana State University, USA

Tsodyks Misha – Weizmann Institute of Science, Rehovot, Israel

Tsoy Yuri – Institut Pasteur Korea, Republic of Korea

Wunsch Donald – Missouri University of Science and Technology

Уважаемые коллеги!

Конференция НЕЙРОИНФОРМАТИКА вновь собирает исследователей, работающих по актуальным направлениям теории и приложений искусственных нейронных сетей. Как и на предыдущих наших собраниях, в этом году на конференции «НЕЙРОИНФОРМАТИКА-2024» представлены доклады по проблемам теории нейронных сетей, нейробиологии, моделям адаптивного поведения, прикладным задачам нейронинформатики.

Более 200 российских ученых и наших зарубежных коллег представляют на конференции результаты своих исследований.

По сложившейся традиции конференцию открывают доклады приглашенных участников. В рамках школы-семинара участники конференции прослушают лекции известных специалистов по актуальным проблемам нейронинформатики. На мастер-классах от ведущих компаний можно познакомиться с различными приложениями искусственных нейронных сетей.

Особое внимание уделяется работам студентов, аспирантов и молодых специалистов, которые примут участие в творческом конкурсе.

За прошедшие годы научно-техническая конференция «НЕЙРОИНФОРМАТИКА» сложилась как представительный и многоплановый по тематике научный форум, в работе которого принимают участие и известные ученые, и молодые специалисты, аспиранты и студенты.

Желаем всем участникам конференции плодотворной работы, активного сотрудничества и новых творческих идей!

Оргкомитет

СОДЕРЖАНИЕ

ТЕОРИЯ НЕЙРОННЫХ СЕТЕЙ

А.Ю. ДОРОГОВ

Нейросетевой комитет компетентных классификаторов повторного входа 11

НЕЙРОСЕТЕВЫЕ ПАРАДИГМЫ И АРХИТЕКТУРЫ

В.В. КЛИНЫШОВ, А.В. КОВАЛЬЧУК, И.А. СОЛОВЬЕВ

Модель сверточной нейронной сети с динамически генерируемым ядром 24

НЕЙРОКОМПЬЮТЕРЫ, НЕЙРОМОРФНЫЕ ВЫЧИСЛЕНИЯ И ИМПУЛЬСНЫЕ НЕЙРОННЫЕ СЕТИ

D.I. ANTONOV

Systems theory principles for investigating the spiking neural network trained with the Hebbian rule 33

ГЛУБОКОЕ ОБУЧЕНИЕ

V.V. KUZNETSOV, V.A. MOSKALENKO, N.Yu. ZOLOTYKH

Diagnosis of cardiovascular diseases using recurrent and convolutional neural networks 43

Е.П. ВАСИЛЬЕВ, Д.А. ЕРМОЛАЕВ, А.Н. ЕРЕМИН

Разработка инструментария для исследования X-хромосомного мозаицизма на препаратах ядер клеток крови по данным конфокального микроскопа 53

ОБУЧЕНИЕ С ПОДКРЕПЛЕНИЕМ В НЕЙРОБИОЛОГИИ И В СИСТЕМАХ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА

Е.И. ЗАЙЦЕВ, Е.В. НУРМАТОВА

Мультиагентная система управления распределенными информационно-вычислительными ресурсами 62

ВЫЧИСЛИТЕЛЬНЫЕ ИССЛЕДОВАНИЯ ПРИНЦИПОВ И МЕХАНИЗМОВ РАБОТЫ ЕСТЕСТВЕННЫХ НЕЙРОННЫХ СИСТЕМ

К.Д. СЛАДКОВ, С.С. КОЛЕСНИКОВ

Вкусовая клетка типа III как нейрональная 72

С.В. БОЖОКИН, А.А. РЯБОКОНЬ, Т.Д. ШОХИН, А.М. ФИРСОВ

Корреляция спектральных интегралов нестационарной ЭЭГ	80
S. SUKHOV, B. BATUEV	

Predictive coding models within the neural engineering framework.....	90
---	----

ПРИКЛАДНЫЕ НЕЙРОСЕТЕВЫЕ СИСТЕМЫ

И.В. КОЛОМИЙЧУК, А.И. КАНЕВ	
-----------------------------	--

Нечеткий поиск не словарных слов	100
--	-----

A.A. ЯКОВЕНКО, М.Р. ЛОЙКОНЕН, М.М. ЗАСЛАВСКИЙ	
---	--

Методы выделения признаков в задаче распознавания эмоций по голосу.....	108
---	-----

M.YU. МАЛЬСАГОВ, Е.В. МИХАЛЬЧЕНКО, Я.М. КАРАНДАШЕВ, Л.И. СТАМОВ	
---	--

Аппроксимация процесса детонации водорода в воздушной среде искусственной нейронной сетью.....	118
--	-----

A.B. КУЗЬМЕНКО, В.С. КИРЕЕВ	
-----------------------------	--

Методология реконструкции графа знаний в малочисленных коллекциях текстовых документов	128
--	-----

В.И. КУЗНЕЦОВ, Д.А. ЮДИН	
--------------------------	--

Распознавание трёхмерных объектов на последовательности изображений с использованием нейронных полей яркости	139
--	-----

L. SAQQOUR, I. KARABULATOVA, Yu. GAPANYUK	
---	--

Exploring the shortcomings of open translation services for Arabic and Russian languages.....	146
---	-----

НЕЙРОИНФОРМАТИКА И ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ

A.P. ИСХАКОВ, М.Р. БОГДАНОВ	
-----------------------------	--

Параметрическая идентификация систем машинного зрения как задача обучения	155
---	-----

В.Н. ШАЦ	
----------	--

Упорядочивание признаков как способ снижения энтропии обучающей выборки и основа простейших алгоритмов классификации	164
--	-----

НЕЙРОННЫЕ СЕТИ И КОГНИТИВНЫЕ НАУКИ. АДАПТИВНОЕ ПОВЕДЕНИЕ И ЭВОЛЮЦИОННОЕ МОДЕЛИРОВАНИЕ

В.Е. КУРЬЯН	
-------------	--

Построение графа модели мира без учителя.....	174
---	-----

С.И. ФОКИН

Новый подход к математическому описанию потенциалов действия, отличающийся от алгоритма Ходжкина–Хаксли 183

А.С. РОМАНОВА

Моделирование автономных систем управления корпорациями на основе синтетических данных..... 193

Г.С. ВОРОНКОВ

В поиске нейрофизиологического фундамента, обуславливающего характеристики субъективного зрительного образа пространства.. 203

А.С. РАТУШНЯК, А.Л. ПРОСКУРА, В.А.ШЕВЫРИНА

Моделирование формирования протобиологических молекулярных функциональных систем 213

А.Л. ПРОСКУРА, С.О. ВЕЧКАПОВА, А.С. РАТУШНЯК

Регуляторные сети эндоцитоза глутаматных рецепторов в интерактоме дендритного шипика гиппокампа..... 221

С.А. ПОЛЕВАЯ, И.С. ЕГОРОВ, А.И. ФЕДОТЧЕВ, Е.А. МУХИНА, С.Б. ПАРИН

Управление нейропластичностью с помощью неинвазивной сенсорной стимуляции..... 230

С.В. СТАСЕНКО, А.А. ЛЕБЕДЕВ, О.В. ШЕМАГИНА, И.В. НУЙДЕЛЬ, А.В. КОВАЛЬЧУК, В.Г. ЯХНО

Адаптивная коррекция многокаскадного детектора биоморфной системы искусственного интеллекта для задач распознавания образов..... 240

ИНТЕРФЕЙС "МОЗГ-КОМПЬЮТЕР". НЕЙРОТЕХНОЛОГИИ**О.А. ГРИГОРЬЕВА, Д.Ф. КЛЕЕВА**

Соотношение физиологических характеристик стадий сна между одновременно регистрируемыми ЭЭГ И МЭГ 249

ТЕОРИЯ НЕЙРОННЫХ СЕТЕЙ

А.Ю. ДОРОГОВ

Санкт-Петербургский государственный электротехнический университет «ЛЭТИ»
vaksa2006@yandex.ru

НЕЙРОСЕТЕВОЙ КОМИТЕТ КОМПЕТЕНТНЫХ КЛАССИФИКАТОРОВ ПОВТОРНОГО ВХОДА

В статье предложен новый метод обучения частных классификаторов, а также способ агрегации их прогнозов в составе комитета. Обучение основано на гипотезе итерационного повторного входа биологических нейронных сетей и использует для своей реализации метод главных компонент. Предложена модель непрерывного обучения комитета классификаторов, пригодная для построения самообучающихся систем распознавания. Рассмотрена нейросетевая реализация классификаторов в классе самоподобных нейронных сетей.

Ключевые слова: *частный классификатор, комитет классификаторов, коллективное распознавание, область компетенции, повторный вход, непрерывное обучение, самоподобная нейронная сеть.*

Введение

Нобелевский лауреат Джеральд М. Эдельман в 1977 году выдвинул гипотезу [1,2] о том, что ресентрабельная (reentry – повторно-вводимая) передача сигналов служит общим механизмом для накопления знаний в коре головного мозга. Повторный вход по Эдельману – это динамический процесс постоянной пространственно-временной корреляции, происходящий между функционально разделёнными нейронными областями, который опосредуется повторной передачей сигналов через параллельные нервные волокна. Механизмом этого явления считается ритмическая активность мозга. С точки зрения теории динамических систем повторный вход подобен итерационному процессу, сходящемуся к некоторому устойчивому аттрактору, обобщающему накопленные знания.

Развивая собственную теорию нейродарвинизма (Neural Darwinism) [3], Эдельман отмечал, что в коре головного мозга на один и тот же сенсорный сигнал одновременно реагируют несколько конкурирующих нейронных групп, реакция которых в дальнейшем конвергируется в более глубоких слоях коры, реализуя принцип коллективного распознавания и одновременно стимулируя процесс эволюционного развития нейронных структур коры головного мозга.

В технических приложениях метод коллективного решения задачи классификации заключается в объединении моделей одиночных классификаторов в комитет с некоторым общим принципом агрегации прогнозов частных классификаторов. В настоящей работе будут рассмотрены технические решения, объединяющие методы коллективного распознавания и усиление классификаторов за счет реализации принципа повторного входа. В заключительном разделе будет представлена нейросетевая реализация предлагаемых решений.

Методы коллективного распознавания

При формировании комитета классификаторов необходимо оптимизировать два критерия – качество обучения частных классификаторов и оптимальность их объединения [4, 5]. Общее название методов построения коллективных классификаторов выражается термином «boosting» – усиление [6]. Наиболее известными являются методы bagging и adaboost.

Метод bagging, предложенный Лео Брейманом в 1994 году [7], предполагает обучение частных классификаторов на случайных выборках, формируемых из обучающего множества. Агрегация классификаторов выполняется простым голосованием.

Метод adaboost, предложенный Йоавом Фройндом и Робертом Шапире [8] в 1995 году, предполагает обучение частных классификаторов на одной выборке с пошаговой адаптацией по ошибкам обучения как самих классификаторов, так и комитета классификаторов. Агрегация частных прогнозов выполняется взвешенным голосованием. Оптимальные веса для классификаторов рассчитываются на каждом шаге алгоритма.

В 70-х и 80-х годах советские ученые Л.А. Растринин и Р.Х. Эренштейн в ряде работ разработали метод коллективного распознавания на основе компетентных классификаторов [9, 10, 11]. Идея таких методов базируется на том, что каждый базовый классификатор может работать хорошо в некоторой области пространства признаков (эта область интерпретируется как область компетентности классификатора), превосходя в этой области остальные классификаторы по точности и достоверности решений. Авторы предложили несколько способов определения областей компетенции классификаторов и доказали применимость используемого подхода на примерах.

Компетентный классификатор повторного входа

Рассмотрим классификатор, основанный на оценке длины проекции образа на подпространства классов. Подпространства класса порождаются линейной комбинацией векторов обучающей выборки, принадлежащих одному классу. Проекцию вектора-образа легко определить, если в подпространстве класса задать ортогональный базис. Сумма квадратов коэффициентов разложения вектора образа по этому базису будет определять квадрат

проекция образа на подпространство класса. Классификатор относит нормированный образ к классу, для которого проекция на подпространство класса оказывается максимальной.

Среди возможных ортогональных базисов можно подобрать оптимальный базис, который обеспечивает наиболее быструю сходимость коэффициентов разложения. Известно, что таким свойством обладает ортогональный базис, полученный методом главных компонент (англ. Principal Component Analysis, PCA) [12] из ковариационной матрицы векторов класса. Эмпирическая или выборочная ковариационная матрица класса определяется выражением:

$$C = \frac{1}{m-1} \sum_{i=1}^m (x_i - \bar{x})^T (x_i - \bar{x}),$$

где m – число примеров в классе, x_i – вектор-строка представителя класса, $\bar{x}$ – вектор-строка среднее значение примеров в классе. Ковариационная матрица является симметричной и положительно определённой. Для такой матрицы можно найти собственные ортогональные вектора v такие что $Cv = \lambda v$.

Скаляр λ называется собственным числом, соответствующим собственному вектору v . Все собственные числа ковариационной матрицы положительные, и чем большее значение они имеют, тем большей значимостью обладает собственный вектор. Для построения модели классификатора отбирают наиболее значимые векторы.

Воспользуемся теперь принципом повторного входа для обучения PCA-классификатора. В обучающей выборке выберем представительный фрагмент, который далее будем называть областью компетенции классификатора. В области компетенции произвольно выделим базовое множество примеров, которое будем использовать для построения модели PCA-классификатора. На каждом шаге эволюции по результатам классификации выполняется контроль ошибок модели на области компетенции. Примеры, на которых произошли ошибки, добавляются к базовому множеству с дифференциацией по классам и модель PCA-классификатора строится заново. Процесс многократно повторяется до полного устранения ошибок на области компетенции. Построенный классификатор будем называть компетентным.

Покажем эффективность метода обучения классификатора на примере набора данных MNIST [13]. Набор данных содержит двумерные образы рукописных цифр от 0 до 9 в виде пиксельных изображений размером 28×28 . Объем обучающей выборки равен 60 000 образов и 10 000 изображений

содержит тестовая выборка. На рис. 1 показана выборка с десятью представителями для каждого класса. При проведении экспериментов все примеры выборок нормировались по энергии к единичному уровню.

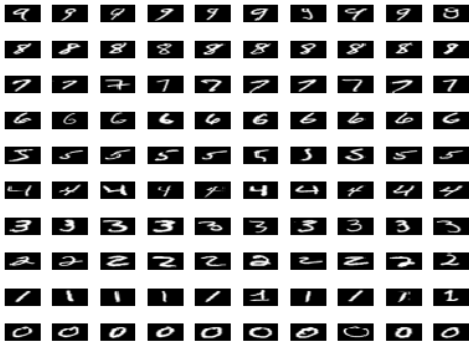


Рис. 1. Выборка из набора рукописных цифр MNIST

На рис. 2 показана зависимость точности классификатора примеров с областью компетенции из 5000 примеров и начальной базовой выборки 1000 примеров. Как видно из рисунка, итерационный процесс сходится примерно за 6 шагов. Размер моделей в этом эксперименте составлял 30 собственных векторов. В процессе эволюции базовое множество возросло с 1000 до 1385 примеров. Обученный классификатор проверялся на тестовой выборке размером 10 000 примеров. На рис. 3 показана таблица ошибок классификатора. Достигнутая точность на тестовой выборке составила 94.6%.

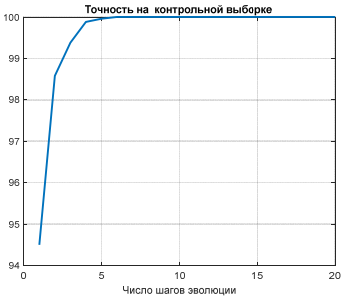


Рис. 2. Зависимость точности классификатора на области компетенции от числа шагов эволюции

1	961	2	2	2	1	5	2	5	2	961	19	
2		1123	4	2	1	1	3	3	1	1123	12	
3	8	4	969	14	4		4	11	14	4	969	63
4	1	1	14	939		28	1	7	14	5	939	71
5		2	3		951		4	8		14	951	31
6	3	2		29	3	825	5	5	17	3	825	67
7	8	2	1		4	15	924		3		924	33
8		10	16	1	2	2		974	1	22	974	54
9	6	6	8	35	5	14	6	8	871	15	871	103
10	4	5	3	9	31	3	7	23	9	922	922	87
	1	2	3	4	5	6	7	8	9	10		
Фактический класс	Предсказанный класс											

Рис. 3. Таблица ошибок PCA-классификатора повторного входа на тестовой выборке

Комитет компетентных классификаторов

Для построения комитета классификаторов разделим всю обучающую выборку на представительные области компетенции и на каждой области обучим РСА-классификатор с повторным входом. Для формирования агрегированного результата комитета будем использовать принцип максимума.

На рис. 4 показана матрица ошибок для полного обучающего множества, состоящего из 60000 примеров. Точность классификации на обучающей выборке составляет чуть более 99 процентов при размере комитета, равным 12. На рис. 5 показана матрица ошибок этого же коллективного классификатора на тестовой выборке, состоящей из 10 000 примеров. Точность классификации составляет примерно 97 процентов.



Рис. 4. Таблица ошибок комитета эволюционных классификаторов для обучающей выборки



Рис. 5. Таблица ошибок комитета эволюционных классификаторов для тестовой выборки

Для эксперимента использовались классификаторы повторного входа с областью компетенции 5000 примеров и начальным размером базового множества 1000 примеров. На текущий момент мировое достижение [14] по точности классификации тестового набора данных MNIST составляет 99.87%.

Коллективный классификатор непрерывного обучения

Природа не разделяет данные на обучающую и тестовую выборки, биологический мозг обучается непрерывным потоком поступающих данных. С этой точки зрения имеет смысл отказаться от общепринятого деления выборки на обучающую и тестовую. Принцип разделения классификаторов по областям компетенции, с агрегированием по максимуму, позволяет неограниченно наращивать размер комитета, покрывая всю выборку и обучающую, и тестовую, а также все последующие поступившие данные. На

рис. 6 приведена матрица ошибок для объединённой выборки коллективного классификатора для объединённой выборки с числом примеров 70 000.

В комитете использовались эволюционные классификаторы с областью компетенции 5000 примеров и начальным базовым множеством 1000 примеров. Число компетентных классификаторов для объединённой выборки из 70 000 примеров составило 14. Точность коллективного классификатора составила 99.49%. Модель непрерывного обучения может быть использована для построения самообучающихся систем распознавания реального времени, за счет добавления новых классификаторов в процессе функционирования системы.

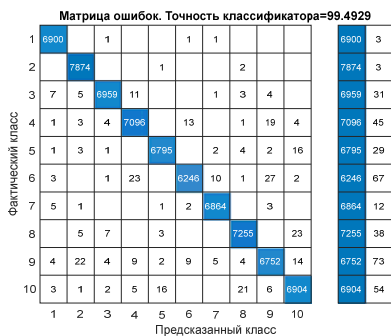


Рис. 6. Таблица ошибок
непрерывного классификатора на
объединённой выборке

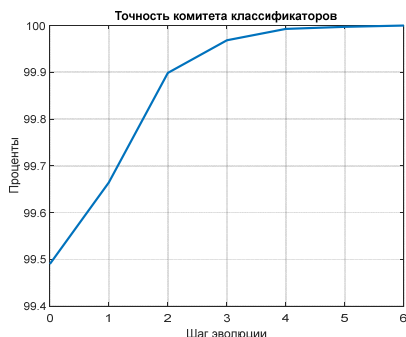


Рис. 7. Зависимость точности
комитета классификаторов от числа
шагов эволюции

Эволюционный принцип повторного входа можно распространить также на весь комитет классификаторов. С этой целью в модели каждого компетентного классификатора запоминаются номера примеров базового множества на момент завершения обучения (т.е. достижения 100% точности на области компетенции). Алгоритм эволюции для комитета, на каждом шаге фиксирует номера примеров в объединённой выборке, на которых произошли ошибки. Эти примеры разделяются по областям компетенции частных классификаторов и добавляются в списки примеров базовых множеств. Далее заново выполняется обучение частных классификаторов с обновлёнными базовыми множествами до достижения 100% точности на областях компетенции и вновь контролируются ошибки комитета классификаторов, далее итерации комитета повторяются.

На рис. 7 показана зависимость точности комитета классификаторов от числа шагов эволюции. Мета параметры комитета соответствуют предыдущему примеру. Эволюционный процесс сходится к 100% процентной точности.

Быстрая нейронная сеть

Классификатор повторного входа вычисляет проекцию вектора-образа на подпространства классов. Проекция определяется через разложение вектора в усеченных ортогональных базисах. Это линейная матричная операция, которая реализуется через вычисление скалярных произведений векторов. Переход к нейросетевой реализации связан с факторизацией линейного преобразования в параллельно-последовательную матричную форму.

В основе построения схемы факторизации лежат классические алгоритмы быстрых преобразований. В работе [15] показано, что они являются частным случаем многослойных самоподобных нейронных сетей. На рис. 8 представлен граф быстрого преобразования Фурье размерности 8 с топологией Кули – Тьюки «с прореживанием по времени».

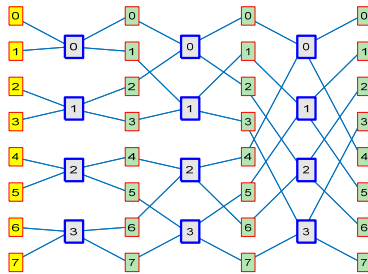


Рис. 8. Граф быстрого преобразования Фурье

В каждом слое выделены четыре базовых операции типа «Бабочка». Для преобразования Фурье параметры базовых операций полностью определены, однако, если допустить, что их параметры можно изменять, то мы приходим к варианту быстрых нейронных сетей (БНС) [16]. В этом контексте базовые операции уместно назвать нейронными ядрами. Сетевая модель данной топологии и описывается набором кортежей [15]:

$$\begin{aligned} U^m &= \langle u_{n-1} u_{n-2} \cdots u_{m+1} u_m v_{m-1} v_{m-2} \cdots v_1 v_0 \rangle, \\ V^m &= \langle u_{n-1} u_{n-2} \cdots u_{m+1} v_m v_{m-1} v_{m-2} \cdots v_1 v_0 \rangle, \\ z^m &= \langle u_{n-1} u_{n-2} \cdots u_{m+1} v_{m-1} v_{m-2} \cdots v_1 v_0 \rangle, \end{aligned} \quad (1)$$

где z^m – порядковый номер нейронного ядра, U^m, V^m – порядковые номера рецепторов и аксонов в слое m , u_m, v_m – локальные номера рецепторов и аксонов в пределах нейронного ядра слоя m . Каждый кортеж представляет собой поразрядную форму представления порядкового номера через разрядные переменные u_m, v_m . Основания разрядных переменных u_m, v_m определяются целыми положительными числами p_m, g_m , и по слоям могут быть заданы соответствиями:

$$\begin{pmatrix} u_0 & u_1 & \cdots & u_{n-2} & u_{n-1} \\ p_0 & p_1 & \cdots & p_{n-2} & p_{n-1} \end{pmatrix}, \quad \begin{pmatrix} v_0 & v_1 & \cdots & v_{n-2} & v_{n-1} \\ g_0 & g_1 & \cdots & g_{n-2} & g_{n-1} \end{pmatrix}.$$

Для ядер размерности 2×2 имеем $p_m = 2, g_m = 2$, и все разрядные переменные принимают значения $\{0, 1\}$. Координатные направления U^m и V^m в дальнейшем будем называть входной и выходной плоскостью нейронного слоя. Для терминальных плоскостей (относящихся к начальному и конечному слою сети) будем использовать обозначения U и V . Для построения быстрых алгоритмов размерность преобразования должна быть составным числом, и чем больше множителей в разложении размерности, тем выше вычислительная эффективность быстрого алгоритма. Размерности быстрой нейронной сети по входу и выходу вычисляются через произведения оснований разрядных переменных:

$$N = p_{n-1} \cdots p_1 p_0, \quad M = g_{n-1} \cdots g_1 g_0.$$

Число слоёв в быстром преобразовании равно числу сомножителей в этих произведениях. Несмотря на большое разнообразие быстрых алгоритмов, конфигурации их структур удовлетворяют системному инварианту самоподобия [15]. Как известно таким же свойством самоподобия обладают фракталы. Свойство структурной фрактальности позволяет решить одновременно две задачи: реализовать быструю обработку данных и выполнить быстрое обучение преобразования. Обучение быстрого преобразования заключается в выборе значений элементов нейронных ядер, так чтобы в столбцах матрицы преобразования содержался заданный набор эталонных функций, это могут быть, например, ортогональные собственные векторы выборочной ковариационной матрицы. В [16] показано, что для самоподобных нейронных сетей элементы матрицы быстрого преобразования могут быть выражены через произведения элементов нейронных ядер:

$$h(U, V) = w_{z^{n-1}}^{n-1}(u_{n-1}, v_{n-1}) w_{z^{n-2}}^{n-2}(u_{n-2}, v_{n-2}) \cdots w_{z^0}^0(u_0, v_0).$$

Там же доказано, что произвольная функция, заданная на дискретном интервале длиной $N = p_{n-1} \cdots p_1 p_0$, может быть представлена в мультипликативной форме:

$$f(u) = \phi_{i^0}(u_0)\phi_{i^1}(u_1)\dots\phi_{i^{n-2}}(u_{n-2})\phi_{i^{n-1}}(u_{n-1}),$$

где $i^m = \langle u_{n-1}u_{n-2}\dots u_{m+1} \rangle$. Отсюда следует правило настройки нейронных ядер:

$$w_{z^m}^m(u_m, v_m) = \phi_{i^m}^k(u_m). \quad (2)$$

Здесь k – номер опорной функции. Зададим точку привязки эталонной функции в выходной плоскости числом, представленным в поразрядной форме:

$$x = \langle x_{n-1}x_{n-2}\dots x_0 \rangle,$$

тогда номера настраиваемых ядер по слоям будут определяться выражением:

$$z^m = \langle u_{n-1}u_{n-2}\dots u_{m+1}x_{m-1}x_{m-2}\dots x_1x_0 \rangle.$$

Для $m = 0$ имеем $z^0 = \langle u_{n-1}u_{n-2}\dots u_1 \rangle$, это означает, что независимо от выбора точки привязки, все ядра слоя будут настраиваться, причём номер ядра определяется из условия: $z^0 = i^0$. Настройка элементов ядер этого слоя выполняется по правилу:

$$w_{z^0}^0(u_0, v_0) = \phi_{i^0}^k(u_0).$$

Очевидно, должно быть задано взаимно-однозначное соответствие $k \leftrightarrow v_0$ между номером опорной функции k и разрядной переменной v_0 . Эта разрядная переменная принимает значения $0, 1, \dots, g_0 - 1$. Отсюда следует вывод, что число эталонных функций не может быть больше, чем g_0 , а если ещё потребовать выполнения условия ортогональности ядер, то матрица такого преобразования будет содержать только одну произвольную функцию в качестве столбца. Этого явно недостаточно для реализации РСА-классификатора. В следующем разделе будет показана модификация топологии БНС, которая позволяет решить поставленную задачу.

Самоподобные нейронные сети с дополнительными плоскостями

Дополним топологическую модель быстрого преобразования (1) дополнительными плоскостями нейронных ядер [17], номера которых определим переменной π_m , новая топологическая модель в этом случае будет описываться набором кортежей:

$$\begin{aligned}
 U^m &= \langle u_{n-1}u_{n-2} \cdots u_{m+1}u_mv_{m-1}v_{m-2} \cdots v_1v_0 \rangle, \\
 V^m &= \langle u_{n-1}u_{n-2} \cdots u_{m+1}v_mv_{m-1}v_{m-2} \cdots v_1v_0 \rangle, \\
 z^m &= \langle u_{n-1}u_{n-2} \cdots u_{m+1}v_{m-1}v_{m-2} \cdots v_1v_0 \rangle, \\
 \pi_m &= \langle v_{n-1}v_{n-2} \cdots v_{m+2}v_{m+1} \rangle.
 \end{aligned}$$

Максимальное количество дополнительных плоскостей появится в нулевом слое. Номер плоскости в нулевом слое будет определяться кортежем $\pi_0 = \langle v_{n-1}v_{n-2} \cdots v_2v_1 \rangle$. Плоскость с номером $\pi_0 = 0$ будем считать плоскостью исходной базовой топологической структуры. Число плоскостей в нулевом слое будет равно произведению оснований: $g_{n-1}g_{n-2} \cdots g_2g_1$. По мере движения к выходному слою число дополнительных плоскостей будет уменьшаться и для последнего слоя $\pi_{n-1} = \langle \rangle$, т.е. их не будет совсем, останется только одна плоскость базовой топологии. Таким образом, в новой топологии плоскость последнего слоя останется прежней, а в младших слоях появятся дополнительные плоскости. Номер ядра, теперь следует уточнять его размещением в дополнительной плоскости. На рис. 9 показана новая топология, построенная на базе трёхслойной БНС с основаниями 2.

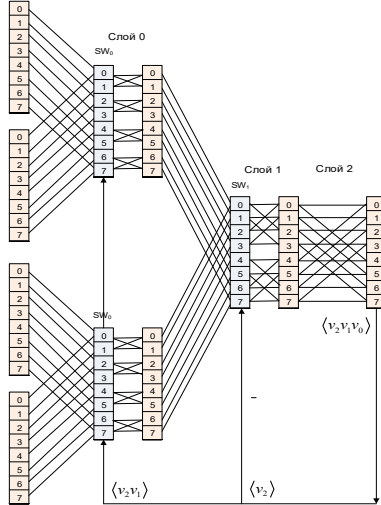


Рис. 9. Топология самоподобной нейронной сети с дополнительными плоскостями

В сети условно показаны коммутаторы (SW), которые служат для пояснения принципа работы сети. Коммутаторы управляются разрядными переменными точек выходной плоскости. Фактически коммутаторов нет, они реализуются в составе алгоритма и не препятствуют параллельной обработке входных данных. В контексте спектрального преобразования точки выходной плоскости точки ассоциируются со спектральными коэффициентами, поэтому привязка спектрального коэффициента к координатам выходной плоскости предопределяет выбор дополнительных плоскостей, которые используются для обработки входных образов. При построении полного спектрального анализатора все векторные входы сети параллельно объединяются.

Поскольку правило порождения новых плоскостей не противоречит базовой топологической модели, то для настройки ядер преобразования к эталону можно использовать прежнее правило (2), расширив его аргументом π_m для дополнительных плоскостей:

$$w_{z^m}^m \langle \pi_m \rangle (u_m, x_m) = \phi_{i^m}^k (u_m),$$

здесь k – номер эталонной функции, $x = \langle x_{n-1} x_{n-2} \dots x_0 \rangle$ – точка привязки, $z^m = \langle u_{n-1} u_{n-2} \dots u_{m+1} x_{m-1} x_{m-2} \dots x_1 x_0 \rangle$ – номер настраиваемых ядер по слоям, $\pi_m = \langle x_{n-1} x_{n-2} \dots x_{m-1} \rangle$ – номер плоскости размещения ядер. Индекс k в правой части нумерует точку приспособления. Для $m=0$ имеем $z^0 = i^0 = \langle u_{n-1} u_{n-2} \dots u_1 \rangle$, а варьируемыми переменными в левой части являются номер плоскости $\pi_0 = \langle x_{n-1} x_{n-2} \dots x_1 \rangle$ и разряд x_0 . Вместе они покрывают весь диапазон координат выходной плоскости. Этому диапазону отвечают возможные значения индекса k в правой части, отсюда следует, что каждая точка выходной плоскости может быть связана с собственной эталонной функцией. Т.е. построенная сеть обладает максимально возможным числом точек привязки, покрывающих всю выходную плоскость и, таким образом может быть использована для реализации произвольного линейного преобразования размерности $N \times M$.

В модели РСА-классификатора для каждого класса сохраняется только часть ортогональных векторов, которые обладают наибольшей значимостью. Размерность выходной плоскости следует выбрать так чтобы значение M было не меньше, чем общее число всех векторов модели классификатора. Подобным образом можно построить все частные классификаторы комитета. Нейросетевая реализация обеспечивает максимальное распараллеливание вычислительных операций, что позволяет получить высокое быстроедействие классификатора на специализированных процессорах.

Заключение

Рассмотренный метод построения коллективных классификаторов отличается от методов bagging и adaboost использованием классификаторов с заданными областями компетенции и способом получения агрегативного решения комитета.

В отличие от метода коллективного распознавания Л.А. Растригина и Р.Х. Эренштейна, области компетенции классификаторов не локализуются расчётным путём, а назначаются исходя из условия достаточной представительности классов в базовых множествах и областях компетенции. Благодаря принципу повторного входа частные классификаторы идеально обучаются в пределах областей их компетенций. Это позволяет сделать процедуру классификации одноступенчатой, без предварительного определения области компетенции. Предложенный способ агрегации решений частных классификаторов по принципу максимума, позволяет наращивать классификатор в процессе функционирования и реализовать принцип непрерывного обучения.

В статье показано, что РСА-классификаторы можно реализовать в классе линейных самоподобных нейронных сетей с дополнительными плоскостями, при этом число реализуемых эталонных функций по сравнению с БНС кардинально возрастает и покрывает все элементы выходной плоскости нейронной сети. Такое расширение топологии не нарушает принципа построения обучающего алгоритма и позволяет строить и обучать РСА- классификаторы с произвольной размерностью модели. Алгоритмы обучения к модельным векторам являются абсолютно устойчивыми и завершаются за конечное число шагов.

Список литературы

1. Edelman G.M. Group selection and phasic reentrant signaling: a theory of higher brain function // The Mindful Brain: cortical organization and the group-selective theory of higher brain function. Eds. Edelman G.M., Mountcastle V.B. Boston: MIT Press. 1978. P. 51–98.
2. Эделман Дж., Маунткасл В. Разумный мозг. Москва : Мир. 1981.
3. Edelman G.M. Neural Darwinism: The theory of neuronal group selection. New York : Basic Books. 1987. 240 p.
4. Городецкий В.И., Серебряков С.В. Методы и алгоритмы коллективного распознавания: обзор // Труды СПИИРАН. СПб. : Наука. 2006. Вып. 3, Т. 1.
5. Терехов С.А. Гениальные комитеты умных машин // IX Всероссийская научно-техническая конференция «Нейроинформатика-2007» : Лекции по нейроинформатике. Ч. 2. Москва : МИФИ. 2007. С. 11–43.
6. Bauer E., Kohavi R. An empirical comparison of voting classification algorithms: Bagging, boosting, and variants // Machine Learning. 1999. 36. P. 105–142.

7. Breiman L. Bagging Predictors. Department of Statistics University of California Berkeley, California. Technical Report. 1994. N 421.
8. Freund Y., Schapire R.E. A decision-theoretic generalization of online learning and an application to boosting // Second European Conference on Computational Learning Theory. 1995.
9. Растринин Л.А., Эренштейн Р.Х. Обучение коллектива решающих правил // Адаптивные системы. 1974. Вып. 4. С. 8–20.
10. Растринин Л.А., Эренштейн Р.Х. Принятие решений коллективом решающих правил в задачах распознавания образов // Автоматика и телемеханика. 1975. № 9. С. 133–144.
11. Растринин Л.А., Эренштейн Р.Х. Метод коллективного распознавания. Москва: Энергоиздат. 1981. 80 с.
12. Лагутин М.В. Наглядная математическая статистика : учеб. пособие. Москва : БИНОМ. Лаборатория знаний. 2007. 472 с.
13. <http://yann.lecun.com/exdb/mnist/>
14. Byerly A., Kalganova T., Dear I. N Routing Needed Between Capsules // Neurocomputing. 2021. V. 463. P. 545–553.
15. Дорогов А.Ю. Быстрые преобразования и самоподобные нейронные сети глубокого обучения. Ч. 1. Стратифицированные модели самоподобных нейронных сетей и быстрых преобразований // Информационные и математические технологии в науке и управлении. 2023. № 4(32). С. 5–20.
16. Дорогов А.Ю. Быстрые преобразования и самоподобные нейронные сети глубокого обучения. Ч. 2. Методы обучения быстрых нейронных сетей // Информационные и математические технологии в науке и управлении. 2024. № 1(33). С.5–19.
17. Дорогов А.Ю. Быстрые нейронные сети глубокого обучения III Международная научная конференция по проблемам управления в технических системах (CTS'2019). Сборник докладов. Санкт-Петербург. 30 октября – 1 ноября 2019 г. СПб. : СПбГЭТУ «ЛЭТИ». С. 275–280.

НЕЙРОСЕТЕВЫЕ ПАРАДИГМЫ И АРХИТЕКТУРЫ

В.В. КЛИНЬШОВ, А.В. КОВАЛЬЧУК, И.А. СОЛОВЬЕВ

Институт прикладной физики РАН, Нижний Новгород

vladimir.klinshov@gmail.com, solocool46@gmail.com

МОДЕЛЬ СВЕРТОЧНОЙ НЕЙРОННОЙ СЕТИ С ДИНАМИЧЕСКИ ГЕНЕРИРУЕМЫМ ЯДРОМ*

Одним из типов слоев в искусственных нейросетях являются сверточные слои, основанные на свертке входа с некоторым ядром. В классическом случае ядро не зависит от входных данных и настраивается в процессе обучения сети. Альтернативой являются динамические свертки, в которых ядро зависит от входной информации. В данной работе предложен новый вариант реализации динамической свертки. Показано ее преимущество по скорости обучения и метрике $f1$ по сравнению с нединамической в задаче сравнения изображений при одинаковом числе обучаемых параметров.

Ключевые слова: *нелинейная динамика, нейронные сети, динамическая свертка.*

Введение

В последние годы наблюдается рост использования систем искусственного интеллекта, многие из которых основываются на искусственных нейронных сетях. Важнейшим этапом подготовки такой сети для решения какой-либо задачи является ее обучение, то есть настройка параметров для оптимизации ее эффективности [1]. Как правило, процесс обучения является итерационным и основан на последовательном предъявлении сети серии обучающих примеров, после каждого из которых следует изменение параметров сети в соответствии с некоторым правилом обучения. Таким образом, в процессе обучения нейронная сеть представляет собой динамическую систему, фазовое пространство которой представляет собой пространство обучаемых параметров.

Помимо классических полносвязных сетей [2], широкое распространение получили сверточные сети [3], в которых входные данные сворачиваются с некоторыми обучаемыми ядрами, формируя так называемые карты признаков. В реализации свертки, называемой далее «классической», для всех входных данных используется одно и то же ядро, которое настраива-

* Работа выполнена при финансовой поддержке РФФИ, грант № 23-72-10088.

ется в процессе обучения. В нескольких работах [4, 5], однако, была рассмотрена идея так называемых динамических сверток, в которых для каждого обрабатываемого примера ядро подбирается на основе самого этого примера. Другими словами, ядро динамической свертки, в отличие от классической, зависит от входной информации. Помимо классификации изображений, динамические свертки успешно были применены в задачах обработки аудиосигналов [6], детектирования звуков [7] и даже в языковой обработке [8, 9].

В настоящей работе предложена нейронная сеть с динамическими свертками для решения задачи сопоставления изображений (image matching). Стандартным методом решения таких задач является поиск на изображениях некоторых ключевых точек (обычно углов объектов) и вычисления для них определенного дескриптора [10]. Для более быстрой обработки в реальном времени был придуман алгоритм FAST, основанный на анализе круговой области вокруг исследуемого пикселя на яркость относительно центра [11].

Важным отличием подхода, предлагаемого в настоящей работе, от других работ по динамическим сверткам является способ генерации динамических ядер. Как правило, эти ядра в других работах генерируются как взвешенная сумма из некоего множества фиксированных ядер [12], весовые коэффициенты (иногда они называются сигналом внимания) генерируются специальной обучаемой подсетью. В настоящей работе предложен альтернативный вариант реализации динамической свертки, при котором вместо взвешенной суммы фиксированных ядер используется специальная подсеть, генерирующая ядро целиком из сжатых представлений входных данных.

Показано, что предложенная архитектура с динамической сверткой имеет более высокое качество работы по метрике $f1$, чем нейронная сеть со стандартной сверткой при равном числе обучаемых параметров. Помимо этого, предложенная архитектура имеет другие важные преимущества, например возможность сравнивать изображения разных размеров.

Модель динамической свертки

Работа посвящена решению задачи сопоставления изображений, которая формулируется следующим образом. Пусть есть два изображения $g \in G$ (галерея) и $q \in Q$ (запрос), каждому из которых соответствует число (метка) $l_g \in L_g$ и $l_q \in L_q$, указывающая класс изображенного объекта. Задача состоит в ответе на вопрос о том, относятся ли объекты с обоих изображений к одному классу. Целевое значение ответа $l_p \in L_p = \{0, 1\}$ для пары изображений g и q определяется как

$$l_p = \begin{cases} 1, l_g = l_q \\ 0, l_g \neq l_q \end{cases}. \quad (1)$$

Для решения поставленной задачи в работе используется модель в виде искусственной нейронной сети прямого распространения, которая на основе подаваемых на нее изображений g и q вычисляет ответ

$$y = S(g, q, \theta), \quad (2)$$

где θ – совокупность обучаемых параметров сети. Задача обучения сети состоит в настройке их таким образом, чтобы минимизировать ошибку

$$\varepsilon = -l_p \ln(y) + (1 - l_p) \ln(1 - y) \quad (3)$$

на некоторой выборке изображений.

В настоящей работе используются сверточные нейронные сети, то есть сети, включающие в себя слои, которые выполняют дискретную свертку поступающих на них сигналов (изображений) x с некоторым ядром w .

Классическая реализация сверточного слоя предполагает использование для всех входных данных одного ядра, при этом значения его параметров настраиваются в процессе обучения сети. В настоящей работе предложена архитектура с использованием динамических сверток, в которых ядро свертки зависит от входного сигнала – $w = w(x)$, причем для генерации ядер используются обучаемая подсеть.

Для выявления преимуществ сети с динамическими сверточными слоями в работе проводится ее сравнение с базовой (общепринятой) архитектурой на основе стандартных сверточных слоев, называемой далее «базовой» моделью. В ней оба входных изображения объединяются в многоканальный тензор, после чего последовательно подаются на сверточные слои (с последующим сжатием), и далее на полносвязный слой-классификатор. Предложенная в работе модель (далее «динамическая») может быть представлена как две связанные подсети G и Q , где сеть Q получает сжатые представления e_q изображения q , а сеть G осуществляет сравнение входного изображения g с этим представлением. Обе подсети являются сверточными, однако, в отличие от базовой модели, ядра сверток, используемых в сети G , зависят от входных данных, поскольку генерируются сетью Q . Это означает, что динамическая модель может быть формально описана как:

$$y = G(g, q, \theta_G), \quad (4)$$

$$e_q = Q(q, \theta_Q), \quad (5)$$

где θ_G и θ_q – обучаемые параметры подсетей G и Q , соответственно.

Использование сжатого представления имеет ряд преимуществ. Во-первых, оно отражает основные свойства изображения, но имеет фиксированный размер, не зависящий от размера исходного изображения. Во-вторых, сжатые представления всех изображений одного класса близки, что позволяет в случае известного класса запроса использовать представление e_{l_q} типичного представителя данного класса. Это позволяет отключать подсеть Q и экономить вычислительные ресурсы в задачах детектирования принадлежности изображения g заранее заданному классу l_q .

Рассмотрим реализацию базовой и динамической модели (рис. 1, а), б) соответственно). Базовая модель на первом шаге осуществляет конкатенацию входных изображений g и q , после чего полученное двухканальное изображение последовательно подается на два обычных сверточных блока. Эти блоки включают стандартные сверточные слои с фиксированными ядрами, их детализированная структура будет описана ниже. Затем изображение подается на блок адаптивного сжатия, на котором изображение сжимается до фиксированного размера 2×2 пикселей. Наконец, полученные данные подаются на полносвязный слой с последующей сигмоидной нелинейностью, этот полносвязный слой выполняет роль классификатора.

Динамическая модель же состоит из двух подсетей G и Q . Сеть Q состоит из тех же блоков и имеет ту же структуру, что и базовая модель, однако на вход принимает только изображение q , а на выходе вычисляет его сжатое представление e_q , которое затем подается на сеть G . Количество нейронов в выходном слое при этом уже равно не 1, а размеру сжатого представления. Структура подсети G также похожа на структуру базовой модели, но вместо стандартных сверточных блоков в ней используются динамические сверточные блоки, в которых на основе полученного сжатого представления происходит вычисление ядер $w(e_q)$. Эти ядра далее используется в сверточном слое. Размер сжатого представления и количество вычисляемых ядер являются гиперпараметрами модели, влияющими на число ее обучаемых параметров.

Рассмотрим теперь подробнее обычный и динамический сверточные блоки (рис. 2 а), б) соответственно). Блок обычной свертки состоит из следующей последовательности слоев: 1) расширяющий сверточный слой, в котором используются ядра размера 1×1 по вертикали и горизонтали, однако имеющие отличный от единицы размер по каналам, что приводит к перемешиванию и увеличению числа каналов; 2) последовательность слоев батч-нормализации и нелинейности ReLU; 3) depth-wise сверточный слой, в котором используются фиксированные обучаемые ядра 3×3 без перемешивания каналов; 4) еще одна последовательность из батч-нормализации и

ReLU; 5) сжимающий сверточный слой с ядром размера 1×1 , который перемешивает каналы и уменьшает их число; 6) третья последовательность батч-нормализации и ReLU; 7) и, наконец, слой, производящий сжатие изображения в 2 раза по обоим измерениям. Количество входных и выходных каналов в блоке, а также множитель, определяющий степень поканального растяжения, являются задаваемыми гиперпараметрами блока.

Динамический сверточный блок состоит из двух ветвей. Первая ветвь по структуре близка к такой у обычного, но ядра для depth-wise слоя не фиксированы, а вычисляются с помощью второй ветви на основе получаемого ей сжатого представления. Ветка для генерации ядер состоит из двух полносвязных слоев со слоем нелинейности ReLU между ними. Размеры этих слоев являются задаваемыми гиперпараметрами, при этом скрытый слой может быть удален для облегчения модели.

(а)



(б)

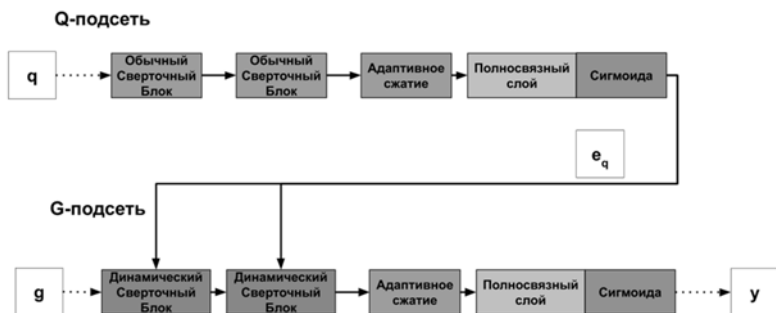


Рис. 1. а) общая схема «базовой» сети, б) общая схема «динамической» сети

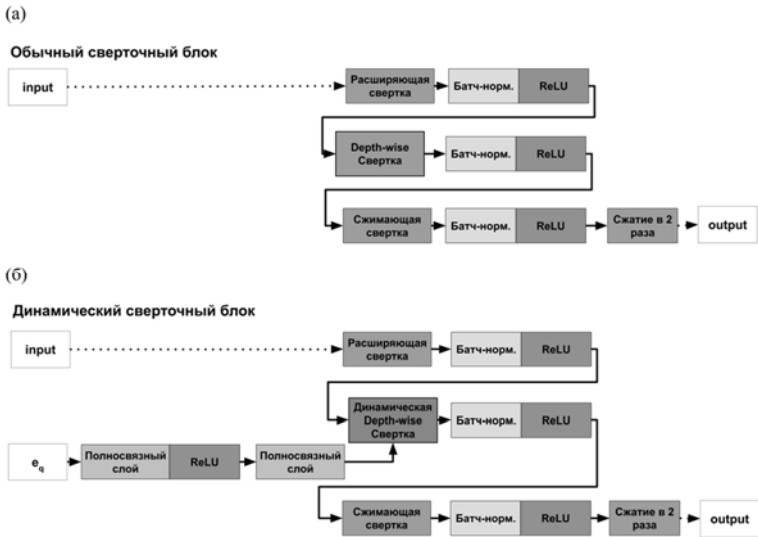


Рис. 2. а) схема обычного сверточного блока, б) схема динамического сверточного блока

Эксперименты

В рамках экспериментов рассматривались модели с различными наборами гиперпараметров, причем для корректности сравнения базовой и динамической моделей они выбирались таким образом, чтобы число обучаемых параметров в обеих моделях было приблизительно одинаковым, причем для каждой было сделано два варианта – «легкий» и «тяжелый». Сведения о количестве обучаемых параметров во всех моделях приведены в Таблице 1.

Для обучения моделей использовалась база данных образцов рукописного написания цифр MNIST [3].

Таблица 1

Сравнение используемых моделей

Сеть	Кол-во параметров	Среднее f1 за последние 10 эпох
Динамическая тяжелая	14602	0,96
Базовая тяжелая	15301	0,93
Динамическая легкая	9874	0,94
Базовая легкая	9376	0,91

В эксперименте использовалась полная база с изображениями всех десяти цифр, из нее был случайным образом сформирован набор пар изображений g и q , на половине из которых значения цифр совпадали, а на другой половине не совпадали. Количество выходных каналов из двух сверточных блоков у обеих динамических моделей равнялось соответственно 5 и 10, коэффициенты поканального растяжения были равны 2. Размеры сжатых представлений в тяжелой и легкой динамических моделях были соответственно 16 и 8, при этом скрытый полносвязный слой в ветвях генерации ядер отсутствовал. Для тяжелой базовой модели количество выходных каналов из блоков равнялось 20 и 40, а для легкой – 15 и 30. Коэффициент растяжения в обеих моделях для обоих блоков был равен 4.

Инициализация моделей во всех экспериментах происходила по Xavier [13], а обновление весов – по алгоритму Adam [14] с базовой скоростью обучения $\lambda = 10^{-4}$. При обучении использовался размер батча, равный 100. Изменения функции потерь на обучающей и валидационной базе в зависимости от времени обучения представлены на рис. 3(а) и б) соответственно). Видно, что при сопоставимом количестве обучаемых параметров динамические модели обучаются быстрее и результативнее, чем базовые.

В качестве характеристики качества работы сети после обучения была выбрана метрика $f1$. Графики зависимости $f1$ от эпохи обучения представлены на рис. 4. Видно, что при близком числе обучаемых параметров динамическая модель демонстрирует более высокое качество распознавания, чем базовая.

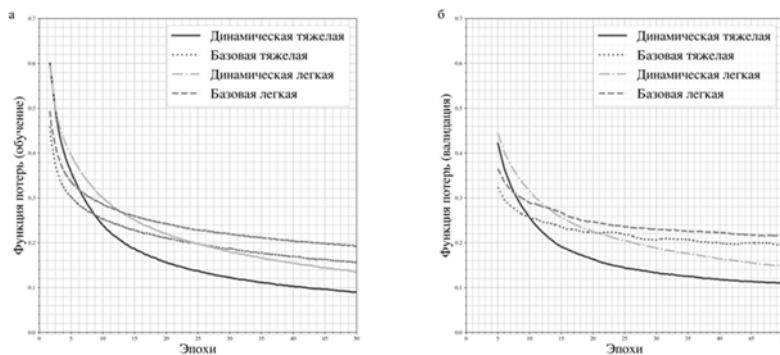
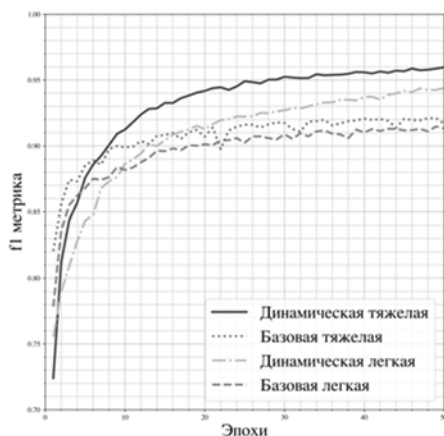


Рис. 3. а) Функции потерь при обучении б) Функции потерь при валидации

Рис. 4. Метрика $f1$ при валидации

Заключение

В работе предложена новая архитектура сверточной нейронной сети для решения задачи сопоставления изображений, использующая динамически генерируемые на основе входных данных ядра. Предложенная модель обучается быстрее, чем общепринятая «базовая» архитектура с тем же числом обучаемых параметров, а при валидации демонстрирует более высокие показатели по метрике $f1$.

Предложенная архитектура обладает и другими преимуществами. Первое из них — возможность за счет использования сжатых представлений сравнивать изображения разного размера. Второе преимущество заключается в возможности отключения части сети при решении задачи сопоставления изображений g с изображением заранее известного класса l_q . В этом случае вместо подсети Q можно использовать заранее сгенерированное представление типичного представителя класса l_q , что позволит сэкономить память и вычислительные мощности.

Архитектура протестирована на базе изображений MNIST. Для демонстрации практической применимости необходимо дальнейшее исследование ее обучения на более сложных данных. Также важным вопросом явля-

ется изучение способности предложенной модели к генерализации – способности сопоставлять изображения из классов, не использованных при ее обучении. Эти направления являются приоритетными для дальнейших исследований.

Список литературы

1. Rumelhart D. E., Hinton G. E., Williams R. J. Learning representations by back-propagating errors // *Nature*. – 1986. – V. 323. – N 6088. – P. 533–536.
2. Block H. D. The perceptron: A model for brain functioning. i // *Reviews of Modern Physics*. – 1962. – V. 34. – N 1. – P. 123.
3. LeCun Y. [et al.]. Gradient-based learning applied to document recognition // *Proceedings of the IEEE*. – 1998. – V. 86. – N 11. – P. 2278–2324.
4. Jia X. [et al.]. Dynamic filter networks // *Advances in neural information processing systems*. – 2016. – V. 29.
5. Zhang Y. [et al.]. Dynet: Dynamic convolution for accelerating convolutional neural networks // *arXiv preprint arXiv:2004.10694*. – 2020.
6. Kim S. H., Nam H., Park Y. H. Temporal dynamic convolutional neural network for text-independent speaker verification and phonemic analysis // *ICASSP 2022–2022 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. – IEEE, 2022. – P. 6742–6746.
7. Nam H. [et al.]. Frequency dynamic convolution: Frequency-adaptive pattern recognition for sound event detection // *arXiv preprint arXiv:2203.15296*. – 2022.
8. Wu F. [et al.]. Pay less attention with lightweight and dynamic convolutions // *arXiv preprint arXiv:1901.10430*. – 2019.
9. Jiang Z. H. [et al.]. Convbert: Improving bert with span-based dynamic convolution // *Advances in Neural Information Processing Systems*. – 2020. – V. 33. – P. 12837–12848.
10. Rodríguez M., Facciolo G., Morel J. M. Robust homography estimation from local affine maps // *Image Processing on Line*. – 2023. – V. 13. – P. 65–89.
11. Rosten E., Drummond T. Machine learning for high-speed corner detection // *Computer Vision–ECCV 2006: 9th European Conference on Computer Vision, Graz, Austria, May 7–13, 2006. Proceedings, Part I 9*. – Springer Berlin Heidelberg, 2006. – P. 430–443.
12. Chen Y. [et al.]. Dynamic convolution: Attention over convolution kernels // *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. – 2020. – P. 11030–11039.
13. Glorot X., Bengio Y. Understanding the difficulty of training deep feedforward neural networks // *Proceedings of the thirteenth international conference on artificial intelligence and statistics*. – *JMLR Workshop and Conference Proceedings*, 2010. – P. 249–256.
14. Kingma D. P., Ba J. Adam: A method for stochastic optimization // *arXiv preprint arXiv:1412.6980*. – 2014.

НЕЙРОКОМПЬЮТЕРЫ, НЕЙРОМОΡФНЫЕ ВЫЧИСЛЕНИЯ И ИМПУЛЬСНЫЕ НЕЙРОННЫЕ СЕТИ

D.I. ANTONOV

Ulyanovsk State Technical University
d.antonov@ulireran.ru

SYSTEMS THEORY PRINCIPLES FOR INVESTIGATING THE SPIKING NEURAL NETWORK TRAINED WITH THE HEBBIAN RULE*

The article presents the results of research on the synaptic weights distribution obtained by applying a new local learning rule for spiking neural networks (SNNs). The developed method uses a combination of Hebbian rules: spike-time dependent plasticity (STDP) and long-term depression (all-LTD). The synaptic weights distribution of a 3-layer SNN trained according to the combined rule allowed, after training, to restructure the synaptic weights distribution of SNN based on the principles of Systems Theory.

Keywords: *spiking neural network, supervised learning, spike-timing-dependent plasticity, all-time long-term depression, learning algorithm, Hebbian rule, Systems Theory, information entropy.*

1. Introduction

1.1. Systems Theory principles

A neural network is a very relevant example of a system. The latest generation of neural networks, spiking neural networks (SNN), demonstrates the behavior that corresponds to the main concepts of Systems Theory and, as L. von Bertalanffy wrote [1], to the provisions of disciplines that have common goals and methods with Systems Theory (Combinatorics, Information Theory).

Cl. Shannon introduced the concept of information entropy [2] as a measure of the uncertainty of a system for a random variable x taking n independent random values x_i with probabilities p_i :

$$H(x) = - \sum_{i=1}^n p_i \log_2 p_i, \quad (1)$$

where $H(x)$ is the amount of information in one event.

* The reported study was funded by the Russian Science Foundation (project number 24-21-00470).

A.N. Kolmogorov introduced a formal analog of the concept of information entropy for a continuous random variable ξ distributed over $X \subseteq R^n$ ($n < \infty$) [3], differential entropy, a functional defined on a set of absolutely continuous probability distributions:

$$H(\xi) = - \int_X f(x) \log_2 f(x) dx, \quad (2)$$

where $f(x)$ is the distribution density of a random variable.

Formulas (1) and (2) show that the minimum information is in a system with a simple structure, the ideal order corresponds to the minimum entropy.

In this way, a connection is established between the degree of information content and orderliness: the greater the degree of orderliness of information, the more it can be compressed. If it is possible to identify a regularity of data distribution, then you can significantly compress the amount of data.

One of the classical definitions describes a system (in Systems Theory) as a set of elements and their relationships, that is, a structure. The structure has a certain hierarchy: each component of the system can also be considered as a system (subsystem), and the properties of the individual components in total are not equivalent to the properties of the entire system.

Information in neural networks is contained in synaptic weights distributed in a certain structure. If the structure of synaptic weights is such that it can be described (in a whole or in a part) using some distribution function, then this means that it can be simplified without loss of system information.

1.2. Spiking Neural Networks learning rules

Research on SNN and the development of their learning rules are relevant today. Neuromorphic systems are bio-inspired and generally operate on the principles of SNN.

Neural network learning methods that have shown good performance for artificial neural networks (ANNs) cannot be directly used for spiking neural networks. The error backpropagation algorithm [4], which has become standard for ANNs, in the case of SNN requires special surrogates.

As we know, biological neurons learn locally rather than globally, hence gradient methods are biologically implausible. And what is more important, gradient methods require additional time for the correction signal to pass through.

Local learning rules based on Hebb's rule [5] are more suitable for training SNNs. Hebb's rule defines a learning process in which connections between neurons that fire one after another should be strengthened, and connections between neurons that fire independently of each other should be weakened. Local methods are bio-inspired but typically provide Unsupervised rather than Supervised Learning. Unfortunately, Unsupervised Learning is more suited to clustering

problems and is poorly suited to distinguishing between data with fairly similar features.

We developed a Supervised Learning method for SNNs on combined Hebbian learning rules [6], in which no correction signal is required.

2. The combined Hebbian learning rule and SNN architecture

2.1. Combined Hebbian learning rule, theory

The developed method for SNNs [9] is based on two bio-plausible learning rules: canonical Hebbian spike-timing-dependent plasticity (STDP) and all-time long-term depression (all-LTD).

The STDP rule adjusts the strength of synaptic connection w changes depending on the time difference in the excitation of spikes in the pre- and postsynaptic neurons $\Delta t = t_{post} - t_{pre}$ (Fig. 1a):

- if the presynaptic spike precedes the postsynaptic one (in a period of less than 50 ms) $\Delta t > 0$ and the connection is strengthened,
- if, on the contrary, the postsynaptic spike precedes the presynaptic one (in the same period) $\Delta t < 0$ and the synaptic connection decreases.

The change in synaptic connection is determined by the following expression:

$$\Delta w(\Delta t) = \begin{cases} A_{pre} \cdot \exp(-\Delta t/\tau_{pre}), & \Delta t > 0 \\ A_{post} \cdot \exp(\Delta t/\tau_{post}), & \Delta t < 0 \end{cases} \quad (3)$$

where $A_{pre} > 0$, $A_{post} < 0$.

The STDP mechanism provides the sensitivity of postsynaptic neurons to the specific features of the input signal and thus teaches a neuron to respond to a specific input.

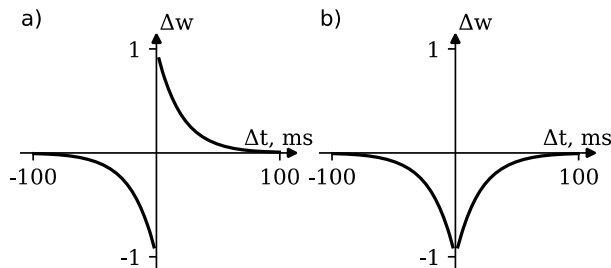


Fig. 1. Canonical Hebbian STDP (a) and all-LTD (b)

On the contrary, the all-LTD rule provides the ignoring the specific features of the input signal and thus teaches a neuron to remain silent to a specific input:

if pre- and postsynaptic spikes both fall within the interval $|\Delta t|$ (within the same period of 50 ms), the connection strength between the neurons decreases. This mechanism of synaptic connection change is determined by the following expression (Fig. 1b):

$$\Delta w(\Delta t) = \begin{cases} -A_{pre} \cdot \exp(-\Delta t/\tau_{pre}), & \Delta t > 0 \\ A_{post} \cdot \exp(\Delta t/\tau_{post}), & \Delta t < 0 \end{cases} \quad (4)$$

2.2. SNN's architecture

The rule was initially tested on a network based on the architecture proposed in [7] and later revised in [8].

The developed SNN uses the leaky integrate-and-fire model of neuron.

The original network in [6] has two layers: the first input layer contains Poisson neurons and the second layer consists of an equal number of excitatory and inhibitory neurons. To improve classification accuracy, a third layer was added to the network (Fig.2) with an equal number of excitatory and inhibitory neurons.

Each training image at the SNN input causes the next steps (Fig.2):

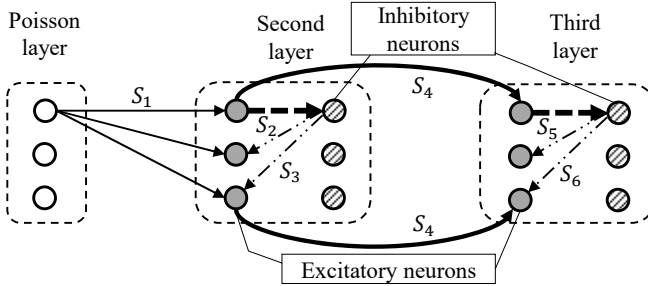


Fig. 2. SNN architecture: the first layer contains 784 Poisson neurons, the second layer contains 100 excitatory and 100 inhibitory neurons, the third layer has 30 excitatory and 30 inhibitory neurons; S1, S2, S3, S4, S5, and S6 are the synaptic connections

- Each Poisson layer neuron receives one pixel of input data, after which the neuron generates a spike train with an average frequency equal to the intensity of a given pixel (rate coding).
- The signal passes from Poisson neurons to excitatory neurons of the second layer through synapses of the S_1 group, connecting Poisson neurons with excitatory neurons of the second layer in one-to-all manner.
- The spike trains arrive at the excitatory neurons of the second layer and increase their membrane potential V , when potential V reaches the threshold level V_{thres} , the excitatory neuron generates spikes. Through the synapses

of the S_2 group, the spikes pass to the inhibitory neurons of the second layer, inducing spikes in them, since the synapses of the S_2 group connect them in a one-to-one manner. The weights of synapses in group S_2 remain constant.

- Through synapses of the S_3 group, the spikes generated by inhibitory neurons of the second layer return to excitatory neurons of the second layer according to the one-to-all-except-initiator manner. Synaptic connections of the S_3 group provide lateral inhibition. The weights of group S_3 remain constant.

- The spike trains arrive from the excitatory neurons of the second layer to the excitatory neurons of the third layer (increasing their membrane potential V to the threshold level V_{thres}) through the synapses of the S_4 group connecting them in a one-to-all manner.

- Subsequently, the signal goes through the synapses of groups S_5 and S_6 , providing communication similarly to the synapses of groups S_2 and S_3 , respectively:

- through the synapses of the S_5 group, the spikes pass to the inhibitory neurons of the third layer, connected in a one-to-one manner (with constant synapse weights),

- through synapses of the S_6 group, the spikes of inhibitory neurons of the third layer return to excitatory neurons of the third layer, connected in a one-to-all-except-initiator manner (with constant synapse weights).

During training, images are presented sequentially, after the spike train by a 350 ms (triggered by a single image) followed by a 150 ms period with no inputs. In the 150 ms resting period, the potential of the excitatory neurons drops to a lower threshold, after which the loop is repeated in the next image.

During the testing process, the class is determined by the highest activity of the neuron populations assigned to each class.

2.3. Training process

Training is carried out on S_1 and S_4 synaptic groups: S_1 group is trained according to the canonical STDP rule, which leads to the clustering of data features, and group S_4 is trained according to the combined Hebbian rule ' $STDP + all-LTD$ '.

To implement Supervised Learning, we divided the excitatory neurons of the third layer into subsets, the number of subsets is equal to the number of data classes, and the sizes of these subsets are the same. The subset number coincides with the class number that this subset classifies. If the spike reaches a subset of excitatory neurons in the third layer with the same number as the data class number, then this action leads to a correct classification.

Synapses of the S_4 group associated with excitatory neurons of the third layer are trained according to the rule ' $STDP + all-LTD$ '. When input data of a certain class (target) are given, synapses associated with a subset of neurons recognizing the target class are trained according to the STDP rule. The remaining synapses associated with neurons recognizing other classes are trained according to the all-LTD rule.

Two details are an important part of the learning rule implementation: lateral inhibition and adaptive threshold. Inhibitory neurons carry out lateral inhibition and thereby reveal a variety of class features: synapses associated with neurons corresponding to unique features are strengthened, and synapses associated with neurons corresponding to general features are weakened. The adaptive threshold allows the activity of excitatory neurons to be coordinated with the strength of the signs; it changes in proportion to the number of activations of a given neuron. Due to the adaptive threshold, not only strongly expressed features are highlighted, but also less pronounced features.

3. Experiment

3.1. Conditions of experiments

We used a computer with an Intel Core i9 processor (3.1 GHz), 16 GB RAM, PyTorch 1.8.0 and Ubuntu 18.04 to run the code. The SNN was implemented in the Brian 2.0 package, an open-source framework for SNN modeling.

The network was trained to classify data from the MNIST set. Due to the low speed of training of SNN on complete MNIST set, we used just 5000 images for the 1-st training step (S_1 training), 900 images for the 2-nd training step (S_4 training), and 300 images for the testing. In the experiments, we used only 3 classes of images out of 10 possible.

The SNN used consists of three layers, the first layer contains 784 Poisson neurons (one neuron per pixel of a 28×28 image), the second layer consists of 100 excitatory and 100 inhibitory neurons, the third layer has 30 excitatory and 30 inhibitory neurons (10 excitatory neurons per class).

In [6], the following classification accuracy was obtained on 2-layer SNN:

- (74.9 $\pm$ 19.8)% without adaptive threshold,
- (80.0 $\pm$ 13.0)% with an adaptive threshold.

A significant feature of such a network is its ability to learn rapidly: classification accuracy of about 80% is achieved when training the network with just 100 images.

To improve the quality of the network, we decided to add another layer. The first testing of the new 3-layer network showed classification accuracy close to the accuracy obtained as a result of random guessing of equally likely outcomes (for the 3-class option, the accuracy of random guessing is 1/3). The selection of

hyperparameters did not lead to a significant increase in accuracy, which was the reason for an in-depth analysis.

3.2. Preliminary analysis

Based on the results of the initial analysis, the main object of subsequent research was the synaptic weights of the S_4 group. The weight distribution of S_4 groups allows a clear visual representation on the graph (Fig. 3).

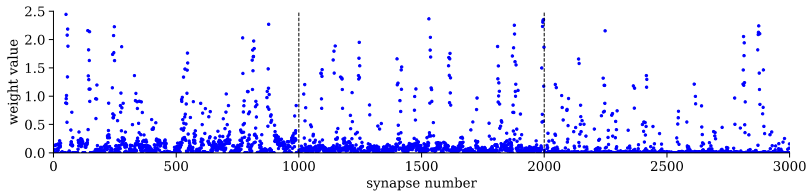


Fig. 3. Distribution of S_4 synaptic weights group after training; weight subgroups are separated by the vertical dashed lines

In the experiments, excitatory neurons of the third layer, according to the '*STDP + all-LTD*' rule, were divided into 3 subsets, each of which was responsible for classifying one of the 3 classes of data (the subset number coincided with the class number).

The synapses of S_4 group are associated to the three 10-element subsets ($10 \times 3 = 30$) of the third layer excitatory neurons.

Visually, Fig. 3 shows a densely distributed subsystems of lighter weights ('light' weights at the bottom of the graph) and a less densely distributed subsystems of heavier weights ('heavy' weights at the top).

The densities of the 'light' weight subsystem and the 'heavy' weight subsystem within the subgroup are quite constant. The analysis showed that different distributions of weights in subgroups are directly related to different degrees of class definition accuracy.

Each weights subgroup is associated with the definition of one class, and the best classification accuracy was shown by the subgroup in which:

- 'light' weights are distributed so that they occupy a relatively wider range,
- the distribution of 'light' weights (without taking into account weights with 0 value) is close to a uniform distribution,
- the sum of the 'heavy' weights is maximal,
- dividing the weights into 'heavy' and 'light' ones is most effective if the number of 'heavy' weights is about 10% of the number of all weights (the

rest are the ‘light’ ones). This result is consistent with our previous research regarding the importance of weights for classification accuracy [9].

The best value for the point separating ‘light’ and ‘heavy’ weights was determined as the inflection point of the weight density distribution curve of the subgroup that had the maximum sum of ‘heavy’ weights (the inflection point mentioned below refers specifically to this class).

3.3. Confirming assumptions based on the principles of Systems Theory

Based on the results of the analysis, we suggested that ‘light’ and ‘heavy’ weights play different roles in the classification process (taking into account the probabilistic nature of the process): ‘light’ weights maintain the potential level in excitatory neurons, ‘heavy’ weights directly fire excitatory neurons.

The suggestion about the different roles of weights allowed us to draw an analogy and assume using the base concepts of Systems Theory: the division of weights into ‘light’ and ‘heavy’ corresponds to their level of information content. The ‘heavy’ weights subsystem contains almost all the information of the weights system; in the ‘light’ weights subsystem, information entropy is close to 0, which makes it possible to radically simplify the ‘light’ weights subsystem.

To confirm the assumption of information redundancy (on 3-layer SNN with adaptive threshold), two series of experiments were carried out, during which, on the one hand, the number of ‘heavy’ weights in subgroups and the total sum of ‘heavy’ weights in subgroups were equalized (by compressing their structure upward), and on the other hand, a simplification of ‘light’ weight subsystems was carried out:

- in the 1-st series, ‘light’ weights (obtained as a result of training) were replaced by random variables uniformly distributed from 0 to the inflection point (Fig. 4); on the graph, ‘light’ weights are presented in the form of an equal-wide ‘foaming’ strip obtained accuracy $(68.4 \pm 4.9)\%$,

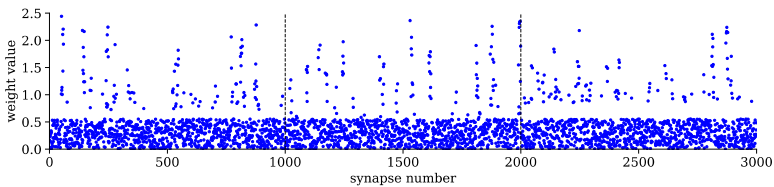


Fig. 4. Distribution of S4 group, where ‘heavy’ weights were equalized, ‘light’ weights were replaced by uniformly distributed random variables

- in the 2-nd series, ‘light’ weights were replaced by a single value (Fig. 5); on the graph, ‘light’ weights are represented as a horizontal line accuracy $(67.8 \pm 3.1)\%$.

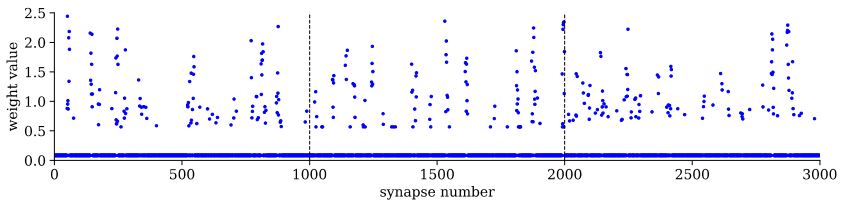


Fig. 5. Distribution of S_4 group, where ‘heavy’ weights were equalized, values of ‘light’ weights were replaced with a single value

Simplification of the ‘light’ weights structure together with the compression of the ‘heavy’ weights structure led to an increase in accuracy, which shows the correctness of the assumption about the information content levels of subgroups.

We think that this indicates that ‘light’ weights, having minimal information content, play a significant role in maintaining a certain level of potential in the excitatory neurons of the third layer. To confirm this assumption about the role of ‘light’ weights, we conducted two more series of experiments (on SNN with adaptive threshold) in which:

- the values of the ‘light’ weights were replaced with the 0-th value, and the structure of the ‘heavy’ was compressed, which led to a drop in accuracy to $(64.0 \pm 4.5)\%$,
- all ‘light’ weights were replaced by one value, and all ‘heavy’ weights were also replaced by one value, which led to a drop in accuracy to $(50.5 \pm 6.7)\%$.

The results of the new experiment series, in addition to confirming the role of ‘light’ weights, showed that the ‘heavy’ weights subsystem contains a significant part of system information, but the help of ‘light’ weights in maintaining a certain level of potential of excitatory neurons significantly increases classification accuracy.

4. Conclusion

The system of synaptic weights was obtained as a result of training a 3-layer SNN with the ‘*STDP + all-LTD*’ learning rule. Only the first part of the 3-layer SNN research was carried out. For the present this approach has allowed us to obtain only a quality understanding of the synaptic weights system operation and, based on the principles of systems theory, to develop the general scheme. Further research will focus on practical aspects of network design (improving SNN accuracy).

The author would like to thank Dr. S. Sukhov for careful reading of this manuscript and for constructive comments.

References

1. Bertalanffy L.V. General System Theory — A Critical Review // General Systems. 1962. V. VII. P. 1–20.
2. Shannon C.E. A Mathematical Theory of Communication // Bell System Technical Journal. 1948. V. 27, N 3. P. 379–423.
3. Kolmogorov A.N. Information theory and algorithm theory. Moscow: Nauka. 1987.
4. Werbos P.J. Applications of advances in nonlinear sensitivity analysis // Drenick R.F., Kozin F. (eds) System Modeling and Optimization. Lecture Notes in Control and Information Sciences. Springer. Berlin. Heidelberg. 1982. V. 38. P. 762–770.
5. Hebb D.O. Organization of behavior. New York : Wiley and Sons. 1949.
6. Antonov D.I., Sukhov S. Spiking neural networks-classifiers training with local rules // Proceedings of XXV International Conference on Artificial Neural Networks Neuroinformatics-2023. M.: National Research Nuclear University MEPhI. 2023. P. 116–125.
7. Diehl P.U., Cook M. Unsupervised learning of digit recognition using spike-timing-dependent plasticity // Front. Comput. Neurosci. 2015. V. 9, N 99.
8. <https://www.kaggle.com/code/dlarionov/mnist-spiking-neural-network/notebook>
9. Antonov D.I., Sviatov K.V., Sukhov S. Continuous learning of spiking networks trained with local rules // Neural Networks. 2022. V. 155. P. 512–522.

ГЛУБОКОЕ ОБУЧЕНИЕ

V.V. KUZNETSOV, V.A. MOSKALENKO, N.YU. ZOLOTYKH

Institute of Information Technologies, Mathematics and Mechanics

Lobachevsky State University of Nizhni Novgorod

{vladislav.kuznetsov,viktor.moskalenko,nikolai.zolotikh}@itmm.unn.ru

DIAGNOSIS OF CARDIOVASCULAR DISEASES USING RECURRENT AND CONVOLUTIONAL NEURAL NETWORKS

In this work, we explore methods for diagnosing cardiovascular diseases based on electrocardiogram (ECG) data, using neural networks, as well as machine learning methods. We will use convolutional and recurrent neural networks, as well as the XGBoost algorithm. Our goal is to build a model that can identify a large set of cardiovascular diseases using a signal and segmentation built on that signal with a high degree of accuracy, as well as additionally calculated features.

Keywords: *diagnostics, convolutional neural network, recurrent neural network, ECG, XGBoost, electrocardiography, cardiovascular system.*

1. Introduction

We want to offer automated methods for determining heart pathologies to help and correct the work of cardiologists.

We will examine the patient's condition on the basis of an ECG. On each ECG, 5 main waves can be distinguished: P, Q, R, S, T, which describe the heart-beat [1].

By using the ECG as a sequence of signal values, we want neural networks to learn to make diagnoses just like doctors do. Machine learning (deep learning) methods are widely used to study ECG diagnostics; see recent review by [2].

We believe that this direction is quite promising, because If this problem is solved, diagnosing diseases will become easier and faster. In addition, neural networks will help in determining diagnoses that require large financial and human resources, thereby making medicine more accessible to people.

2. Algorithm

2.1. Preprocessing

The raw data consists of 18,885 12-lead ECG recordings of patients [3]. They also contain diagnoses for each disease. By extracting each, we find that our sample is 21,837 12-lead, 10-s ECGs, which we normalize to a sampling frequency of 500 Hz. Preprocessing consists of normalizing the raw data. To do

this, we calculate the arithmetic mean of the signals for each lead. We calculate the difference between the raw signal values and their mean value and divide the result by the standard deviation for each lead. This way, we get normalized data ready to work with models.

2.2. Convolutional neural network architecture

Our convolutional neural network has 2 inputs. The first input is a sequence of size $(N, 12)$, where N is the length of the signal entering the input, and 12 is the number of leads. The second input is the output from another neural network, whose task is ECG segmentation [4], has a size of $(N, 4)$, where 4 is a value reflecting whether a given signal value belongs to the P, QRS, T segments or none of the above. Specifically, in our case, we chose the value of N - 4500 and 3000, i.e. 6 and 9 second signal with a sampling rate of 500 Hz.

Next, from the first inputs follow 2 blocks, each of which contains a convolutional layer, with the help of which we strive to identify some features and properties in our sequences of signals and segmentations. Each convolutional layer has the same kernel size parameter: 9. It is convenient to write the number of filters in layers in the form of a list. For a 12-lead input, this list looks like this:

$filters = [32, 32]$, where $filters$ reflect the dimension of the output space.

For an input belonging to one or another ECG segment (segmentation), this list looks like $filters = [32, 32]$.

We obtain the output data tensors for each filter. The depth of the set of tensors will be the number of filters.

Next comes *Dropout* – the main tool for combating overfitting in our network. In our case, the value of discarded neurons is 0,25 of the number of all available ones.

This is followed by another convolutional layer with the number of filters from $filters$. Next, in each of the 2 blocks, after the convolutional layer, there is a layer with normalization.

And the last element of each of the 2 blocks: the ReLU activation function. This is a kind of rectifier function.

And after these blocks comes the *Flatten* layer, which reduces the sequence we received to a vector. Next, the resulting vectors are concatenated.

Then branching occurs depending on the task at hand. In our case, from the concatenation of two vectors we have 6 branches to determine:

1. Sinus rhythm, fibrillation, flutter, or other non-sinus rhythm;
2. The presence or absence of any hypertrophy;
3. The presence or absence of any extrasystole;
4. The presence or absence of any AV block;
5. Presence or absence of bundle branch block;

6. Determination of the electrical axis of the heart.

Each branch contains 2 fully connected layers with certain weights. The first layer gives us the output of 512 neurons with the ReLU activation function, and the last regular dense connected layer allows us to organize the output of a certain number of neurons taking into account the activation function. The final result is formed based on the kernel - the weight matrix and the displacement vector.

The activation functions in this case on the last layer may differ. If there is 1 neuron at the output, then we use *Sigmoid*, if there are more (4 for determining rhythms, 9 for electrical axes), then *Softmax*.

In total, the network has 6 outputs for determining a particular diagnosis. Let's list them:

- 1. Presence of sinus rhythm, fibrillation, flutter, or other non-sinus rhythm;
- 2. The presence or absence of any hypertrophy;
- 3. The presence or absence of any extrasystole;
- 4. The presence or absence of any AV block;
- 5. Presence or absence of bundle branch block;
- 6. Determination of the electrical axis of the heart.

The first output for rhythms has 4 neurons, the last one for electrical axes has 9 neurons. In other cases, the output is 1 neuron. Representation of convolution neural network architecture is shown in Figure 1.

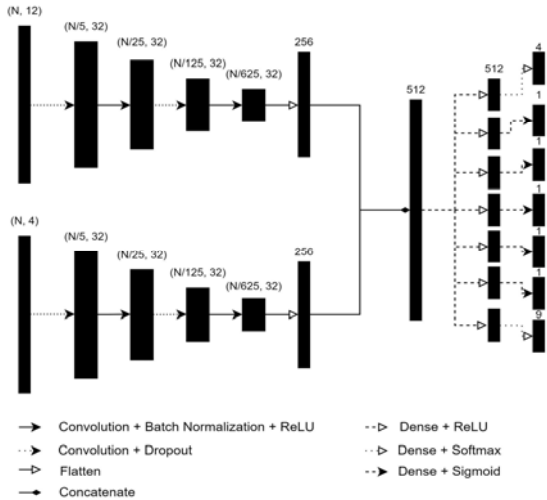


Fig. 1. Convolutional neural network architecture

2.3. Recurrent neural network architecture

Since we are working with time series, where elements of a sequence are related to each other, we hypothesize the use of recurrent levels of a neural network in diagnostics.

A recurrent neural network has 3 inputs. The first two inputs are identical to the convolutional neural network. The third input required to determine the electrical axis of the heart is an array of 60 elements (maximum, minimum, mean, standard deviation and sum for each lead).

From the first two inputs there are 9 blocks, each of which contains a convolutional layer. Each convolutional layer had the same kernel size parameter: 3. We again indicate the number of filters in the form of a list: *filters* = [128,256,256,256,256,512,512,512,512].

For an input belonging to a particular ECG segment, this list looks like: *filters* = [16,32,32,32,32,64,64,64,64].

In each of the 9 blocks, after the convolutional layer, there is a normalization layer. Next comes Max-pooling, which is a sampling process based on our sample. Max-pooling in our case always divides the sample into blocks of two elements and selects the maximum value from them.

Next comes dropout. In this network, the number of discarded neurons is 0,3 of the number of all available ones.

And the last element of each of the 9 blocks: the ReLU activation function. After all 9 blocks are executed, the recurrent LSTM layer follows.

After this, we normalize the resulting vectors using Batch Normalization.

Next, the resulting vectors are concatenated. To determine the electrical axis, a third input from a 60-element vector is added to these vectors.

Then there is a branching depending on how we pose the problem. In the most general case, from the concatenation of two vectors, we have 5 branches for determining the diseases described above for the convolutional network, with the exception of the electrical axes of the heart.

The first 2 layers give us an output of 576 and 278 neurons, respectively, with the ReLU activation function. Similarly with a convolutional network, the last regular dense connected layer allows us to organize the output of a certain number of neurons taking into account the activation function. The final result is formed based on the kernel - the weight matrix and the displacement vector. The activation function also differs from the task as before.

The concatenation of 3 vectors also results in 3 fully connected layers. However, the first 2 layers now give us an output of 600 and 300 neurons, respectively, with the ReLU activation function. The last layer gives us the output of 9 neurons with the Softmax activation function. In this case, the electrical axis of the heart is determined.

Representation of convolution neural network architecture is shown in Fig. 2.

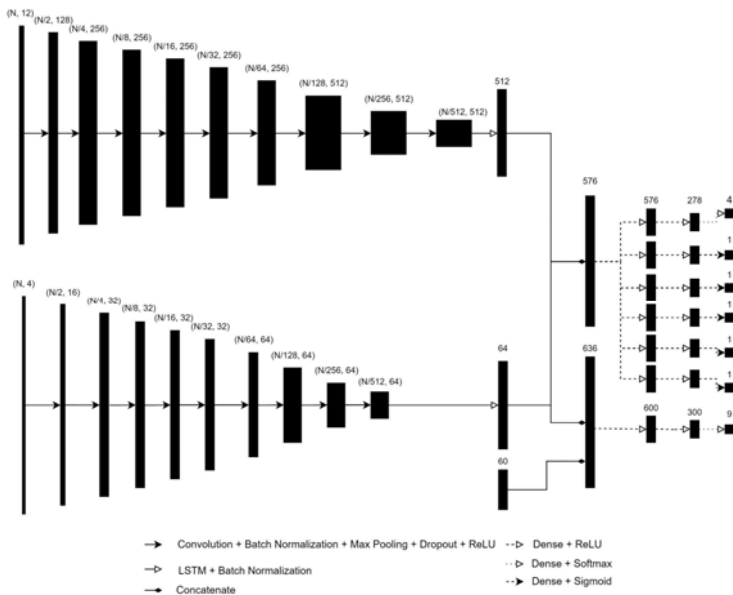


Fig. 2. Recurrent neural network architecture

2.4. XGBoost preprocessing

We used a ready-made machine learning algorithm from the library xgboost to analyze and evaluate our results and the results of the ready-made algorithm. In this experiment, unlike all previous ones, we will not use a sequence of signals, but a set of characteristics that was calculated in advance. The characteristics are vectors of 696 values. Characteristics are not just numbers, but special indicators that reflect the work of the heart. These include lead averages, wave maxima, wave minima, standard deviations and much more.

3. Experimental results

3.1. Preparing training data

The original data set consists of 21837 12-lead ECGs with a duration of 10 seconds. For each ECG there is a marking that reflects whether each signal belongs to the P, QRS, T segments or none of the above [4]. After the preliminary data processing described above, we divide our total sample into training and testing in a percentage ratio of 70 to 30.

Augmentation was used to prevent overfitting [5]. We will not submit the entire 10 second signal to the input of the neural network, but only part of it and draw conclusions based on this part. We will “cut out” 6 and 9 second signals from the source data. To determine the diagnoses we have planned, ECG trimming data is acceptable. At 6 and 9 seconds the result will be similar to the result at 10 seconds. Thus, we expand the maximum allowable set of input sequences. Hence, it will be quite problematic for us to retrain.

Let us add that many tasks have a large imbalance between classes. Therefore, our task is to distribute weights to classes in order to “equalize” them. We do this by proportionally dividing the total number of instances by the number of classes and the number of representatives of a particular class for which the weight is calculated.

3.2. Quality metrics

The null hypothesis must be defined. In our case, this is: "The patient has some (depending on the task) disease." In the case of a binary problem, we accept or reject this hypothesis.

If we accept the hypothesis and the patient is really sick, then this situation is called True Positive (TP). If we reject the hypothesis and the patient is truly healthy, then this is a case of True Negative (TN). If we accepted the hypothesis and the patient is healthy, then this is False Positive (FP) (or a type 2 error). If the hypothesis is rejected and the patient is sick, then this is the worst case of False Negative (FN) (type 1 error) [5].

To evaluate the results of diagnostic effectiveness, we used the following statistical indicators: accuracy, precision, recall (sensitivity), specificity, F_1 -score.

Additionally, it is worth noting that statistical indicators of diagnostic efficiency in multi-class classification problems have a slightly different form. They are divided into micro-average, where the contribution of each class is taken into account and then the average is calculated, and macro-average, where the metric is calculated independently for each class and the average is taken.

3.3. Experiments

To train the model we used 300 epochs of Adaptive Moment Estimation (Adam) algorithm and implemented in TensorFlow Framework with data augmentation.

Let's start by looking at heart rhythms. This is a multiclass task with 4 classes: the presence of sinus rhythm, fibrillation, flutter, or other non-sinus rhythm. As

a result of network training for this task, there are 4 mutually exclusive diagnoses. The results are in Table 1.

Table 1

Rhythm diagnostic results

Metrics	Convolution NN				Recurrent NN				XGBoost	
	6 sec.		9 sec.		6 sec.		9 sec.			
	Train	Test	Train	Test	Train	Test	Train	Test	Train	Test
Accuracy	0,936	0,91	0,978	0,91	0,652	0,729	0,838	0,774	1	0,928
Precision (micro)	0,936	0,91	0,978	0,91	0,652	0,729	0,838	0,774	1	0,928
Recall (micro)	0,936	0,91	0,978	0,91	0,652	0,729	0,838	0,774	1	0,928
Specificity (micro)	0,979	0,97	0,993	0,97	0,884	0,91	0,946	0,925	1	0,976
F-measure (micro)	0,936	0,91	0,978	0,91	0,652	0,729	0,838	0,774	1	0,928
Precision (macro)	0,882	0,72	0,955	0,725	0,434	0,402	0,68	0,506	1	0,899
Recall (macro)	0,936	0,76	0,986	0,648	0,544	0,538	0,85	0,635	1	0,74
Specificity (macro)	0,974	0,956	0,992	0,672	0,892	0,362	0,949	0,92	1	0,945
F-measure (macro)	0,907	0,673	0,97	0,955	0,38	0,908	0,729	0,517	1	0,794

Table 1 shows that the best results were obtained from a convolutional neural network with a signal length of 6 seconds. The values range from 75 to 95%. These are very high indicators, which indicate a good diagnosis of 4 classes of heart rhythm. Let us note right away that the results of the XGBoost algorithm are close to the results of a convolutional neural network. Note that the classes of sinus rhythm and other non-sinus rhythms are most confused. Classifier errors in determining these diagnoses lead to the greatest decrease in metric values.

Table 2

Results of hypertrophy diagnostics

Metrics	Convolution NN				Recurrent NN				XGBoost	
	6 sec.		9 sec.		6 sec.		9 sec.			
	Train	Test	Train	Test	Train	Test	Train	Test	Train	Test
Accuracy	0,81	0,732	0,953	0,714	0,676	0,636	0,788	0,639	0,935	0,768
Precision	0,846	0,81	0,965	0,795	0,709	0,758	0,88	0,761	0,912	0,844
Recall	0,836	0,813	0,957	0,805	0,785	0,716	0,749	0,716	0,986	0,826
F-measure	0,812	0,91	0,961	0,8	0,745	0,736	0,81	0,738	0,948	0,835
Specificity	0,77	0,531	0,947	0,49	0,51	0,438	0,846	0,448	0,856	0,625

Table 2 presents the results of diagnosing hypertrophies. It is worth noting that the best performance was again achieved by a convolutional neural network with a signal length of 6 seconds. On average, diagnostic indicators give an 80% success rate for predictions. This is a fairly high result. Also, the result of XGBoost is very close to the value of the convolutional neural network.

Table 3

Results of diagnostics of extrasystole

Metrics	Convolution NN				Recurrent NN				XGBoost	
	6 sec.		9 sec.		6 sec.		9 sec.			
	Train	Test	Train	Test	Train	Test	Train	Test	Train	Test
Accuracy	0,974	0,871	0,992	0,783	0,842	0,675	0,975	0,825	1	0,931
Precision	0,626	0,28	0,845	0,177	0,085	0,122	0,895	0,215	1	0,475
Recall	0,931	0,667	0,986	0,667	0,278	0,667	0,472	0,667	1	0,905
<i>F</i> -measure	0,749	0,394	0,91	0,28	0,13	0,206	0,618	0,326	1	0,623
Specificity	0,975	0,884	0,992	0,791	0,867	0,675	0,998	0,836	1	0,933

Table 3 presents the results of diagnosing ecstasystoles. Here, relatively good results are obtained from a convolutional neural network with a signal length of 6 seconds. Unfortunately, in this case the precision is low. This may have happened due to inaccuracies in the markup. In any case, our research in this area will continue to improve the quality of individual metrics and improve diagnostics.

Table 4

Results of diagnostics of AV-blocades

Metrics	Convolution NN				Recurrent NN				XGBoost	
	6 sec.		9 sec.		6 sec.		9 sec.			
	Train	Test	Train	Test	Train	Test	Train	Test	Train	Test
Accuracy	0,891	0,91	0,931	0,91	0,442	0,572	0,788	0,798	0,999	0,952
Precision	0,317	0,324	0,424	0,324	0,053	0,057	0,189	0,123	1	0,542
Recall	0,977	0,611	1	0,611	0,586	0,444	0,954	0,444	0,989	0,722
<i>F</i> -measure	0,479	0,423	0,596	0,423	0,097	0,101	0,316	0,193	0,994	0,619
Specificity	0,887	0,927	0,927	0,927	0,434	0,58	0,78	0,819	1	0,965

A convolutional neural network detects AV block better than a recurrent one. This is proven by the good values of the quality metrics in Table 4. Unfortunately, here again the precision is not the highest. However, the presence of diseases is predicted very well.

Table 5

Results of diagnosis of bundle branch block

Metrics	Convolution NN				Recurrent NN				XGBoost	
	6 sec.		9 sec.		6 sec.		9 sec.		Train	Test
	Train	Test	Train	Test	Train	Test	Train	Test		
Accuracy	0,953	0,952	0,982	0,958	0,915	0,922	0,976	0,871	1	0,991
Precision	0,385	0,25	0,617	0,278	0,23	0,167	0,565	0,089	1	0,714
Recall	1	0,833	0,957	1	0,833	0,8	0,78	0,667	1	0,833
<i>F</i> -measure	0,556	0,385	0,763	0,417	0,357	0,278	0,656	0,157	1	0,769
Specificity	0,952	0,954	0,981	0,96	0,919	0,923	0,982	0,874	1	0,994

Bundle branch block is best identified by a convolutional neural network. However, we have a precision. XGboost has similar results, but with a much higher precision. This follows from Table 5.

Table 6

Electric axis diagnostic results

Metrics	Convolution NN				Recurrent NN				XGBoost	
	6 sec.		9 sec.		6 sec.		9 sec.		Train	Test
	Train	Test	Train	Test	Train	Test	Train	Test		
Accuracy	0,938	0,627	0,945	0,6	0,851	0,542	0,896	0,56	0,962	0,654
Precision (micro)	0,938	0,627	0,945	0,6	0,851	0,542	0,896	0,56	0,962	0,654
Recall (micro)	0,938	0,627	0,945	0,6	0,851	0,542	0,896	0,56	0,962	0,654
<i>F</i> -measure (micro)	0,938	0,627	0,945	0,6	0,851	0,542	0,896	0,56	0,962	0,654
Specificity (micro)	0,991	0,947	0,992	0,943	0,979	0,935	0,985	0,937	0,995	0,951
Precision (macro)	0,921	0,437	0,916	0,393	0,786	0,385	0,815	0,431	0,982	0,475
Recall (macro)	0,973	0,462	0,974	0,459	0,859	0,418	0,89	0,451	0,97	0,395
<i>F</i> -measure (macro)	0,944	0,447	0,942	0,415	0,814	0,391	0,837	0,418	0,975	0,413
Specificity (macro)	0,991	0,94	0,992	0,935	0,977	0,929	0,984	0,928	0,993	0,939

Electrical axes are best determined by a convolutional neural network with an input signal length of 6 seconds (Table 6). Of the 9 diagnoses, 4 stand out most clearly here, which make the main contribution to the value of the metrics. The horizontal and normal electrical axes of the heart are most confused, which makes the greatest contribution to the deterioration of metric indicators. In general, they are defined quite well. The XGBoost algorithm gives similar results.

4. Conclusions and further research

As a conclusion, we note that we have created convolutional and recurrent neural networks for diagnosing 6 classes of various diseases with different signal lengths. We can determine rhythm with rates exceeding 90%, hypertrophy with rates exceeding 81%, premature beats with a rate of ~67%, AV heart block with a probability of 61%, electrical axis of the heart with a rate of ~60% and some of the bundle branch blocks with very good performance.

Thus, we can already determine some diagnoses with good probability (rhythms, hypertrophies, some bundle branch blocks). The remaining diagnoses should be determined more accurately by improving the models and adjusting the labeling of the source data.

If we compare the results we obtained, we come to the conclusion that the overall most successful diagnostic neural network is a convolutional one. The diagnosis rates were approximately 63%. We expected that the recurrent neural network would be able to compete with convolutional ones. However, this did not happen and recurrent networks turned out to be an order of magnitude worse than convolutional ones.

Further work will be devoted to solving the problem of diagnosing other cardiovascular diseases and improving the diagnosis of those diseases that are already determined by the network (electrical axes, AV blockades, etc.). Our convolutional neural network will be used as the basis.

References

1. Hampton J., Hampton J. The ECG Made Easy. Elsevier (2019).
2. Hong S., Zhou Y., Shang J., Xiao C., Sun J. Opportunities and challenges of deep learning methods for electrocardiogram data: A systematic review. *Computers in Biology and Medicine*, 103801 (2020). doi: 10.1016/j.combiomed.2020.103801.
3. Wagner P., Strodthoff N., Boussejot R.-D., et al. PTB-XL: A Large Publicly Available ECG Dataset. *Scientific Data* (2020). doi: 10.1038/s41597-020-0495-6.
4. Moskalenko V.A., Zolotykh, N.Yu., Osipov G.V. Deep Learning for ECG Segmentation. *International Conference on Neuroinformatics*, 246–254. Springer (2019). doi: 10.1007/978-3-030-30425-6_29.
5. LeCun Y., Bengio Y., Hinton G. Deep learning. *Nature*, 521, 436–444. Nature Publishing Group (2015). doi: 10.1038/nature14539.
6. Kingma D.P., Ba J. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980* (2014).

Е.П. ВАСИЛЬЕВ, Д.А. ЕРМОЛАЕВ, А.Н. ЕРЕМИН

Национальный исследовательский Нижегородский государственный университет
им. Н.И. Лобачевского
evgeny.vasiliev@itmm.unn.ru

РАЗРАБОТКА ИНСТРУМЕНТАРИЯ ДЛЯ ИССЛЕДОВАНИЯ X-ХРОСОМНОГО МОЗАИЦИЗМА НА ПРЕПАРАТАХ ЯДЕР КЛЕТОК КРОВИ ПО ДАННЫМ КОНФОКАЛЬНОГО МИКРОСКОПА*

Флуоресцентная гибридизация *in situ*, или FISH, представляет собой цитогенетический метод, используемый для обнаружения последовательности ДНК в интерфазных ядрах *in situ*. В данной работе предложен метод автоматической сегментации объектов препаратов ядер клеток крови и метод обнаружения отдельных хромосом, а также разработано программное обеспечение для обработки данных, позволяющее как интерактивную обработку изображений с использованием удобного графического интерфейса, так и пакетную обработку данных.

Ключевые слова: *флуоресцентная гибридизация, сегментация, глубокое обучение, обработка изображений.*

Введение

В настоящее время все чаще возникают проблемы, требующие разработки программного обеспечения, способного эффективно анализировать цифровые изображения. Одним из вариантов решения этой проблемы является подход, заключающийся в четком разделении объектов на изображении – сегментации с целью анализа каждого объекта отдельно. Использование сегментации полезно в различных областях: медицинской диагностике [1], анализе структурных изменений образований, например, при сравнении различных генных проявлений мутаций внутри определенной популяции. Сегментация объектов (Instance segmentation) – это подход компьютерного зрения, который позволяет автоматически выбирать отдельные объекты на изображении и определять их границы. Сегментация объектов является более продвинутой методикой, чем семантическая сегментация, поскольку позволяет разделять объекты на изображении не только по цвету и яркости, но также по форме и текстуре, а также распознавать и классифицировать эти объекты.

* Работа выполнена при финансовой поддержке Минобрнауки РФ, проект № FSWR-2024-0005.

Перспективной задачей использования сегментации экземпляров является анализ FISH-изображений. Флуоресцентная гибридизация *in situ* (FISH) [2] – это метод, который можно использовать для определения местоположения генетического материала в клетках вплоть до отдельных генов и хромосом и выявления генетических аномалий. Флуоресцентно меченные зонды ДНК или РНК используются для связывания с комплементарными целевыми последовательностями в образце. Используя FISH-метод, исследователи могут визуализировать пространственное расположение конкретных генов или хромосомных областей, получая ценную информацию об экспрессии генов, хромосомных аномалиях и геномных вариациях [3]. Метод широко используется в различных областях, включая генетику, исследования рака, пренатальную диагностику и микробную экологию [4]. FISH может выполняться на фиксированных клетках или срезах тканей и обеспечивает детальное понимание генетического состава и расположения биологического образца. Важным аспектом FISH-анализа является выявление и изучение хромосомного мозаицизма – наличия в организме двух и более различных хромосомных популяций клеток, возникающих в результате генетических изменений, затрагивающих только часть клеток [5]. Хромосомный мозаицизм может возникать на ранних стадиях эмбрионального развития или на более позднем этапе жизни из-за генетических мутаций, хромосомных перестроек или ошибок деления клеток. Понимание хромосомного мозаицизма важно для медицинского ведения, генетического консультирования и индивидуального лечения.

Описание исходных данных

Для FISH-анализа были использованы цитогенетические препараты лимфоцитов периферической крови и клеток эпителия. Снимки были сделаны с помощью микроскопа в лаборатории Института биологии и биомедицины ННГУ им. Н.И. Лобачевского. Изображения в основном имеют размер 512x512 пикселей и разрешение 0.12–0.69 нанометров на пиксель. Доступный набор изображений содержит около 100 изображений в форматах файлов CZI, BMP и JPG. Формат CZI был разработан компанией Zeiss для хранения и обработки микроскопических изображений [6]. Он содержит массив изображений и метаданные (модель и тип микроскопа, разрешение изображения, размер и масштаб, контрастность и используемые линзы и фильтры, параметры камеры и детектора (например, значения экспозиции, битовая глубина, время пребывания пикселя и т. д.). Данная информация также необходима при обработке изображений, полученных с помощью FISH-анализа. Пример получаемых данных представлен на рис. 1.

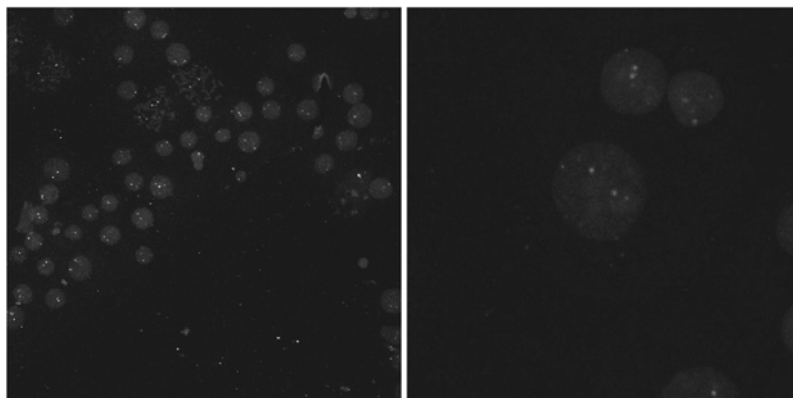


Рис. 1. Пример изображений FISH-анализа. Светлые области – целые и разрушенные клетки, яркие точки – X-хромосомы и 21 хромосомы

Оригинальное изображение с конфокального микроскопа является цветным, и содержит три канала (красный, зеленый, синий), на оригинальном цветном изображении яркие точки в красном канале представляют расположение 21 хромосомы, а яркие точки в зеленом канале представляют собой X-хромосомы.

Постановка задачи и описание алгоритма

В данной работе рассматривается задача разработки программного обеспечения для исследования X-хромосомного мозаицизма на препаратах ядер клеток крови по данным конфокального микроскопа. Данная задача включает в себя несколько этапов:

1. Сегментация изображений FISH-анализа на отдельные клетки.
2. Детектирование X-хромосом и 21 хромосом на изображении.
3. Построение взаимосвязей между клетками и хромосомами, подсчет численных характеристик для изображения.
4. Вывод численных характеристик в формате, удобном для исследователя.

Сегментация целых и разрушенных ядер клеток

В качестве алгоритма для решения задачи сегментации экземпляров будет использоваться семейство алгоритмов глубокого обучения YOLO [7].

YOLO (You Only Look Once) – это семейство одноэтапных моделей обнаружения объектов на изображении, которые за один проход через сеть выполняют и генерацию предположений о расположении объектов, и классификацию полученных предположений в отличие от двухэтапных моделей, таких как Faster R-CNN [8], которые сначала генерируют предложения по регионам, а затем классифицируют их. Модели YOLOv5 и YOLOv8 делят изображение на сетку и одновременно предсказывают ограничивающие прямоугольники и классы объектов для каждой ячейки сетки. Такой подход обеспечивает значительное увеличение скорости работы модели без существенного снижения точности.

Архитектура модели YOLOv5 приведена на рис. 2, модель состоит из следующих блоков:

1. Сверточные блоки, которые уменьшают размер изображения и обучаются признакам на изображении. Они состоят из слоев Conv, пакетной нормализации (BN) и функции активации Swish Linear Unit (SiLU).
2. Блоки C3, включающие несколько модулей Bottleneck, улучшающих способность модели к обобщению. Модули Bottleneck содержат как skip connections, так и свертки, разные типы слоев в модуле помогают предотвратить проблемы с градиентным затуханием и улучшить передачу информации между блоками сети.
3. Блок SPPF (SPP, Spatial Pyramid Pooling-Fast, быстрое объединение пространственных пирамид). Последний уровень архитектуры состоит из блока SPPF, который объединяет функции из разных слоев путем применения операций max pooling с разными размерами окон. Это позволяет модели YOLOv5 эффективно интегрировать информацию из изображения с разных масштабов и улучшает способность модели обнаруживать объекты разных размеров.
4. Блоки Upsample и Concat. Эти блоки позволяют объединять выходы других блоков с разных уровней сети, улучшая контекстное понимание изображения и повышая точность прогнозов.

Развитием семейства YOLO является модель YOLOv8 [9]. Архитектура модели YOLOv8 схожа с архитектурой модели YOLOv5, изменения касаются внутреннего строения блоков с целью увеличения производительности модели без ухудшения качества ее работы.

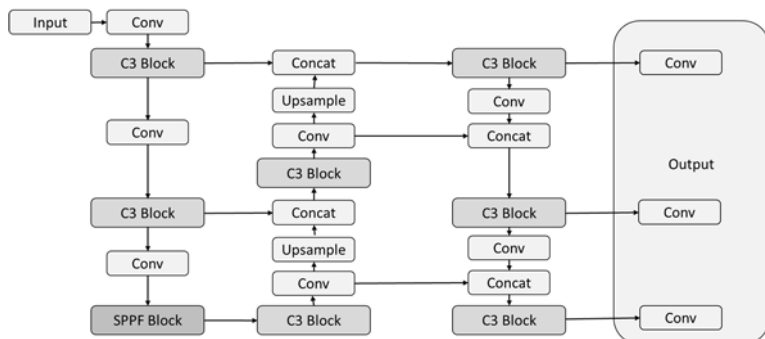


Рис. 2. Архитектура модели YOLOv5

Для обучения модели детектирования клеток была выполнена ручная разметка снимков с использованием сервиса Roboflow [10]. Для создания датасета и разметки ядер на каждой фотографии было выбрано 36 изображений, содержащих целые и разрушенные ядра клеток. В графическом интерфейсе Roboflow выделялись целые ядра клеток и разрушенные ядра клеток (рис. 3).

Далее данные распределяются между обучающим, валидационным и тестовым наборами. В документации YOLO рекомендуется, чтобы валидационный набор содержал не менее 20% от общего количества изображений.

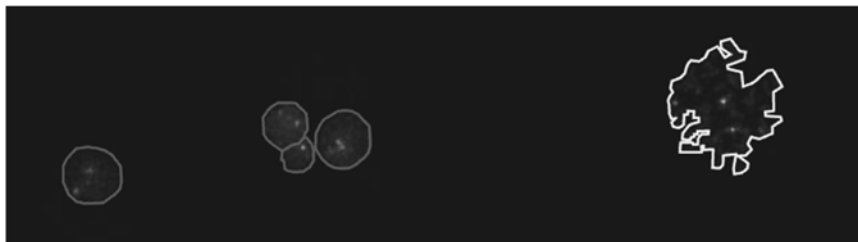


Рис. 3. Пример ручной разметки целых и разрушенных ядер в Roboflow

Для данных тренировочной и валидационной выборки была выполнена аугментация изображений средствами Roboflow: поворот изображения, изменение яркости изображения, Exposure, перевод в оттенки серого, перевод изображения в HSV и вариации в Hue, Saturation каналах, мозаичную аугментацию [11]. Аугментации применяются при создании датасета, после их применения датасет включает 87 изображений в тренировочной части, 20 в валидационной и 20 в тестовой.

Обучение сегментационных YOLOv5 и YOLOv8 моделей выполнялось в Google Colab при помощи GPU Tesla T4. Было выполнено 100 итераций обучения, для подсчета метрик качества выбрана итерация с наилучшими метриками на валидационном датасете. Метрики качества представлены в таблице ниже.

Таблица 1

Метрики качества работы моделей YOLO на тестовой части датасета

		Preci- sion	Recall	mAP[0.5]	mAP [0.5:0.95]
YOLOv5	Целые ядра	0.953	0.915	0.954	0.673
	Разрушенные	0.865	0.717	0.725	0.4
	Все вместе	0.909	0.816	0.84	0.536
YOLOv8	Целые ядра	0.973	0.902	0.968	0.703
	Разрушенные	0.828	0.696	0.749	0.46
	Все вместе	0.9	0.799	0.859	0.582

Качество работы алгоритмов сегментации оценивалось как по результатам полученных метрик, так и по визуальному оцениванию на тестовых изображениях (рис. 4).

Методом визуального сравнения сложно выделить лучшую модель, так как обе модели дают сегментацию, подходящую для встраивания в алгоритм исследования X-хромосомного мозаицизма. Тем не менее, модель YOLOv8 работает в среднем на треть быстрее (44.5 мсек проход по сети + 89.7 мсек постобработка для YOLOv5 и 37.2 мсек проход по сети + 49.3 мсек постобработка для YOLOv8). В разработанной системе имеется возможность использовать обе модели для сегментации ядер клеток.

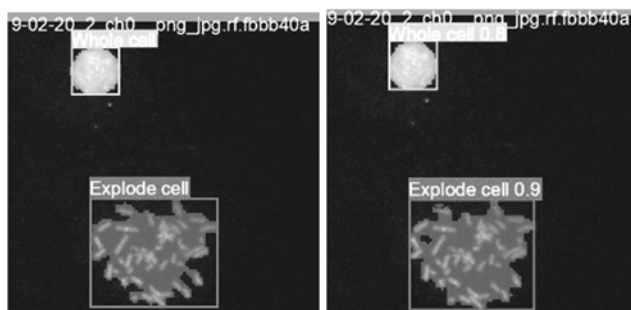


Рис. 4. Пример разметки и предсказания модели для целого и разрушенного ядер моделью YOLOv8

Детектирование хромосом

Хромосомы представляют собой на изображении ярко-красные или зеленые точки в зависимости от их типа. Они выглядят как одиночные или двойные пиксели с большой долей красного и зеленого в их окраске.

Разработанный алгоритм обнаружения хромосом состоит из следующих этапов:

1. Применение фильтра «нерезкая маска» (Unsharp Masking) [13] к красному и зеленому каналам изображения для устранения размытия;
2. Отбираются потенциальные кандидаты в хромосомы:
 - а. Производится бинаризация каждого канала по порогу T ($T = 80$);
 - б. Применяется операция морфологического градиента [14];
 - с. К полученным областям применяется маркировка связанных компонентов.
3. Применить фильтрацию потенциальных хромосом-кандидатов:
 - а. Для каждой хромосомы выявляют принадлежащее ей ядро клетки, если хромосома не лежит внутри него, то исключается из рассмотрения;
 - б. Проверяется близость хромосомы к границе клетки, если она лежит на ней (или очень близко по порогу $C = 1$), то хромосома исключается из рассмотрения.

Результатом работы алгоритма является список обнаруженных X-хромосом и 21 хромосом с указанием принадлежности к определенным клеткам на изображении.

Разработка графического инструментария

Предложенные алгоритмы детектирования клеток и обнаружения хромосом реализованы на языке Python. Для использования данных алгоритмов исследователями с помощью графической библиотеки PyQt было реализовано приложение с графическим интерфейсом, позволяющее открывать FISH-изображения в различных форматах, выполнять сегментацию ядер клеток и детектирование хромосом, визуализировать оригинальные изображения и результаты сегментации, выполнять подсчет количества целых и разрушенных клеток, подсчет количества хромосом внутри и вне ядер, сохранять результаты сегментации и подсчетов. Также приложение содержит функционал для автоматической обработки большого количества FISH-изображений, поддерживается обработка папки с FISH-изображениями либо из графического интерфейса, либо из командной строки. Исходный код приложения находится в открытом доступе: <https://github.com/CG-UNN-LAB/FISH-data-processing>.

Выводы

Предложен и реализован алгоритм и разработано программное обеспечение для исследования X-хромосомного мозаицизма на препаратах ядер клеток крови по данным конфокального микроскопа. Для этого был вручную размечен набор изображений и обучены модели глубокого обучения YOLOv5 и YOLOv8, выполнено сравнение их работы. Предложенный алгоритм показал свою высокую эффективность в задаче обработки данных с конфокального микроскопа, а разработанное программное обеспечение удобно в использовании для исследователя.

Программа дальнейших исследований будет включать расширение тренировочного набора данных, дальнейшее улучшение качества выделения оболочек клеток, развитие графического интерфейса.

Список литературы

1. Lachinov, D., Vasiliev, E., Turlapov, V. (2019). Glioma Segmentation with Cascaded UNet. In: Crimi, A., Bakas, S., Kuijf, H., Keyvan, F., Reyes, M., van Walsum, T. (eds) Brainlesion: Glioma, Multiple Sclerosis, Stroke and Traumatic Brain Injuries. BrainLes 2018. Lecture Notes in Computer Science, V. 11384. Springer, Cham. https://doi.org/10.1007/978-3-030-11726-9_17.
2. Gall J.G., Pardue M.L., Formation and detection of rna-dna hybrid molecules in cytological preparations, Proceedings of the National Academy of Sciences 63 (1969) P. 378–383. doi:10.1073/pnas.63.2.378.
3. Langer-Safer, P.R., Levine, M., Ward, D.C. Immunological method for mapping genes on Drosophila polytene chromosomes // Proc. Natl. Acad. Sci. USA, 1982. 79(14), P. 4381–4385.
4. Volpi, E.V. FISH glossary: an overview of the fluorescence in situ hybridization technique // Biotechniques / Bridger, J.M. 2008, 45(4), 385–409.
5. Knouse, K.A., Assessment of megabase-scale somatic copy number variation using single-cell sequencing // Genome Research / Wu, J., Amon, A. 2016, 26(3), P. 376–384.
6. CZI Image Format for Microscopes: <https://www.zeiss.com/microscopy/en/products/software/zeiss-zen/czi-image-file-format.html>
7. J. Redmon, S. Divvala, R. Girshick and A. Farhadi, "You Only Look Once: Unified, Real-Time Object Detection," 2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Las Vegas, NV, USA, 2016, P. 779–788, doi: 10.1109/CVPR.2016.91.
8. Ren, S. et al. 2015. Faster r-cnn: Towards real-time object detection with region proposal networks. Advances in neural information processing systems, 28.
9. ultralytics/yolov5: v7.0 - YOLOv5 SOTA Realtime Instance Segmentation doi.org/10.5281/zenodo.3908559

10. Ezzeddini L, Ktari J, Frikha T, Alsharabi N, Alayba A, J Alzahrani A, Jadi A, Alkholidi A, Hamam H. Analysis of the performance of Faster R-CNN and YOLOv8 in detecting fishing vessels and fishes in real time. PeerJ Comput Sci. 2024 May 31;10:e2033. doi: 10.7717/peerj-cs.2033.
11. Roboflow: computer vision for developers and enterprises: <https://roboflow.com/>
12. Mosaic augmentation: <https://wiki.cloudfactory.com/docs/mp-wiki/augmentations/mosaic>
13. SHARPENING: UNSHARP MASK: <https://www.cambridgeincolour.com/tutorials/unsharp-mask.htm>
14. Morphological transformations in OpenCV: https://docs.opencv.org/4.x/d9/d61/tutorial_py_morphological_ops.html

ОБУЧЕНИЕ С ПОДКРЕПЛЕНИЕМ В НЕЙРОБИОЛОГИИ И В СИСТЕМАХ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА

Е.И. ЗАЙЦЕВ, Е.В. НУРМАТОВА

МИРЭА – Российский технологический университет, Москва
zajcev@mirea.ru, nurmatova@mirea.ru

МУЛЬТИАГЕНТНАЯ СИСТЕМА УПРАВЛЕНИЯ РАСПРЕДЕЛЕННЫМИ ИНФОРМАЦИОННО- ВЫЧИСЛИТЕЛЬНЫМИ РЕСУРСАМИ

В статье рассматривается мультиагентная система управления распределенными информационно-вычислительными ресурсами, включающая системных и прикладных программных агентов, которые способны действовать рационально в неопределенной и динамичной внешней среде. Искусственный интеллект прикладных программных агентов, реализованный с использованием когнитивных структур данных, методов логического вывода и машинного обучения на базе нейронных сетей, позволяет создать высокоэффективную самоорганизующуюся распределенную систему управления. Приведена схема управления распределенными информационно-вычислительными ресурсами, которую совместно с операционной системой реализуют системные и прикладные программные агенты.

Ключевые слова: *распределенная система, мультиагентная система, база знаний, программные агенты, обучение с подкреплением.*

Введение

Мультиагентная система управления распределенными информационно-вычислительными ресурсами имеет многоуровневую гибридную архитектуру, в которой на самом верхнем уровне функционируют прикладные (интеллектуальные) программные агенты. Прикладные программные агенты реализуют свои решения посредством обращения к системным программным агентам, которые непосредственно взаимодействуют с операционной системой и событийно-управляемыми микросервисами. Агенты и микросервисы реализованы как автономные процессы, контролируемые операционной системой локального вычислительного узла. Агенты и микросервисы выполняют определённые функции и взаимодействуют друг с другом через интерфейсы прикладного программирования (API,

Application Programming Interface) с внутрипроцессными и с внепроцессными вызовами API. Для повышения производительности и надёжности многоагентной системы состояния микросервисов анализируются системными программными агентами [1, 2].

Прикладные программные агенты взаимодействуют с системой управления базой данных/знаний, осуществляют логические выводы, а также поиск оптимальной стратегии поведения. Для обучения прикладных программных агентов в мультиагентной системе управления используются методы обучения с подкреплением (Reinforcement Learning, RL) [4–8].

Первоначально в системе был реализован алгоритм Multi-Agent Proximal Policy Optimization (MAPPO), который является расширением достаточно эффективного и относительно простого в реализации алгоритма PPO (Proximal Policy Optimization). Однако, для сложной нестационарной распределенной среды, на которую одновременно воздействуют сразу несколько меняющих свои стратегии агентов, требуются более совершенные алгоритмы [9–23].

Такие алгоритмы, как COMA (Counterfactual Multi-Agent Policy Gradients), представляющий собой модификацию алгоритма A2C (Advantage Actor-Critic) для мультиагентного обучения, алгоритм SMMAE (Stochastic Monotonic Mixing Actor-Critic for Multi-Agent Environments), включающий стохастическую монотонную сеть смешивания, а также алгоритм MADDPG (Multi-Agent Deep Deterministic Policy Gradient) позволяют повысить эффективность обучения агентов. В новой версии мультиагентной системы управления распределенными информационно-вычислительными ресурсами реализуется подход, сочетающий в себе Model-free и Model-based алгоритмы, в которых учитываются не только модель и состояния (наблюдения) среды, но также состояния самих агентов.

Структура мультиагентной системы управления

Решение сложных проблем в мультиагентной системе осуществляется путем декомпозиции проблем на подзадачи, которые совместно решают программные агенты. Используется как горизонтальная декомпозиция, приводящая к созданию многосвязной системы с плоской структурой, необходимой для организации распределенных вычислений, так и вертикальная декомпозиция, создающая на вычислительных узлах иерархическую систему с несколькими уровнями.

Архитектура мультиагентной системы управления проиллюстрирована на рис. 1.

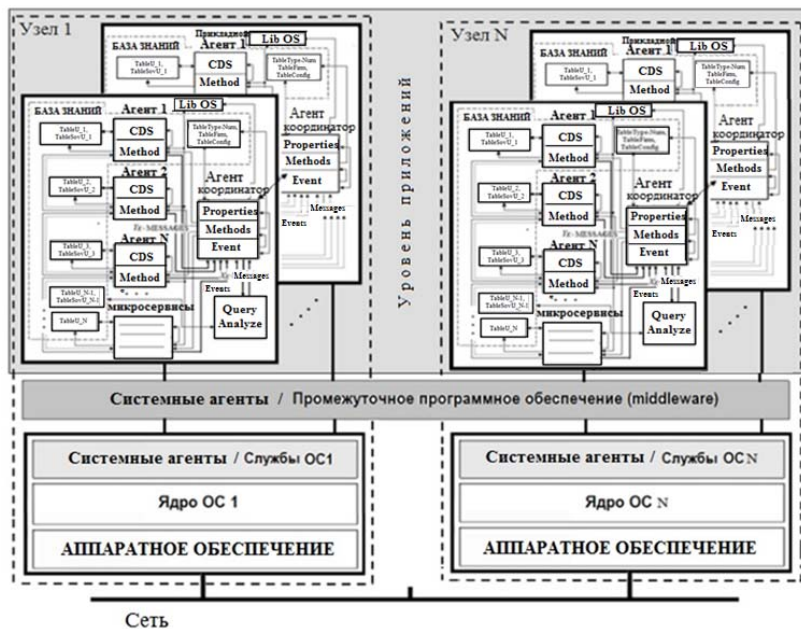


Рис. 1. Архитектура мультиагентной системы

Системные программные агенты реализуются в виде программных модулей адаптированной для приложения системной библиотеки LibOS. Системные агенты планируют работу процессов (потоков) и управляют распределенными информационно-вычислительными ресурсами, распределяя вычислительную нагрузку с учётом информации, которую передают агенты-мониторы.

В процессе решения подзадач программные агенты обращаются к микросервисам, которые с целью повышения производительности и надежности системы дублируются на разных вычислительных узлах (рис. 2).

При взаимодействии с базой знаний программные агенты реализуют четыре типа методов: анализ (ANS), ассоциация (ASS), сравнение (CMP) и спецификация (VAL). Метод ANS используется для логического анализа событий; метод ASS – для получения ответов на запросы о связях между объектами и событиями, метод CMP вызывается при сравнении событий или объектов; с помощью метода VAL осуществляется поиск значений определенных атрибутов. В VAL-методе для спецификации объектов предусмотрена возможность использования нечетких запросов к базе знаний [1–3].

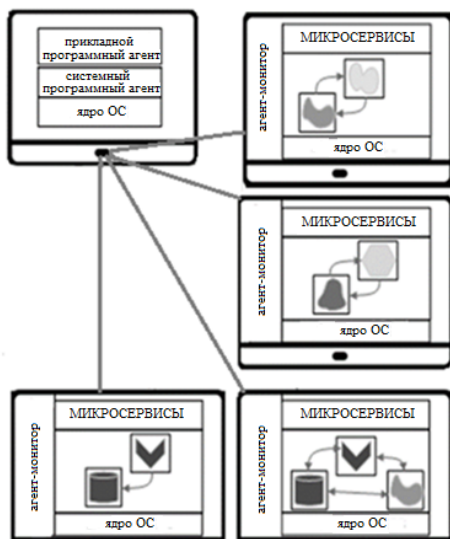


Рис. 2. Взаимодействие агентов с микросервисами

Если среда, в которой действует интеллектуальный агент, известна и полностью предсказуема, то агенту не требуется обучаться, он действует правильно согласно заложенным в него знаниям. Если среда изначально неизвестна, то программный агент должен её исследовать. При этом он не только собирает информацию об окружающей среде, но также обучается на тех данных, которые он воспринимает. Исходная конфигурация программного агента может отражать некоторые предварительные знания о среде, однако по мере приобретения агентом опыта эти знания могут модифицироваться и пополняться.

При многократном решении одной и той же подзадачи поведение интеллектуального агента должно становиться всё более рациональным. При использовании алгоритмов обучения с подкреплением оценка того, является ли поведение прикладного агента рациональным, даётся в соответствии с последствиями его действий. Показатели эффективности деятельности обычно выбираются в соответствии с тем, какие цели требуется достичь в данной среде. Эти показатели бывает достаточно сложно правильно сформулировать, кроме того, некоторые цели могут противоречить друг другу.

Усложнение модели среды и количества обучаемых агентов порождает дополнительные трудности в процессе обучения. В многоагентной среде

поведение прикладного агента описывается как направленное на максимизацию его собственных показателей эффективности, величина которых зависит от поведения других агентов. В многоагентной среде прикладные агенты в процессе обучения могут оказаться в состоянии равновесия Нэша, что приводит к тому, что агенты перестают улучшать свои стратегии. Если многоагентная система оказывается в ситуации, которая одновременно является равновесием Нэша и Парето-оптимальной, то тогда прикладные агенты вырабатывают оптимальные стратегии, находясь при этом в стабильном состоянии, которое не позволяет им увеличить свой выигрыш путем изменения стратегии.

Для обучения прикладных агентов системы управления распределенными ресурсами можно использовать существующие алгоритмы мультиагентного обучения с подкреплением, такие как QMIX (Q-value Mixing Network) с оценкой долгосрочной ценности каждого агента; VDN (Value-Decomposition Networks) с декомпозицией ценности действий; алгоритм SMMAE (Stochastic Monotonic Mixing Actor-Critic for Multi-Agent Environments) со стохастической монотонной сетью смешивания; алгоритм СОМА с общим критиком, передающим информацию от агента к агенту, либо ещё одну модификацию алгоритма Actor-Critic – MADDPG, позволяющую агентам при принятии собственных решений учитывать действия других агентов. Однако использование указанных алгоритмов в многоагентной системе управления распределенными ресурсами не делает процесс обучения эффективным.

В настоящее время для многоагентной системы управления распределенными ресурсами разрабатывается новый алгоритм, в котором используются модель среды и данные экспертов. Этот алгоритм позволяет быстрее обучать прикладных агентов в среде с большим пространством действий. В новом методе обучения в классическую схему Reinforcement Learning, использующую понятие "состояние среды" (S_Env), дополнительно вводится понятие "состояние агента" (S_Agent). Таким образом, в новом алгоритме принимаемые агентом решения о выполнении того или иного действия зависят не только от наблюдений за средой и сигналов подкрепления, но и от состояний самих агентов.

Подсистема проектирования и оптимизации логической структуры распределенной базы знаний мультиагентной системы

Для создания высокоэффективной самоорганизующейся распределенной системы управления ресурсами требуется разработать подсистему

управления распределенной базой знаний (РБЗ), которая должна выполнять синтез оптимальной логической структуры РБЗ или адаптировать к меняющимся условиям нестационарной внешней среды существующую структуру, а также своевременно и адекватно реагировать на ситуации, возникающие при дефиците вычислительных или временных ресурсов. В целях повышения управляемости и доступности РБЗ используется разделение хранимых объектов баз знаний на партии – разделы с разными параметрами физического хранения.

Основой такой подсистемы являются алгоритмы для решения задачи синтеза оптимальной логической структуры РБЗ по выбранному критерию эффективности. В качестве основного критерия эффективности, переменные которого приведены в табл. 1, рассматривался минимум общего времени последовательного выполнения множества транзакций, важность которого определяется тем, что в мультиагентной распределённой системе время является наиболее ценным ресурсом, а последовательная фиксация транзакций актуальна в системах со сложными процессами управления данными и знаниями, отдельные этапы которых выполняются различными пользователями. В дополнение, при параллельной работе транзакций могут возникать аномалии одновременного выполнения.

Математическая постановка задачи и имеет вид [2]:

$$\min_{\{x_{it} y_{tr} y_{jm}\}} \sum_{k=1}^{K_0} \sum_{p=1}^{P_0} \xi_{kp}^Q \varphi_{kp}^Q \left\{ \sum_{t=1}^T w_{st}^n \left[\sum_{m=1}^M (1 - Z_{sr}^m \delta_{pr}) (t^{acc} + t_{rm}^{transf}) + \sum_{r=1}^R Z_{sr} (t^{task} + t_{rm}^{ser} + t_{rm}^{transf}) \right] \right\} \quad (1)$$

Здесь t^{task} – средняя длительность формирования одного запроса в составе транзакции;

t_{rm}^{ser} – средняя длительность выбора маршрута, установления логических соединений между узлом-клиентом (r) и узлом-сервером (m);

t_{rm}^{transf} – средняя длительность передачи раздела данных (логической записи) с узла клиента (r) на узел сервера (m) по оптимальному пути;

t^{acc} – среднее время доступа к РБЗ и поиска записей;

ξ_{kp}^Q – частота использования запроса (p) транзакцией (k);

φ_{kp}^Q – бинарный признак использования запроса (p) транзакцией (k).

В качестве структурных и системных ограничений в задаче синтеза используются ограничения на число разделов данных и время обработки множества транзакций [2].

Таблица 1

Критерии эффективности задачи синтеза

Переменная	Значения
X_{it}	Определяет распределение логической записи i РБЗ в раздел t : $X_{it} = 1$, если запись попала в раздел; $X_{it} = 0$, в противном случае
Y_{tr}	Определяет размещение раздела t на сервере узла r распределённой системы
W_{st}^n	Определяет обновление транзакцией s , состоящей из n запросов, раздела t $W_{st}^n = 1$, если раздел обновлен
Z_{sr}	Показатель P_{sr} определяет множество узлов r серверов РБЗ, фиксируемых транзакцией s : $P_{sr} = \sum_{t=1}^T \sum_{i=1}^I w_{st}^n x_{it} y_{tr}$ $Z_{sr} = 1$, если $P_{sr} \geq 1$
Z_{sm}^m	Показатель P_{sm} определяет множество узлов m серверов РБЗ, содержащих данные по транзакции s раздела j : $P_{sm} = \sum_{j=1}^J w_{sj}^n y_{jm}$ $Z_{sr} = 1$, если $P_{sr} \geq 1$;

Задача синтеза оптимизации логической структуры РБЗ по указанному выше критерию принадлежит к классу задач дискретной оптимизации и является NP-трудной нелинейной задачей целочисленного программирования. Для решения задачи применён метод ветвей и границ, который является вариацией полного перебора с отсевом подмножеств допустимых решений, заведомо не содержащих оптимальных решений.

Применённый алгоритм синтеза (рис. 3), с помощью которого получен граф, позволил дать рекомендации по выбору и оценке параметров оптимальной структуры РБЗ.

В результате решения задачи распределения БЗ определяются следующие основные характеристики:

- состав разделов РБЗ, формализованный в виде матрицы X_{it} ;

- структура размещения разделов по серверам узлов МСПОЗ, формализованная в виде матрицы Y_{lr} ;
- граф распределения РБЗ на узлах-серверах Y_{jm} , содержащих данные по транзакциям конкретного раздела.

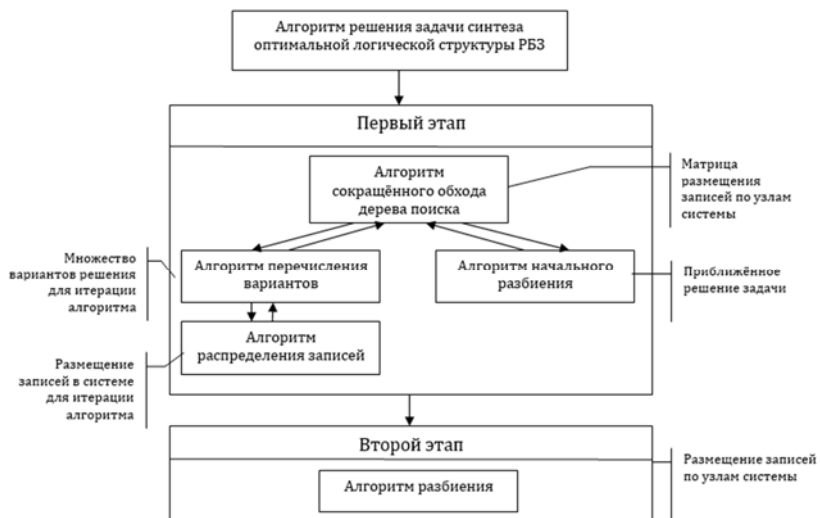


Рис.3. Сводная схема этапов решения задачи синтеза

Оптимально спроектированная модель БЗ позволяет избежать проблем со скоростью доступа к данным и знаниям системы, а также обеспечивает возможность дальнейшего масштабирования БЗ и подключения дополнительных источников данных.

Заключение

Рассмотрена архитектура мультиагентной системы управления, в которой прикладные и системные программные агенты совместно с операционной системой управляют распределенными информационно-вычислительными ресурсами в условиях неопределенности и нестационарности внешней среды. Для эффективного управления прикладные программные агенты используют адаптивную базу знаний и методы машинного обучения с подкреплением. Это позволяет создать высокоэффективную самоорганизующуюся распределенную систему управления, которая совершенствует свою стратегию в процессе работы.

Список литературы

1. Zaytsev E.I., Khalabiya R.F., Stepanova I.V., Bunina L.V.: Multi-Agent System of Knowledge Representation and Processing // Proceedings of the Fourth International Scientific Conference “Intelligent Information Technologies for Industry” (IITI’19), Springer, 2020, P. 131–141.
2. Zaytsev E.I., Nurmatova E.V. Approach to knowledge management and the development of a multi-agent knowledge representation and processing system. Russian Technological Journal, 2023, V. 11, N 4, P. 16–25.
3. Nurmatova E.V., Zaitsev E.I. Modeling of objects state estimation on the basis of queries with fuzzy conditions // Engineering Bulletin of Don. 2024. N 1(109). P. 115–12.
4. Sutton R.S., Barto A.G. Reinforcement Learning: An Introduction. Second Edition. MIT Press, Cambridge, MA, 2018. 552 p.
5. Barreto [et al.]. Successor features for transfer in reinforcement learning. In Advances in Neural Information Processing Systems, pages 4055–4065, 2017.
6. Jurgenson, T., Avner, O., Groshev, E., and Tamar, A. Sub-goal trees a framework for goal-based reinforcement learning. In International Conference on Machine Learning, 2020.
7. Zhao, R., Sun, X., and Tresp, V. Maximum entropy-regularized multi-goal reinforcement learning. In International Conference on Machine Learning, 2019.
8. Pong, V., Gu, S., Dalal, M., and Levine, S. Temporal difference models: Model-free deep RL for model-based control // International Conference on Learning Representations, 2018.
9. Jakob Foerster, Gregory Farquhar, Triantafyllos Afouras, Nantas Nardelli, Shimon Whiteson. Counterfactual Multi-Agent Policy Gradients // Proceedings of the Thirty-Second AAAI Conference on Artificial Intelligence, 2018.
10. Bozhidar Vasilev, Tarun Gupta, Bei Peng, Shimon Whiteson. Semi-On-Policy Training for Sample Efficient Multi-Agent Policy Gradients. arXiv preprint arXiv:2104.13446, 2021.
11. Ryan Lowe, Yi Wu, Aviv Tamar, Jean Harb, Pieter Abbeel, Igor Mordatch. Multi-agent actor-critic for mixed cooperative-competitive environments. arXiv preprint arXiv:1706.02275, 2017.
12. Wang H., Liu Z., Yi J., Pu Z. Multiagent hierarchical cognition difference policy for multiagent cooperation // Algorithms. 2021. V. 14, N 3.
13. Silva F.L. Da, Nishida C.E.H., Roijers D.M., Costa A.H.R. Coordination of Electric Vehicle Charging through Multiagent Reinforcement Learning // IEEE Trans. Smart Grid. 2020. V. 11, N 3.
14. Cui J., Liu Y., Nallanathan A. Multi-Agent Reinforcement Learning-Based Resource Allocation for UAV Networks // IEEE Trans. Wirel. Commun. 2020. V. 19, N 2.
15. Shamsoshoara A., Khaledi M., Afghah F., Razi A., Ashdown J. Distributed cooperative spectrum sharing in UAV networks using multi-agent reinforcement learning // arXiv. 2018.
16. Jie H. [et al.]. Joint Optimization of Multi-UAV Target Assignment and Path Planning Based on Multi-Agent Reinforcement Learning // IEEE Access. 2019. V.

17. Fang X. et al. Multi-agent reinforcement learning approach for residential microgrid energy scheduling // *Energies*. 2019. V. 13, N 1.
18. Hernandez-Leal P., Kartal B., Taylor M.E. A survey and critique of multiagent deep reinforcement learning // *Auton. Agent. Multi. Agent. Syst.* 2019. V. 33, N 6.
19. Zhang K., Yang Z., Başar T. Multi-agent reinforcement learning: A selective overview of theories and algorithms // *arXiv*. 2019.
20. Da Silva F.L., Reali Costa A.H. A survey on transfer learning for multiagent reinforcement learning systems // *J. Artif. Intell. Res.* 2019. V. 64.
21. Nguyen T.T., Nguyen N.D., Nahavandi S. Deep Reinforcement Learning for Multiagent Systems: A Review of Challenges, Solutions, and Applications // *IEEE Trans. Cybern.* 2020. V. 50. N 9.
22. Yang Y. [et al.]. Q-value path decomposition for deep multiagent reinforcement learning // *arXiv*. 2020.
23. Shamsoshoara A., Khaledi M., Afghah F., Razi A., Ashdown J. Distributed Cooperative Spectrum Sharing in UAV Networks Using Multi-Agent Reinforcement Learning // 2019 16th IEEE Annual Consumer Communications and Networking Conference, CCNC 2019.

ВЫЧИСЛИТЕЛЬНЫЕ ИССЛЕДОВАНИЯ ПРИНЦИПОВ И МЕХАНИЗМОВ РАБОТЫ ЕСТЕСТВЕННЫХ НЕЙРОННЫХ СИСТЕМ

К.Д. СЛАДКОВ, С.С. КОЛЕСНИКОВ

Институт биофизики клетки РАН, Пушкино, Московская обл.
klimitrich@ya.ru

ВКУСОВАЯ КЛЕТКА ТИПА III КАК НЕЙРОНАЛЬНАЯ*

Вкусовые клетки типа III функционируют как клеточные сенсоры кислого и с точки зрения электрогенеза являются нейроноподобными: они способны генерировать серию потенциалов действия, которые инициируют выброс афферентного нейротрансммиттера по классическому экзоцитозному механизму. В данной работе рассмотрены имеющиеся электрофизиологические модели вкусовых клеток типа III. В рамках моделей изучены ответы клеток на инъецируемый ток и обнаружено, что вкусовая клетка типа III относится к первому типу возбудимости по Ходжкину.

Ключевые слова: *вкусовая клетка, кислый вкус.*

Введение

Функциональной единицей вкусовой системы млекопитающих является вкусовая почка – объединение 50–100 вкусовых клеток. Апикальные мембраны клеток, выполняющие сенсорную функцию, обращены во вкусовую пору и отделены от базальной мембраны плотными межклеточными контактами. По своему составу среда во вкусовой поре сильно отличается от плазмы крови. Так, слюна содержит около 40 мМ Na^+ и 20 мМ K^+ , тогда как плазма крови содержит 140 мМ Na^+ и 4 мМ K^+ [1]. Поэтому на плазмолемме вкусовых клеток существуют как трансмембранные, так и латеральные градиенты ключевых физиологических ионов. Хотя с точки зрения эмбриологических критериев вкусовые клетки имеют эпителиальное происхождение [2], вкусовые клетки типа II и типа III имеют ряд свойств, свойственных нейрональным клеткам, такие как электрическая возбудимость и синаптическая передача. С точки зрения электрогенеза вкусовые клетки типа III (ВК III), которые специализируются на детекции кислых стимулов,

* Данная работа выполнена при частичной поддержке гранта Российского научного фонда № 22-14-00032.

практически мало отличаются от нейронов. В этих вкусовых клетках функционируют потенциал-зависимые Na^+ - и Ca^{2+} -каналы, потенциал-зависимые K^+ -каналы, включая K^+ -каналы задержанного выпрямления и транзистентные K^+ -каналы, HCN каналы, активируемые гиперполяризацией, K^+ -каналы входного выпрямления [3].

Как элемент системы передачи вкусовой информации ВК III получает на вход информацию о закислении раствора во вкусовой поре, преобразует ее в величину инъецируемого тока, генерирует ПД, приводящие к росту цитозольного кальция, вызывающего экзоцитоз, и выдает на выход афферентному нерву информацию о количестве выброшенного АТФ и серотонина. Нейроны также участвуют в передаче информации по схожему механизму.

Компьютерное моделирование вкусовой трансдукции от молекулярного сенсора вкусового вещества до выброса нейротрансмиттера представляет интерес, поскольку, в частности, позволяет симулировать физиологические условия функционирования вкусовых клеток, которые невозможно воспроизвести в эксперименте *in vitro*. Попытки использовать арсенал методов компьютерной электрофизиологии для описания электрических ответов вкусовых клеток предпринимались ранее [4, 5]. Использование формализма Ходжкина – Хаксли показало неплохую предсказательную силу для описания формы потенциала действия (ПД) и вольт-амперной характеристики клеток, что во многом связано с использованием экспериментальных значений для выбора величин параметров модели. Однако для количественного описания трансдукции вкуса необходимо описать частоту ПД как функцию инъецируемого тока, чего ранее не было сделано.

Согласно Ходжкину, возбудимые клетки делятся на три типа, определяемые типом зависимости частоты ПД, индуцируемых инъекцией тока, от его величины. К первому типу относятся клетки, монотонно меняющие частоту в широком диапазоне, также в них возможна генерация последовательности ПД с околонулевой частотой. Зависимость частоты ПД от величины тока в клетках второго типа носит пороговый характер, испытывая скачок при некоторой пороговой величине тока. Клетки третьего типа возбудимости в ответ на скачок тока генерируют один ПД и не могут генерировать последовательность ПД. В рамках вычислительной электрофизиологии было выяснено, что различия в поведении клетки связаны с разным типом бифуркации, происходящей в поведении динамической модели клетки, при изменении инъецируемого тока.

Следует отметить, что вкусовые клетки не были классифицированы в рамках возбудимости по Ходжкину. Видимо, это было связано с тем, что в литературе не были представлены репрезентативные зависимости частоты

ПД во вкусовых клетках от величины инъецируемого тока, без чего невозможно сделать вывод о принадлежности к одному из типов возбудимости. Однако, анализируя ранние работы нашей лаборатории, возможно построить зависимость частоты ПД от силы инъецируемого тока в одной клетке.

В более ранних работах нашей группы мы утверждали, что ВК III испытывает бифуркацию Андронова – Хопфа, с опорой на литературные результаты о частоте потенциалов действия в клетке в ответ на внеклеточное закисление. Предполагалось, что кислые стимулы вызывают в ВК III входящие протонные токи той же величины, что и в электрофизиологическом эксперименте, в котором клетки исследовались методом patch-clamp в режиме фиксации тока. Однако позже нами было выявлено, что частоты ПД, генерируемые клеткой в ответ на инъекцию тока и на кислые стимулы, вызывающие такие же по величине протонные токи, отличаются в десятки раз. Вероятно, помимо индуцирования протонных входящих токов, кислый внеклеточный раствор приводит к блокированию натриевых и калиевых каналов на базальной мембране клетки. Во вкусовой почке базальная мембрана отделена плотным контактом от апикальной мембраны, обращенной во вкусовую пору, поэтому не подвергается воздействию кислых стимулов.

Постановка задачи и описание модели

Будем предполагать, что ВК III представляет собой однокомпарментный нейрон и будем описывать токи через его мембрану в формализме Ходжкина – Хаксли:

$$C\dot{V} = \sum I_i + I_{st} + I_{hold} \quad (1)$$

$$I_i = gmh(V - E_r) \quad (2)$$

$$\dot{m} = \frac{(m_\infty(V) - m)}{t_a(V)}, \quad \dot{h} = \frac{(h_\infty(V) - h)}{t_i(V)} \quad (3-4)$$

$$m_\infty(V) = \frac{1}{1 + \exp\left(\frac{V_a - V}{s_a}\right)}, \quad h_\infty(V) = \frac{1}{1 + \exp\left(\frac{V_i - V}{s_i}\right)} \quad (5-6).$$

Здесь C – значение емкости ВК III;

$\dot{V}$ – временная производная трансмембранного потенциала ВК III;

$\sum I_i$ – сумма токов через ионные каналы;

I_{st} – величина деполяризующего тока;

I_{hold} – постоянная величина тока, поддерживающая клетку перед стимуляцией при потенциале покоя;

m, h – активационная и инактивационная воротные переменные соответствующего тока, индексом бесконечности обозначены стационарные значения переменных при текущем потенциале;

E_r – потенциал реверсии соответствующего ионного канала;

$t_a(V), t_i(V)$ – времена изменения соответствующих воротных переменных при текущем потенциале.

Параметры каналов выбраны в соответствии с компьютерной моделью Kimura [5]. Помимо этого, на основе литературных данных [6,7] нами были получены параметры входящего и выходящего тока, приведённые в таблице 1.

Таблица 1

Параметры ионных токов, использованные в модели

Kimura [5]	g , нСм	E_{rev} , мВ	V_a , мВ	s_a , мВ	t_a , мс	V_i , мВ	s_i , мВ	t_i , мс*
Na-ток	35	67	-28,0	5,0	0,34	-60	-10	2,2
K-TS-ток	10,4	-87	2,3	11,0	1,1	-	-	-
K-TN-ток	4,1	-87	16,5	11,4	0,5	-	-	-
Ток утечки	1,4	-21						
Ohtubo [6,7]	g , нСм	E_{rev} , мВ	V_a , мВ	s_a , мВ	t_a , мс	V_i , мВ	s_i , мВ	t_i , мс
Na-ток быстрый	4,14	90	-20,6	5,8	0,5	-61,8	9,2	4,9
Na-ток медлен- ный	4,86	90	-20,6	5,8	0,5	-61,8	9,2	980
K-ток	30	-88	17,9	8,1	0,9	-11,0	6,8	32,3

Результаты

В работе изучалась динамика электрического потенциала модельной клетки в ответ на токовые стимулы различной формы и интенсивности. В ответ на инъекцию стимулирующего тока (> 8 пА) модель генерировала пачки ПД практически неизменной частоты, которая зависела от величины

деполяризующего тока. Для сравнения данных моделирования и экспериментальных значений частоты ПД, в модели и в эксперименте определялись времена достижения второго и первого максимума потенциала после подачи токового стимула и вычислялась обратная величина их разности. Следует отметить, что оцененная таким образом частота ПД незначительно отличалась от средней частоты за время подачи стимулирующего тока. Частоты экспериментальных и модельных ПД, представлены на рис. 1. Согласно модели, частота непрерывно зависит от силы тока и может принимать минимальное значение $f = 5,0$ Гц при $I = 8,1$ пА, что говорит в пользу того, что ВК III относится к первому типу возбудимости. Имеющиеся в литературе экспериментальные данные недостаточны, чтобы подтвердить данный вывод, но значительное изменение частоты при изменении возбуждающего тока свидетельствует в пользу этого утверждения. Различия в значениях частоты объяснимы тем, что неизвестны электрическая емкость экспериментальной клетки и величина поддерживающего тока. Изменение этих параметров привело бы к сжатию-растяжению и смещению графика вдоль горизонтальной оси.

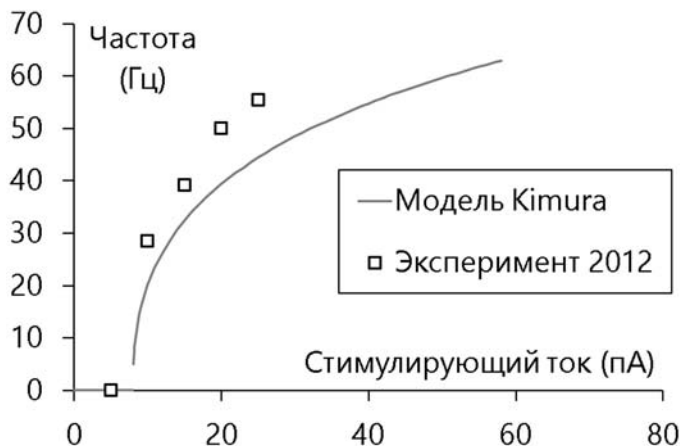


Рис. 1. Зависимость частоты потенциалов действия, генерируемых ВК III, в зависимости от силы токового стимула, полученная в эксперименте [3] и в компьютерной модели [4]

Для оценки частоты выше использовались первые два ПД, так как в эксперименте и модели ВК III неспособна к непрерывной генерации ПД. В ответ на пороговое изменение тока ВК III выдает до 10 ПД со всё уменьшающейся амплитудой [3], — явление, известное как деполяризационный блок. Ограниченное сверху количество ПД мотивировало дополнительно проверить, не принадлежит ли ВК III в рамках компьютерной модели к третьему типу возбудимости по Ходжкину. С этой целью мы изучали ответ модельной клетки на линейно возрастающий ток. Модель генерировала ПД с увеличивающейся частотой, что также свидетельствовало о первом типе возбудимости ВК III (рис. 2), поскольку нейроны третьего типа возбудимости не генерируют ПД в ответ на такой стимул, а нейроны второго типа генерируют ПД с постоянной частотой

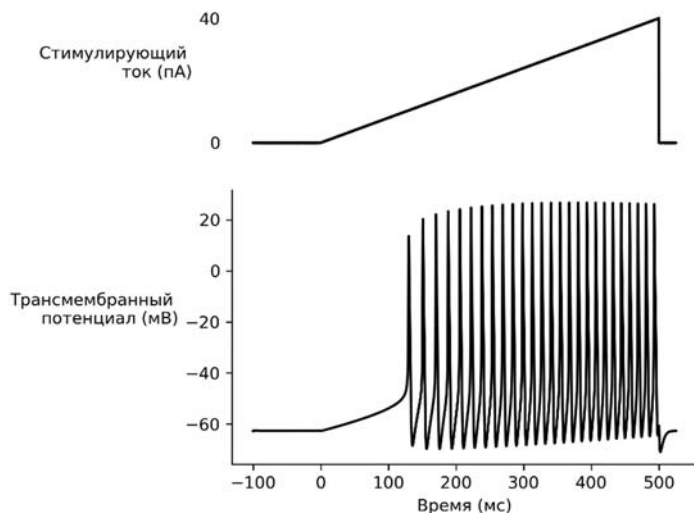


Рис. 2. Генерируемый моделью ВК III ответ на пилообразный токовый стимул.
Частота ПД увеличивается по мере увеличения тока

Уменьшение амплитуды ПД со временем при деполяризационном блоке происходит на масштабах времени, значительно превышающих продолжительность одного ПД. В предыдущих работах [6] это связывалось с наличием в ВК III двух изоформ натриевых каналов, одна из которых обладает крайне медленной инактивацией. Однако, деполяризационный блок наблюдался и в модели без двух различных натриевых каналов. С целью выяснить, какие каналы приводят к убыванию амплитуды ПД, нами был

выполнен расчет динамической проводимости каналов в соответствии с методикой, предложенной в [8]. Методика позволяет пересчитать стационарные проводимости каналов и временные характеристики их потенциал-зависимой модуляции в величины динамической проводимости каналов, описывающие участие каналов в формировании петель обратной связи, формирующих ПД на разных масштабах времени (рис. 3). Положительные значения динамической проводимости соответствуют положительной обратной связи, отрицательные значения – отрицательной

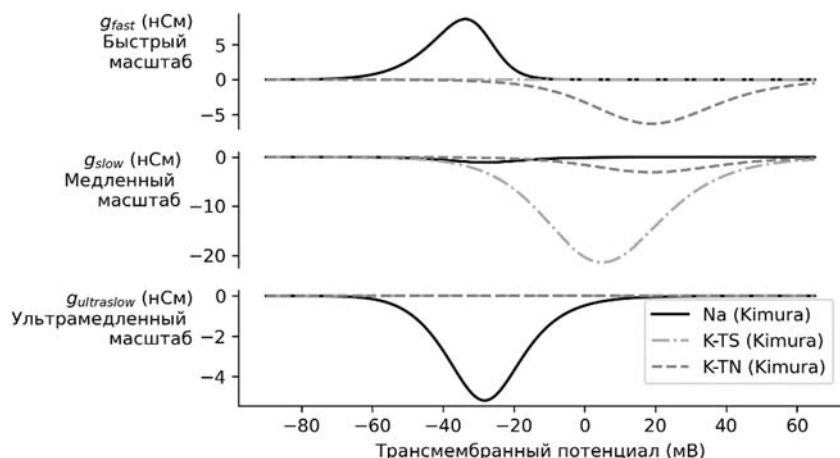


Рис. 3. Динамические проводимости отдельных каналов модели на разных масштабах времени, определенные для компьютерной модели Кимура в соответствии с методикой [8]. Три графика соответствуют трем масштабам времени, разными линиями отмечены динамические проводимости разных каналов как функция потенциала

Анализ полученного графика говорит, что в модели Кимура возбудимость ВК III обеспечивается натриевым каналом (положительная динамическая проводимость на быстром масштабе при потенциалах около потенциала покоя клетки). Отрицательную обратную связь, формирующую задний фронт ПД, обеспечивают на медленном и быстром масштабах времени соответственно ТЕА-чувствительный (K-TS) и ТЕА-нечувствительный каналы (K-TN). На самом медленном масштабе времени, при котором происходит убывание амплитуды ПД, наибольшей динамической проводимостью обладает натриевый канал. Это подтверждает роль инактивации натриевых токов и спорит с необходимостью двух популяций каналов с

разной динамикой инактивации для формирования деполяризационного блока.

Выводы

В работе изучалась предложенная ранее компьютерная электрическая модель вкусовой клетки типа III и изучены свойства её ответов на токовые стимулы. Полученная модель воспроизводит генерацию ПД, токо-зависимость частоты ПД и деполяризационный блок и предсказывает, что ВК III принадлежит к первому типу нейрональной возбудимости по Ходжкину. Результаты позволяют количественно описать блок вкусовой трансдукции кислого, начинающийся на апикальной мембране ВК III, где закисление раствора приводит к инъекции протонного тока, и заканчивающийся на базальной мембране, где расположены ионные каналы, приводящие к генерации ПД.

Список литературы

1. Spielman A.I. Interaction of saliva and taste // Journal of dental research. 1990. V. 69, N 3. P. 838–843.
2. Bradley R.M. and Stern I.B. The development of the human taste bud during the foetal period // Journal of anatomy. 1967. V. 101, N 4.
3. Romanov R.A., Rogachevskaja O.A., Bystrova M.F. and Kolesnikov S.S. Electrical excitability of taste cells. Mechanisms and possible physiological significance // Biochemistry (Moscow) Supplement Series A: Membrane and Cell Biology. 2012. V. 6. P. 169–185.
4. Chen P., Liu X.D., Zhang W., Zhou J., Wang P., Yang W. and Luo J. Modeling and simulation of ion channels and action potentials in taste receptor cells // Science in China Series C: Life Sciences. 2009. V. 52, N 11. P. 1036–1047.
5. Kimura K., Ohtubo Y., Tateno K., Takeuchi K., Kumazawa T., Yoshii K. Cell-type-dependent action potentials and voltage-gated currents in mouse fungiform taste buds // European Journal of Neuroscience. 2014. V. 39. N1. P. 24–34.
6. Ohtubo Y. Slow recovery from the inactivation of voltage-gated sodium channel Nav1.3 in mouse taste receptor cells // Pflügers Archiv-European Journal of Physiology. 2021. V. 473, N 6. P. 953–968.
7. Moribayashi T., Nakao Y. and Ohtubo Y. Characteristics of A-type voltage-gated K⁺ currents expressed on sour-sensing type III taste receptor cells in mice // Cell and Tissue Research. 2024. V. 396, N 3. P. 353–369.
8. Drion G., Franci A., Dethier J. and Sepulchre R. Dynamic input conductances shape neuronal spiking // Eneuro. 2015. V. 2, N 1.

**С.В. БОЖОКИН, А.А. РЯБОКОНЬ, Т.Д. ШОХИН,
А.М. ФИРСОВ**

Санкт-Петербургский политехнический университет Петра Великого
bsvjob@mail.ru

КОРРЕЛЯЦИЯ СПЕКТРАЛЬНЫХ ИНТЕГРАЛОВ НЕСТАЦИОНАРНОЙ ЭЭГ

Нестационарная электроэнцефалограмма (ЭЭГ) представлена как совокупность всплесков электрической активности нейронных ансамблей в различных спектральных диапазонах. Корреляция нестационарных сигналов ЭЭГ различных отведений головного мозга вычислена с помощью спектральных интегралов. Показано, что критерием корреляции является величина среднеквадратичной ошибки между нормированными спектральными интегралами. Проведено сравнение полученных величин корреляции с традиционным коэффициентом корреляции Пирсона.

Ключевые слова: *корреляция всплесков нестационарной ЭЭГ, непрерывное вейвлетное преобразование, спектральные интегралы.*

Введение

Сигнал электроэнцефалограммы (ЭЭГ), отражающий электрическую активность нейронных ансамблей головного мозга, является принципиально нестационарным сигналом, так как его спектральные и статистические свойства изменяются во времени [1–4]. Нестационарность сигнала ЭЭГ проявляется как в состоянии покоя испытуемого, так и во время проведения функциональных проб (фотостимуляция, гипервентиляция). Для анализа изменений спектрального состава нестационарных сигналов часто применяется оконное преобразование Фурье. При этом весь интервал измерений T разбивается на ряд последовательных окон, продолжительность каждого из которых равна $W = T / N$, где N – суммарное количество окон. Согласованное протекание во времени двух различных процессов $Z_J(t)$ и $Z_K(t)$ различных отведений головного мозга J и K характеризуется величиной когерентности. Величина когерентности позволяет оценить слаженность работы различных отведений головного мозга при различных функциональных пробах (фотостимуляция, гипервентиляция, психоэмоциональные пробы и решение когнитивных задач). Предположим, что фурье-преобразования двух сигналов $\hat{Z}_J^{(n)}(f)$ и $\hat{Z}_K^{(n)}(f)$, зависящие от частоты f и представленные через свои модули и фазы

$\hat{Z}_{J,K}^{(n)}(f) = \left| \hat{Z}_{J,K}^{(n)}(f) \right| \exp\left(i\phi_{J,K}^{(n)}(f)\right)$, вычислены для каждого окна n , где $n=1, 2, \dots, N$. Определение функции когерентности $\gamma_{JK}(f)$ содержит усреднение Фурье компонент двух сигналов по многим реализациям [5–10]:

$$\gamma_{JK}^2(f) = \frac{\left| \langle \hat{Z}_J(f) \hat{Z}_K^*(f) \rangle \right|^2}{\langle \left| \hat{Z}_J(f) \right|^2 \rangle \langle \left| \hat{Z}_K(f) \right|^2 \rangle}, \quad (1)$$

$$\langle \hat{Z}_J(f) \hat{Z}_K^*(f) \rangle = \frac{1}{N} \sum_{n=1}^N \left| \hat{Z}_J^{(n)}(f) \right| \left| \hat{Z}_K^{(n)}(f) \right| \exp\left[i\left(\phi_J^{(n)}(f) - \phi_K^{(n)}(f)\right)\right]. \quad (2)$$

Для сильно коррелированных сигналов величина $\gamma_{JK}(f) \approx 1$. Для хаотических сигналов $\gamma_{JK}(f) \approx 0$. Фазовая когерентность двух сигналов характеризуется величиной

$$PLV_{JK}(f) = \sqrt{\langle \cos(\phi_J(f) - \phi_K(f)) \rangle^2 + \langle \sin(\phi_J(f) - \phi_K(f)) \rangle^2}, \quad (3)$$

называемой Phase Locking Value. Угловые скобки в (3) обозначают процедуру усреднения по различным реализациям сигналов с окном продолжительностью W . Если два сигнала не синхронизированы по фазе, то величина $PLV_{JK}(f) \approx 0$. Для синхронизированных сигналов $PLV_{JK}(f) \approx 1$. Заметим, что в [10] показано, что значение когерентности двух сигналов зависит от процедуры усреднения, от выбора величины окна, от функции окна, от величины сдвига шага окна. Следовательно, величину когерентности нельзя рассматривать в качестве строгой количественной меры коррелированности двух сигналов $Z_J(t)$ и $Z_K(t)$, где J и K – два номера отведения ЭЭГ. Кроме того, при компьютерном моделировании используются различные виды окон и количество усреднений, что приводит к сложности сопоставления результатов, полученных в работах многих авторов.

Целью статьи является разработка метода вычисления корреляций нестационарных сигналов различных J и K каналов ЭЭГ. В отличие от ранее используемых величин $\gamma_{JK}(f)$ и $PLV_{JK}(f)$, зависящих от многих параметров оконного преобразования Фурье, в данной работе применяется

непрерывное вейвлетное преобразование сигналов ЭЭГ, которое не зависит от размера окна и параметров его сдвига. Корреляция J и K вычисляется с помощью сравнения спектральных интегралов $E_\mu(J, t)$ и $E_\mu(K, t)$ в различных спектральных диапазонах $\mu = (\delta, \theta, \alpha, \beta)$.

Непрерывное вейвлетное преобразование

В настоящее время большое развитие для количественного описания нестационарной ЭЭГ получил метод непрерывного вейвлет-преобразования (*CWT–Continuous Wavelet Transform*) [3, 4, 11–13]. Для исследуемого сигнала ЭЭГ $Z(t)$ формулы для непрерывного преобразования $V(v, t)$ и спектральных интегралов $E_\mu(t)$ в различных спектральных диапазонах $\mu = \{\delta, \theta, \alpha, \beta, \gamma\}$ приведены в работах [4, 13]. Непрерывное вейвлет-преобразование $V(v, t)$ отображает нестационарный сигнал $Z(t)$ с изменяющейся частотно-временной структурой на плоскость времени t и частоты v . Спектральный интеграл $E_\mu(t)$ представляет собой среднее значение локальной плотности спектра энергии сигнала, проинтегрированное по определенному интервалу частот $\mu = [v_\mu - \Delta v / 2; v_\mu + \Delta v / 2]$, где величина v_μ – середина диапазона частот $\mu = \{\delta, \theta, \alpha, \beta\}$, а Δv – ширина интервала.

Вычислим корреляцию двух нестационарных сигналов ЭЭГ $Z_J(t)$ и $Z_K(t)$ различных затылочных отведений $J = O_z$ и $K = O_2$ с помощью анализа поведения спектральных интегралов $E_\alpha(J = O_z, t)$ и $E_\alpha(J = O_2, t)$ в диапазоне α ритма $\alpha = [7.5 - 14 \text{ Hz}]$. На рис. 1 приведены графики сигналов ЭЭГ, квадратов модулей непрерывного вейвлетного преобразования и спектральных интегралов в диапазоне α ритма.

Запись ЭЭГ (рис. 1) для $E_\alpha(t)$ может быть рассмотрена как совокупность всплесков в диапазоне α ритма, причем каждая вспышка по времени имеет свое начало и свою продолжительность.

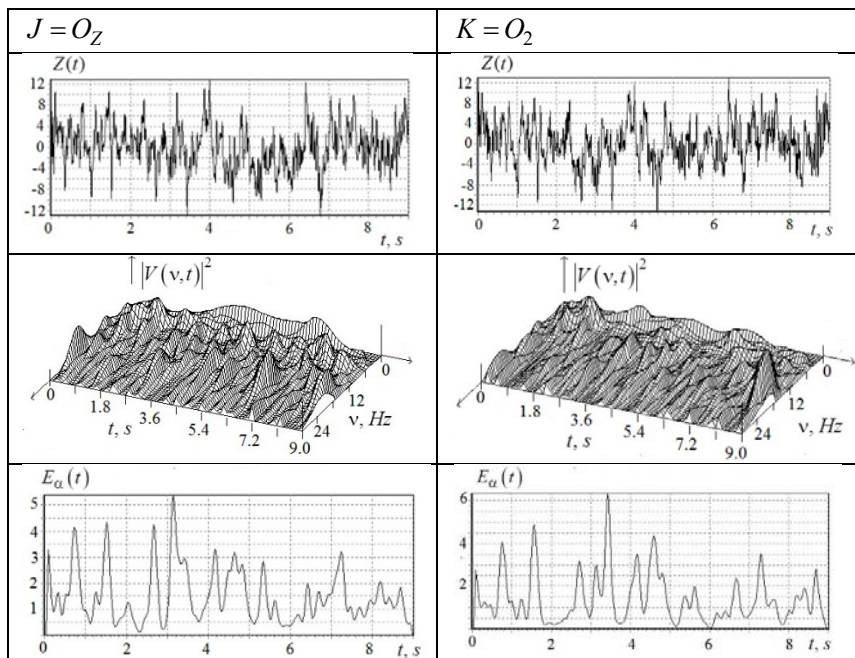


Рис. 1. Графики сигналов $Z(t)$ (верхний рисунок), квадраты модулей CWT $|V(v, t)|^2$ (средний рисунок) и спектральных интегралов $E_\alpha(t)$ (нижний рисунок). Левый столбец: отведение $J = O_Z$. Правый столбец: $K = O_2$

Расчет корреляции нестационарных сигналов

Обзор методов вычисления корреляции сигналов ЭЭГ содержится в работах [14], [15]. В данной работе такую корреляцию мы будем вычислять с помощью сравнения спектральных интегралов $E_\alpha(J = O_Z, t) = y(t)$, $E_\alpha(K = O_2, t) = z(t)$. Произведем нормировку всех входящих величин в интервале $[-1; 1]$. Если любая величина x изменяется в пределах $[x_{\min}; x_{\max}]$, то безразмерная величина $X = (2x - (x_{\max} + x_{\min})) / (x_{\max} - x_{\min})$ будет изменяться в пределах $X = [-1; 1]$. С помощью такого линейного преобразования вместо обычного времени $t \rightarrow u$ мы введем безразмерный параметр $u = [-1; 1]$. Для спектральных интегралов $y(t)$ и $z(t)$ выполним преобразо-

вания, причем $y(t) \rightarrow Y(u) = [-1; 1]$, $z(t) \rightarrow Z(u) = [-1; 1]$. Критерием сходства двух кривых $Y(u)$ и $Z(u)$, заданных на интервале $[-1; 1]$, является величина среднеквадратичной ошибки S_{YZ} , определяемой соотношением $S_{YZ}^2 = \langle [Y(u) - Z(u)]^2 \rangle$, где символом $\langle \dots \rangle$ обозначено выражение $\langle X(u) \rangle = \frac{1}{N} \sum_{i=0}^{N-1} X_i(u)$, в котором величина N представляет собой количество точек, в которых измеряется величина $\{X_0; X_1 \dots X_{N-1}\}$. Заметим, что если функции $Y(u)$ и $Z(u)$ нормированы на интервал $[-1; 1]$, то величина S_{YZ} лежит в пределах $[0; 2]$. Еще одним коэффициентом, количественно оценивающим корреляцию между величинами $Y(u)$ и $Z(u)$, является коэффициент корреляции Пирсона r_{YZ} [16]:

$$r_{YZ} = \frac{\langle Y(u)Z(u) \rangle - \langle Y(u) \rangle \langle Z(u) \rangle}{\sqrt{[\langle Y^2(u) \rangle - \langle Y(u) \rangle^2][\langle Z^2(u) \rangle - \langle Z(u) \rangle^2]}}. \quad (4)$$

Величина r_{YZ} , изменяющаяся в пределах $-1 \leq r_{YZ} \leq 1$, показывает, насколько связь между величинами $Y(u)$ и $Z(u)$ близка к строго линейной. На рис. 2 приведены графики нормированных спектральных интегралов $E_\alpha(J = O_z, t) \rightarrow Y(u)$ и $E_\alpha(K = O_2, t) \rightarrow Z(u)$ в зависимости от безразмерного времени $-1 \leq u \leq 1$.

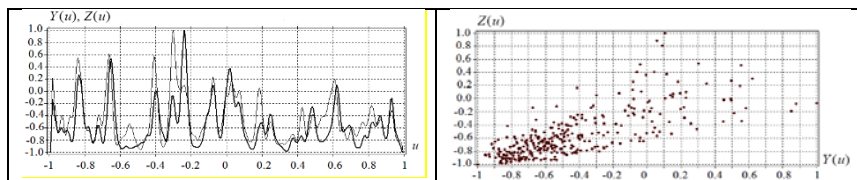


Рис. 2. На левом рисунке в зависимости от безразмерного времени u тонкой линией изображено поведение нормированного спектрального интеграла $E_\alpha(J = O_z, t) \rightarrow Y(u)$. Толстой линией – нормированный спектральный интеграл $E_\alpha(K = O_2, t) \rightarrow Z(u)$. На правом рисунке изображена диаграмма рассеяния точек $Z = F(Y)$

Зависимости спектральных интегралов $E_\alpha(t)$ в диапазоне α ритма для двух затылочных отведений $J = O_z$ и $K = O_2$ показывают сильную корреляцию последовательности таких всплесков. Сильное отличие по форме и продолжительности двух всплесков отмечается в момент времени $t \approx 3.2$ s (безразмерное время $u \approx -0.1$, рис. 1, рис. 2). Среднеквадратичная ошибка между спектральными интегралами для отведений $J = O_z$ и $K = O_2$ составляет величину $S_{YZ} \approx 0.29$. Диаграмма рассеяния точек (рис. 2, правый столбец) соответствует величине $r_{YZ} \approx 0.75$ (4), что говорит о сильной корреляции $E_\alpha(J = O_z, t) \rightarrow Y(u)$ и $E_\alpha(K = O_2, t) \rightarrow Z(u)$.

Неприменимость коэффициента корреляции Пирсона для анализа нестационарных сигналов

Многие исследователи в качестве коэффициента, описывающего корреляцию работы различных отведений головного мозга, используют коэффициент корреляции Пирсона. В качестве сигналов Y и Z , участвующих в вычислении коэффициента корреляции Пирсона, в этом случае выступают спектры мощности $Y \rightarrow P_J(f) = |\hat{Z}_J(f)|^2$, $Z \rightarrow P_K(f) = |\hat{Z}_K(f)|^2$, представляющие собой квадраты модулей коэффициентов Фурье соответствующих отведений J и K . Вычисленные спектры мощности свидетельствуют только о существовании определенных гармоник в сигналах различных отведений, но, к сожалению, не несут никакой информации о времени и продолжительности всплесков в различных спектральных диапазонах $\mu = \{\delta, \theta, \alpha, \beta, \gamma\}$. Отметим также противоречие использования коэффициента корреляции Пирсона. Оказывается, что для многих нормированных

кривых $Y(u)$ и $Z(u)$, между которыми существует большое значение среднеквадратичной ошибки S_{YZ} (величина S_{YZ} находится в интервале $[0; 2]$), коэффициент r_{YZ} неожиданно оказывается по модулю близким к единице $|r_{YZ}| \approx 1$, что говорит о их высокой корреляции (см. таблицу 1).

Левая колонка таблицы 1, состоящая из трех графиков, представляет собой зависимости $Y(u)$ (изображены кружками) и $Z(u)$ (изображен крестиками). Правая колонка – диаграмма рассеяния точек $Z_i = F(Y_i)$.

Таблица 1

Графики функций $Y(u)$, $Z(u)$ и диаграмма рассеяния точек

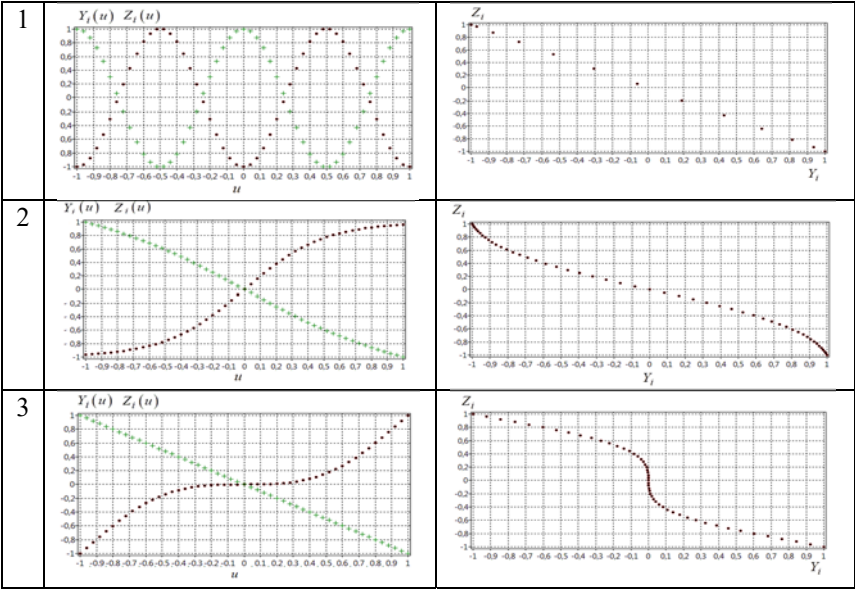


Таблица 2

Функции $Y(u)$ и $Z(u)$, величины S_{YZ} и r_{YZ}

	$Y(u)$	$Z(u)$	S_{YZ}	r_{YZ}
1	$2 \sin^2(\pi u) - 1$	$2 \cos^2(\pi u) - 1$	1.418	-1

2	$\frac{2}{1 + \exp(-4u)} - 1$	$-\frac{4}{\pi} \arctg(u)$	1.39	-0.9922
3	$\frac{4}{\pi} \arctg(u^3)$	$-u$	1.011	-0.9355

В табл. 2 приведен явный вид функций $Y(u)$ и $Z(u)$ (случаи 1–3 табл. 1), значения их среднеквадратичной ошибки S_{YZ} и коэффициента корреляции Пирсона r_{YZ} . Анализ таблицы 1 и 2 показывает, что для многих сильно различающихся функций $Y(u)$ и $Z(u)$, нормированных условиями $-1 \leq Y(u) \leq 1$, $-1 \leq Z(u) \leq 1$ и обладающих большой среднеквадратичной ошибкой $S_{YZ} > 1$, коэффициент корреляции Пирсона $|r_{YZ}| \approx 1$ неправильно указывает на их сильную корреляцию. Следовательно, коэффициент r_{YZ} , описывающий корреляцию двух нормированных функций, применим только в случае, когда величина S_{YZ} является малой.

Для того чтобы установить взаимосвязь величины S_{YZ} с коэффициентом r_{YZ} , были выполнены расчеты для двух функций $y(t)$ и $z(t)$:

$$y_i(t_i) = f_i(t_i) + \Delta(Random_i - 0.5), \quad (5)$$

$$z_i(t_i) = f_i(t_i) + \Delta(Random_i - 0.5). \quad (6)$$

К каждой из этих функций $y(t)$ (5) и $z(t)$ (6) добавлено шумовое воздействие, амплитуда которого равняется Δ , где $Random_i$, случайное число, изменяющееся в интервале $0 \leq Random_i \leq 1$. При увеличении амплитуды шума Δ расхождение между кривыми $y(t)$ и $z(t)$ возрастает. Для каждого значения Δ было выполнено $K=100$ реализаций, вызванных случайным шумом. Для каждой реализации была выполнена процедура нормировки $t \rightarrow u$, $y(t) \rightarrow Y(u)$, $z(t) \rightarrow Z(u)$. После усреднения по всем K реализациям была получена единственная точка, для которой вычисляется усредненный по всем реализациям коэффициент корреляции Пирсона $r_{YZ}(\Delta)$ в зависимости от усредненного по всем реализациям среднеквадратичной ошибки $S_{YZ}(\Delta)$. Затем величина Δ изменяется, и все этапы усреднения по реализациям повторяются. Расчеты, представленные на рис. 3, были выполнены для различных модельных функций $f(t)$: прямые линии (точки на

графике); совокупность гауссовских пиков (крестики); синусоидальный сигнал (треугольники).

Зависимость $r_{YZ} = F(S_{YZ})$ для малых S_{YZ} можно представить в виде простейшей формулы $\bar{r}_{YZ} = 1 - 2.195 \cdot S_{YZ}^2$. Среднеквадратичная ошибка между кривой $\bar{r}_{YZ}$ и всеми кривыми, представленными на рис. 3, составляет величину порядка 0.02. Таким образом, установлено взаимно однозначное соответствие между величиной S_{YZ} и величиной r_{YZ} . Если величина S_{YZ} возрастает в диапазоне $0 \leq S_{YZ} \leq 0.6$, то это означает, что коэффициент r_{YZ} монотонно падает. Диапазон изменений $r_{YZ} = [0.2; 1]$.

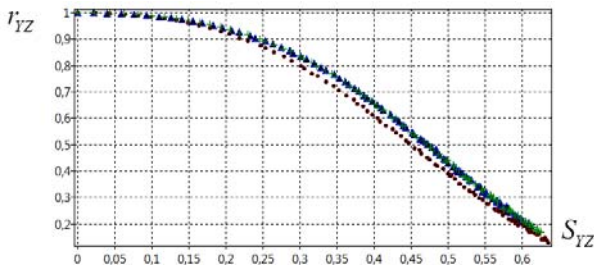


Рис. 3. Зависимость усредненного коэффициента корреляции Пирсона r_{YZ} от усредненной среднеквадратичной ошибки S_{YZ}

Выводы

Предложена модель, согласно которой нестационарный сигнал ЭЭГ каждого отведения J головного мозга $Z_J(t)$ может быть представлен как совокупность всплесков, каждая из которых имеет свое временное положение, продолжительность и спектральный состав. Для анализа свойств нестационарных сигналов использовано непрерывное вейвлетное преобразование. Для нахождения корреляции между сигналами J и K отведений выбраны нормированные спектральные интегралы, которые представляют собой сильно изменяющиеся функции времени для каждого спектрального диапазона $\mu = \{\delta, \theta, \alpha, \beta, \gamma\}$. Показано, что критерием корреляции между различными отведениями может быть величина среднеквадратичной ошибки S_{YZ} между нормированными спектральными интегралами. Указаны случаи, когда традиционный коэффициент корреляции Пирсона r_{YZ}

дает неправильный результат. Получено взаимно однозначное соответствие между r_{YZ} и S_{YZ} . Предлагаемый метод может быть применен для вычисления времен нарастания и спада корреляций сигналов различных отведений головного мозга во время различных функциональных проб, а также для анализа движения возмущений ЭЭГ по поверхности головного мозга.

Список литературы

1. Зенков Л.Р. Клиническая электроэнцефалография с элементами эпилептологии. Москва : МЕДПрессинформ, 2012. 356 с.
2. Кропотов Ю.Д. Количественная ЭЭГ, когнитивные вызванные потенциалы мозга человека и нейротерапия. Донецк, 2010. 512 с.
3. Короновский А.А., Макаров В.А., Павлов А.Н., Ситникова Е.Ю., Храмов А.Е. Вейвлеты в нейродинамике и нейрофизиологии. М. : Физматлит. 2013. 272 с.
4. Bozhokin S.V., Suslova I.B., Wavelet-based analysis of spectral rearrangements of EEG patterns and of non-stationary correlations // Physica A: Statistical Mechanics and its Applications. 2015. V. 421(1). P. 151–160.
5. Казанович Я.Б., Теория временной корреляции и модели сегментации зрительной информации в мозге (обзор) // Математическая биология и биоинформатика. 2010. Т. 5, N 1. С. 43–97.
6. Quiles M.G., Wang D., Zhaoc L., Romeroc R.A.F., Huang D.S. Selecting salient objects in real scenes: An oscillatory correlation model // Neural Networks. 2011. V. 24. P.54–64.
7. Каплан А.Я., Борисов С.В., Динамика сегментных характеристик альфа-активности ЭЭГ человека в покое и при когнитивных нагрузках // Физиология высшей нервной (психической) деятельности человека. 2003. Т. 53. № 1. С. 22–32.
8. Николаев А.Р., Иваницкий Г.А., Иваницкий А.Н., Исследование корковых взаимодействий в коротких интервалах времени при поиске вербальных ассоциаций// Журнал высшей нервной деятельности. 2000. Т. 50, N 1. С. 44-66.
9. Трофимов А.Г., Колодкин И.В., Ушаков В.Л., Величковский Б.М., Агломеративный метод выделения микросостояний ЭЭГ, связанных с характеристиками бегущих волн. XVII Всероссийская научно-техническая конференция «Нейроинформатика-2015»: Сборник научных трудов. В 3-х частях. Ч. 1. Москва : НИЯУ МИФИ, 2015. С. 66–77.
10. Kulaichev A.P., The informativeness of coherence analysis in EEG studies // Neuroscience and behavioral physiology. 2011. V. 41(3). P. 321–328.
11. Chavez M., Czelles B., Detecting dynamic spatial correlation patterns with generalized wavelet coherence and non-stationary surrogate data // Sci. Rep. 2019. V.9. P.7389.
12. Bozhokin S.V., Suvorov N.B, Wavelet analysis of transients of an electroencephalogram at photostimulation // Biomed. Radioelektron. 2008. N 3. P. 21–25.

13. Bozhokin S.V., Zharko S.V., Larionov N.V., Litvinov A.N., Sokolov I.M., Wavelet Correlation of Nonstationary Signals // Technical Physics. 2017. V. 62, N 6. P. 837–845.
14. Islam M.R., Islam M.M., Rahman M.M., Mondal C., Singha S.K., Ahmad M., Moni M.A., EEG channel correlation based model for emotion recognition // Computers in Biology and Medicine, 2021. V. 136. P. 104757.
15. Šverko Z., Vrankić M., Vlahinić S., Rogelj P., Complex Pearson correlation coefficient for EEG connectivity analysis // Sensors, 2022. V. 22, N 4. P.1477.
16. Кобзарь А.И. Прикладная математическая статистика. Для инженеров и научных работников. Москва : Физматлит, 2012. 812 с.

S. SUKHOV¹, B. BATUEV²

¹Ulyanovsk Branch of the Kotelnikov institute of radio engineering and electronics
of Russian academy of sciences

²Kotelnikov institute of radio engineering and electronics
of Russian academy of sciences, Moscow
ssukhov@ulireran.ru

PREDICTIVE CODING MODELS WITHIN THE NEURAL ENGINEERING FRAMEWORK*

The predictive coding (PC) concept postulates that the brain minimizes prediction errors at every hierarchical level of its neural organization. The theory of PC ought to be adapted for more biologically relevant neuronal models. Spiking neural networks (SNNs) can model brain functions with a high level of biological plausibility. However, the specific details of the implementation of PC in SNNs are still unknown. Here we present simple architectures of SNNs incorporating the elements of PC using the neural engineering framework. The dynamic neuronal models were implemented in the freely available Nengo package.

Keywords: *predictive coding, spiking neural networks, Nengo, neural engineering framework.*

Introduction

Predictive coding (PC) proposes a potentially unifying theory of cortical function [1, 2]. The basic idea behind PC is the assumption that the brain is composed of a hierarchy of layers; each layer makes predictions about the activity of the layer immediately below in the hierarchy. Thus, it is assumed that the core function of the brain is the minimization of prediction errors between the predicted and the actual input into one of the layers in the hierarchy.

* The reported study was funded by the Russian Science Foundation, RSF (project N 24-21-00470).

In contrast to modern machine learning algorithms which are trained with a global loss, prediction errors in PC are computed at every layer with local errors only. This property potentially enables predictive coding to learn in a biologically plausible way using only local (Hebbian) learning rules [3].

Since its inception, PC has got elaborated mathematical theory [4]. Different aspects of predictive coding were implemented in machine learning models [5] or in the models of biological neural networks [6, 7]. The single mechanism of predictive coding can explain various perceptual and neurobiological phenomena starting from information processing in the retina [8] and end-stopping in the visual cortex [1] to explanation of consciousness experience [9].

In our paper, we intend to implement simple PC models using spiking neural networks (SNNs). Spiking neural networks allow both modeling of biologically realistic systems and artificial systems suitable for neuromorphic devices. In our work, we use a Python-based Nengo neural simulator implementing the neural engineering framework [10]. Nengo allows to build and simulate large-scale models suitable both for robotic applications and for understanding and explanation of neurobiological experiments. Moreover, with Nengo, it is relatively easy to implement the dynamic behavior in spiking networks. In the present work, we present proof-of-the-principle models of predictive coding that can be used as a base for the construction of more complex dynamic neuronal systems.

Related work

Despite the large popularity of the predictive coding concept and the abundance of relevant literature [2], not much research was done related to implementing PC in spiking neural networks. The paper by N'dri et al. [11] proposes the Predictive Coding Light concept. The proposed network did not have explicit neurons for detecting errors between bottom-up inputs and top-down predictions. Instead, the network learned to inhibit the most predictable spikes from the lower neuronal level via lateral and top-down connections. The authors promoted the idea that the spikes that are easy to predict from the network's ongoing activity represent redundant information. Accordingly, there is no need to propagate these spikes forward. The paper [12] proposes a supervised learning algorithm inspired by predictive coding utilizing only local Hebbian plasticity in a network of integrate-and-fire spiking neurons. Mikulasch et al. [13] propose to implement predictive coding in SNNs using multicompartment neurons. In this approach, prediction errors are computed by the local voltage dynamics in the dendrites. In our paper, we build PC models using single-compartment integrate-and-fire neurons. In contrast to the majority of models of artificial neural networks, we consider recurrent SNNs that possess dynamic behavior. Using recurrent connections at every level of the constructed network, with the help of the PC, we propagate the desired dynamic behavior along the neural network.

Modeling principles of Nengo

Nengo is a neural simulator based on a theoretical concept called the Neural Engineering Framework (NEF) [10]. The NEF provides principles to guide the construction of neural models that incorporate anatomical constraints, functional objectives, and dynamical systems. Nengo has been used to build various neural mechanisms such as working memory [14], motor control [15], decision-making [16], visual attention, etc.

NEF is based on three principles that enable the simple construction of large-scale neural models. The first one is the *Representation principle*. According to this principle, a population of neurons with activities a_i collectively represents the value of time-varying vector $\mathbf{x}$ through a nonlinear encoding and linear decoding. The relation between representation vector $\mathbf{x}$ and underlying activities a_i is determined by the following expression

$$a_i = G(\alpha_i \mathbf{e}_i \mathbf{x} + b_i),$$

where G is the non-linearity specific to a neuron used for simulations, α_i is the gain parameter, $\mathbf{e}_i$ is the encoding vector, and b_i is the constant background bias current. Vice versa, the estimate of vector $\mathbf{x}$ can be found from activities a_i as follows:

$$\hat{\mathbf{x}} = \sum a_i d_i,$$

where d_i are linear decoders. Decoders in Nengo are found through least-squares minimization.

With the *Computation principle*, Nengo is able to quickly compute linear and nonlinear functions $f(\mathbf{x})$ on represented vectors.

Finally, the *Dynamics principle* provides the ability to simulate linear or non-linear dynamic systems through recurrent connections [17]. In particular, Nengo can compute dynamic functions of the form

$$\frac{d\mathbf{x}}{dt} = A\mathbf{x} + B\mathbf{u}, \quad (1)$$

where $\mathbf{u}$ is some (control) input, and A and B are arbitrary matrices. In the context of brain networks control theory, the matrix A is most often chosen to be the structural connectivity matrix [18]. To model eq. (1) with neuronal ensembles according to NEF, the recurrent connection weights should compute function $\tau A(\mathbf{x}) + \mathbf{x}$, and input connection weights should compute function $\tau B(\mathbf{u})$, where τ is the synaptic time constant [10].

With these three principles, Nengo allows simple modeling of dynamic neuronal systems.

Predictive coding with an integrator

One of the particular cases of eq. (1) is integrator:

$$\frac{d\mathbf{x}}{dt} = \mathbf{u}. \quad (2)$$

This equation represents a neural system that maintains its state in the case of no input ($u = 0$). Given a positive input, the represented value will increase, and given a negative input it will decrease. The integrator neural network can be a model of working memory or a decision-making process [19].

We use an integrator to build the simplest predictive model. The architecture of the network is shown in fig. 1. The neuronal population of the first level l_1 represents input signal x_1 . The population of the second level l_2 represents integrator, which dynamically maintains certain value x_2 through recurrent connections. Ensemble ϵ_1 calculates the disagreement (error) between the value of l_1 and the prediction of l_2 . Ensemble l_1 is connected to ϵ_1 by a lateral connection, and l_2 is connected to ϵ_1 by a feedback connection. Feedforward connection from ϵ_1 to l_2 transmits the error value along the neural hierarchy.

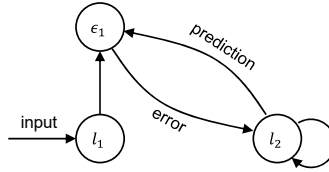


Fig. 1. An architecture of a simple predictive model based on integrator. Neuronal ensemble ϵ_1 calculates the error between representations of ensembles l_1 , l_2 . Each circle represents a separate neuronal ensemble

Eq. (2) for the value represented by the neurons at the second level becomes

$$\frac{dx_2}{dt} = e, \quad (3)$$

with $e = x_1 - x_2$ being the error. The stationary solution of eq. (3) is $x_2 = x_1$ with zero error transmitted upward in the hierarchy. Thus, layer 2 exponentially convergences its dynamics to represent the proper value.

The system represented in fig. 1 was modeled in the Nengo framework. All l_1 , ϵ_1 , and l_2 ensembles contained 100 leaky integrate-and-fire neurons each. For simplicity, all the neuron populations represent scalar values x . The synaptic constant τ was chosen to be equal to 0.2 seconds. Fig. 2 shows the example simulation of x_2 and e dynamics when input x_1 randomly changes its value over time. One can see that the value represented by the ensemble l_2 follows the changes of the value represented by l_1 governed by error e . When the error becomes equal to zero, the layer l_2 maintains its activity representing the proper value of x_1 . Lines jittering in Fig. 2 is determined by neuronal noise.

We should note that similar predictive models based on integrator (2) can be found in several other works. In the paper [20], the authors suggested a predictive estimator that can learn both feedforward and feedback connection weights individually using only locally available information. In contrast to our approach, the

error in [20] was calculated by a separate differential equation. Additionally, the learning of a transformation from l_1 to l_2 was introduced. The paper [21] relates the integrator with the predictive coding to alpha-band brain oscillatory dynamics. The integrator spiking neural network in the context of PC was also presented in [22]. However, the authors understand PC in the sense of predicting a detailed balance between excitation and inhibition in neural networks.

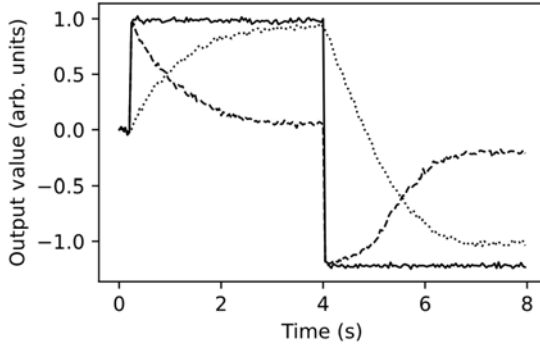


Fig. 2. The evolution of x_1 (solid line), error signal e (dashed line), and x_2 values (dotted line) in predictive integrator

Predictive control of dynamics in spiking neural networks

One of the simplest possible dynamics is the oscillatory one. Oscillations naturally emerge and play a significant role in biological neural networks [23]. For example, oscillations appear to play an important role both during awake perception [24] and during sleep [25]. Cortical oscillations are believed to be useful for sampling from a generative model necessary for predictive coding [26]. Neural oscillations play an important role in cognition and neural communication [27, 28].

An oscillator can be described by the following equation:

$$\ddot{x} + \omega^2 x = 0. \quad (2)$$

Here x is some variable demonstrating oscillatory behavior, ω is the frequency of oscillations. The neuronal ensembles in Nengo can encode only differential equations of the first order. Accordingly, the second order eq. (4) can be transformed into the first-order equation [29] for vector variable $\mathbf{x}$:

$$\begin{aligned} \frac{d\mathbf{x}}{dt} &= \hat{A}\mathbf{x}, \\ \mathbf{x} &= \begin{pmatrix} x \\ v \end{pmatrix}, \hat{A} = \begin{pmatrix} 0 & \omega \\ -\omega & 0 \end{pmatrix}, \end{aligned} \quad (3)$$

where $v = dx/dt$.

A basic functionality that is required in both biological modeling and robotics is the frequency modulation of oscillations [30]. The modulation of frequency can be supplied to an oscillator externally as a control parameter [29]. Combining the predictive integrator (3) and oscillator (5), one can obtain a simple predictive system with controlled dynamics (see fig. 3):

$$\hat{A} = \begin{pmatrix} 0 & x_i \cdot \omega \\ -x_i \cdot \omega & 0 \end{pmatrix},$$

where $x_i = x_{1,2}$ are the values calculated by the predictive integrator. The represented values $x_{1,2}$ of layers $l_{1,2}$ are supplied to the neural populations $o_{1,2}$ to modulate their oscillatory behavior.

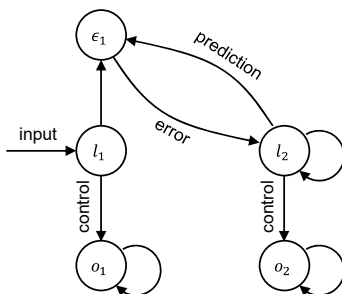


Fig. 3. Predictive control of the dynamics of oscillators o_1, o_2 located at different levels of hierarchy. Neuronal ensemble ϵ_1 calculates the error between representations of ensembles l_1, l_2 . Each circle represents a separate neuronal ensemble

The network depicted in fig. 3 was implemented in Nengo. The parameters of neuronal ensembles l_1, ϵ_1, l_2 and synaptic time constant τ were the same as in the previous section. Neuronal ensembles o_1, o_2 contained 500 leaky integrate-and-fire neurons each. Additional constant input with value 1 and lasting 0,1 seconds was supplied to both oscillators (not shown in fig. 3) to kick-start their activity. The values x_1, x_2 shown in fig. 2 can now be interpreted as oscillation frequencies of oscillators o_1, o_2 .

Learning in the hierarchy of oscillators

Besides fast feedforward information processing, the predictive coding framework assumes a learning stage when the behavior of the upper hierarchical layer is adjusted to better predict the input signal [31]. In this section, we implement the learning stage in a system of two oscillators.

Learning of recurrent networks is a nontrivial task both in machine learning and computational neuroscience. We simplified learning by applying a learning rule to $\epsilon_1 - o_2$ connection (fig. 4). For learning, we used the Prescribed Error

Sensitivity (PES) learning rule [32]. The PES rule applies the following update rule to the decoding weights $\mathbf{d}$ at each timestep:

$$\Delta \mathbf{d} = \kappa e \mathbf{a},$$

where $\kappa > 0$ is the learning rate, e is the error, and vector $\mathbf{a}$ represents the activity of the neurons in a population that we intend to teach. In our realization, the PES rule performed the learning of parameters x , v , and ω of the oscillator o_2 (fig. 4). One set of parameters (x, v, ω) is provided to the population ϵ_1 by the oscillator o_1 ; the second set of (x, v, ω) is provided to ϵ_1 through the feedback connection from the oscillator o_2 . As a result of learning, the oscillator o_2 modifies its dynamics to better correspond to the o_1 dynamics in amplitude, frequency, and phase. During learning, the represented error of ϵ_1 population decreases.

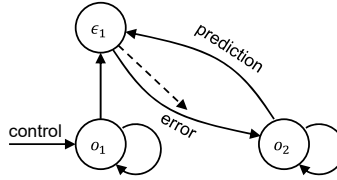


Fig. 4. The architecture of a predictive neural network utilizing learning (indicated by a dashed arrow) in a system of two oscillators o_1 , o_2 . Each circle represents a separate neuronal ensemble

The network depicted in fig. 4 was implemented in the Nengo simulator. The neuronal ensembles o_1 , ϵ_1 , o_2 contained 500 leaky integrate-and-fire neurons each. The synaptic time constant τ was equal to 0,1 seconds. The eigenfrequency of the oscillator o_1 was set to be 8. The initial frequency of the oscillator o_2 was set to 14. For PES, we used default learning parameters. Fig. 5 shows the example simulation of the learning of the receiver oscillator o_2 . One can see that over time o_2 oscillator learns to mimic the behavior of the master oscillator o_1 .

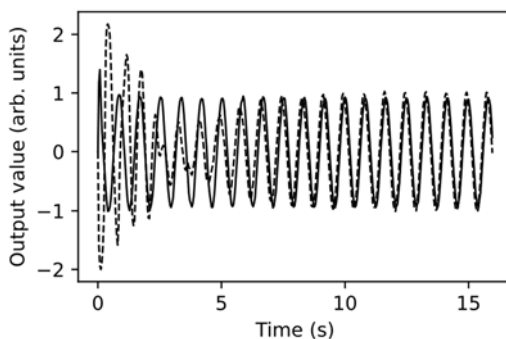


Fig. 5. Learning in a system of two oscillators. The receiver oscillator (dashed line) adjusts its amplitude, frequency, and phase to the corresponding values of the master oscillator (solid line)

Conclusion

In this paper, we presented simple models utilizing the predictive coding concept with the neural engineering framework. NEF allows simple implementation of the predictive coding by explicitly identifying algebraic equations to construct the required behavior and to quickly test hypotheses. However, NEF still does not answer the question, of how predictive coding is realized in biological neural networks where explicit equations cannot be specified.

All our models contain the hierarchy of neuronal levels and a separate error population that provides the error signal to the neurons of the upper hierarchy. In the first two models, we demonstrated a quick hierarchical response to a changing stimulus that corresponds to the perception without learning of weights. The last model includes the learning stage that slowly modifies the neuronal behavior of the neurons in the upper hierarchy.

Although we used leaky integrate-and-fire neurons in all the models, Nengo allows simple adaptation of the developed networks to other neuronal models. The two-layer networks examined in this study can be also easily extended to multilayer structures.

The supplementary Nengo code to reproduce the main results of this work can be found at <https://github.com/ssukhov/predictive-coding-nengo>.

Список литературы

1. Rao R.P.N., Ballard D.H. Predictive coding in the visual cortex: a functional interpretation of some extra-classical receptive-field effects // Nat. Neurosci. 1999. V. 2. P. 79–87.

2. Millidge B., Seth A., Buckley C.L. Predictive coding: a theoretical and experimental review // arXiv preprint arXiv:2107.12979. 2021.
3. Whittington J.C.R., Bogacz R. An approximation of the error backpropagation algorithm in a predictive coding network with local Hebbian synaptic plasticity // *Neural Comput.* 2017. V. 29. P. 1229–1262.
4. Buckley C.L., Kim C.S., McGregor S., Seth A.K. The free energy principle for action and perception: A mathematical review // *J. Math. Psychol.* 2017. V. 81. P. 55–79.
5. Marino J. Predictive coding, variational autoencoders, and biological connections // *Neural Comput.* 2022. V. 34. P. 1–44.
6. Bastos A.M., Usrey W.M., Adams R.A., Mangun G.R., Fries P., Friston K.J. Canonical microcircuits for predictive coding // *Neuron.* 2012. V. 76. P. 695–711.
7. Shipp S., Adams R.A., Friston K.J. Reflections on agranular architecture: predictive coding in the motor cortex // *Trends Neurosci.* 2013. V. 36. P. 706–716.
8. Srinivasan M. V., Laughlin S.B., Dubs A. Predictive coding: A fresh view of inhibition in the retina // *Proceedings of the Royal Society of London - Biological Sciences.* 1982. V.216. P. 427–459.
9. Seth A.K., Suzuki K., Critchley H.D. An interoceptive predictive coding model of conscious presence // *Front. Psychol.* 2012. V. 3. P. 395.
10. Bekolay T., Bergstra J., Hunsberger E., DeWolf T., Stewart T.C., Rasmussen D., Choo X., Voelker A.R., Eliasmith C. Nengo: A Python tool for building large-scale functional brain models // *Front. Neuroinform.* 2014. V.7. P. 48.
11. N'Dri A.W., Barbier T., Teuliere C., Triesch J. Predictive Coding Light: Learning compact visual codes by combining excitatory and inhibitory spike timing-dependent plasticity // *IEEE Computer Society Conference on Computer Vision and Pattern Recognition Workshops.* 2023. P. 3997–4006.
12. Lan M., Xiong X., Jiang Z., Lou Y. PC-SNN: Supervised learning with local Hebbian synaptic plasticity based on predictive coding in spiking neural networks // arXiv preprint arXiv:2211.15386. 2022.
13. Mikulasch F.A., Rudelt L., Wibrall M., Priesemann V. Where is the error? Hierarchical predictive coding through dendritic error computation // *Trends in Neurosciences.* 2023. V. 46. P. 45–59.
14. Singh R., Eliasmith C. Higher-dimensional neurons explain the tuning and dynamics of working memory cells // *Journal of Neuroscience.* 2006. V. 26. P. 3667–3678.
15. DeWolf T., Eliasmith C. The neural optimal control hierarchy for motor control // *J. Neural Eng.* 2011. V. 8. P. 065009.
16. Stewart T.C., Bekolay T., Eliasmith C. Learning to select actions with spiking neurons in the basal ganglia // *Front Neurosci.* 2012. V. 6. P. 2.
17. MacNeil D., Eliasmith C. Fine-tuning and the stability of recurrent neural networks // *PLoS One.* 2011. V. 6. P. e22885.
18. Gu S., Pasqualetti F., Cieslak M., Telesford Q.K., Yu A.B., Kahn A.E., Medaglia J.D., Vettel J.M., Miller M.B., Grafton S.T., Bassett D.S. Controllability of structural brain networks // *Nat. Commun.* 2015. V. 6. P. 8414.
19. Khona M., Fiete I.R. Attractor and integrator networks in the brain // *Nat. Rev. Neurosci.* 2022. V. 23. P. 744–766.

20. Xu D., Clappison A., Seth C., Orchard J. Symmetric Predictive Estimator for Biologically Plausible Neural Learning // *IEEE Trans. Neural Netw. Learn. Syst.* 2018. V. 29. P. 4140–4151.
21. Alamia A., VanRullen R. Alpha oscillations and traveling waves: Signatures of predictive coding? // *PLoS Biol.* 2019. V. 17. P. e3000487.
22. Boerlin M., Machens C.K., Denève S. Predictive Coding of Dynamical Variables in Balanced Spiking Networks // *PLoS Comput. Biol.* 2013. V. 9. P. e1003258.
23. Brunel N. Dynamics of sparsely connected networks of excitatory and inhibitory spiking neurons // *J. Comput. Neurosci.* 2000. V. 8. P. 183–208.
24. Başar E., Başar-Eroglu C., Karakaş S., Schürmann M. Gamma, alpha, delta, and theta oscillations govern cognitive processes // *International Journal of Psychophysiology.* 2001. V. 39. P. 241–248.
25. Adamantidis A.R., Gutierrez Herrera C., Gent T.C. Oscillating circuitries in the sleeping brain // *Nat. Rev. Neurosci.* 2019. V. 20. P. 746–762.
26. Korcsak-Gorzo A., Muller M.G., Baumbach A., Leng L., Breitwieser O.J., Van Albeda S.J., Senn W., Meier K., Legenstein R., Petrovici M.A. Cortical oscillations support sampling-based computations in spiking neural networks // *PLoS Comput. Biol.* 2022. V. 18. P. e1009753.
27. Wang X.J. Neurophysiological and computational principles of cortical rhythms in cognition // *Physiol. Rev.* 2010. V. 90. P. 1195–1268.
28. Womelsdorf T., Schoffelen J.M., Oostenveld R., Singer W., Desimone R., Engel A.K., Fries P. Modulation of neuronal interactions through neuronal synchronization // *Science* 2007. V. 316. P. 1609–1612.
29. Stewart T.C. A Technical Overview of the Neural Engineering Framework // *The Newsletter of the Society for the Study of Artificial Intelligence and Simulation of Behaviour.* 2012. P. 135.
30. Wyffels F., Li J., Waegeman T., Schrauwen B., Jaeger H. Frequency modulation of large oscillatory neural networks // *Biol. Cybern.* 2014. V. 108. P. 145–157.
31. Friston K. Hierarchical models in the brain // *PLoS Comput. Biol.* 2008. V. 4. P. e1000211.
32. Bekolay T., Kolbeck C., Eliasmith C. Simultaneous unsupervised and supervised learning of cognitive functions in biologically plausible spiking neural networks // *Cooperative Minds: Social Interaction and Group Dynamics - Proceedings of the 35th Annual Meeting of the Cognitive Science Society.* 2013. P. 169–174.

ПРИКЛАДНЫЕ НЕЙРОСЕТЕВЫЕ СИСТЕМЫ

И.В. КОЛОМИЙЧУК, А.И. КАНЕВ

Московский государственный технический университет им. Н.Э. Баумана
ilyakolomichyk@gmail.com, aikanev@bmstu.ru

НЕЧЕТКИЙ ПОИСК НЕ СЛОВАРНЫХ СЛОВ

Рассматривается проблема кодирования не словарных слов в отсутствие контекста. Был создан ансамбль моделей, обученных на n -граммах, для кодирования слова в вектор на основе его написания. Проанализирована корректирующая способность этого подхода при различном количестве ошибок в написании слова. Рассчитана точность нахождения правильной формы слова с разным количеством ошибок и разным количеством кандидатов на правильную форму. Результирующий ансамбль моделей сравнивается с классическим методом исправления ошибок в словах, таким как расстояние Левенштейна. Рассмотрен общий конвейер для построения базовой модели поиска слов с орфографическими ошибками в векторной базе знаний. Выделены сильные и слабые стороны подхода.

Ключевые слова: *векторный поиск, FastText, FAISS, расстояние Левенштейна, кодирование не словарных слов.*

Введение

Нечеткий поиск [1], также известный как нечеткое сопоставление, – это метод текстового поиска, который ищет совпадения с близким по значению термином, а не точное совпадение. Этот подход позволяет находить релевантные результаты, даже если поисковые запросы написаны неправильно или не полностью соответствуют ожидаемому содержанию. Количество слов, которых нет в словаре, растет с каждым днем: будь то научный термин, название лекарства или название телефона. Люди часто допускают ошибки при поиске таких слов, поэтому задача поиска термина, названия товара или услуги в базе знаний решается с учетом его неправильного написания. Важно также отметить, что люди с врожденными дефектами речи, а также люди с врожденной глухотой часто допускают не только обычные ошибки при наборе текста на клавиатуре, связанные с когнитивным переутомлением, но и ошибки, связанные с незнанием правильного произношения или написания того или иного слова. Таким образом, поиск информации в определенной предметной области с орфографическими

ошибками особенно важен для современного общества, поскольку этот инструмент позволяет иметь гибкий и надежный доступ к большому объему информации.

Коррекция орфографии – задача обработки естественного языка, направленная на исправление орфографических ошибок в тексте или документе. Эта задача играет важную роль в повышении удобства ввода информации и повышении устойчивости к помехам (опечаткам) в информации.

В исследовании [2] были изучены основные ошибки и частота их возникновения, которые допускаются при наборе текста на английской раскладке клавиатуры. Всего существует пять основных типов ошибок: пропуски, вставки, перестановки, замены и дублирование.

1.Пропуски: ошибки, при которых в тексте пропущены символы. Это наиболее распространенный тип ошибок. Эти опечатки чаще встречаются в более быстрых последовательностях.

2.Вставка: ошибка, при которой вставляются дополнительные символы, отсутствующие в тексте.

3.Замена: ошибка, при которой вместо нужного символа используется символ, который находится рядом на клавиатуре.

4.Дублирование: ошибка, связанная с тем, что определенная клавиша набирается дважды.

5.Перестановка: ошибка, при которой пара клавиш набирается в обратном порядке.

Сегодня существует множество подходов к кодированию слов и предложений, таких как Glove [3] и Word2vec [4], использующих статистическую эвристику. С появлением transformers [5] векторное кодирование стало более осмысленным: модели стали понимать контекст гораздо лучше, чем при статистическом подходе. Например, модель BERT хорошо подходит для решения задачи исправления ошибок [6]. В связи с этим открытием направление семантического поиска стало популярным [7]. Семантический поиск основан на построении и использовании базы знаний, которая используется для поиска информации на основе векторного представления слова или предложения. Простейший пример семантического поиска – поиск книг в библиотеке: клиент приходит с желанием купить книгу, похожую на ту, которую он недавно прочитал, и описывает детали сюжета. Сотрудник библиотеки понимает запрос пользователя и, основываясь на своих знаниях, ищет на библиотечных полках что-то похожее на книгу клиента. Также мы хотели бы упомянуть недавние исследования по семантическому поиску в базе знаний, основанной на метаграфах [8]. Но подход семантического поиска имеет одно слабое место: он плохо подходит для задачи кодирования слова не из словаря, то есть для кодирования

не словарных слов, в отсутствие контекста. В нашем исследовании мы расскажем вам, как искать необходимую информацию в отсутствие контекста, основываясь только на приблизительном написании предмета поиска.

Стоит отметить, что при исследовании мы будем опираться на классический подход в поиске похожих слов по написанию, на расстояние Левенштейна. Наша задача состоит в том, что мы хотим получить качество поиска не хуже, чем у расстояния Левенштейна, при этом ускорив время обработки одного запроса.

Методы

Кодирование векторов. Перед обучением модели необходимо привести слова в нормальную форму, если она существует, например, применить лемминг и стемминг. Стемминг – самый простой метод приведения разных форм одного и того же слова к общей основе. Она заключается в отсечении конца слова. Лемматизация – процесс приведения слова к его нормальной (канонической) форме.

В нашей работе мы использовали библиотеку Fast Text, которая позволяет нам создавать векторное представление слов на основе N-грамм символов. Для этого мы разделили обработанные слова на символы и обучили две модели на биграммах (парах символов) и униграммах (один символ) с 3-символьным контекстным окном для обеих моделей, что позволяет нам оценить вероятность появления текущего символа при условии, что вокруг него есть другие символы (пары символов). Мы решили ограничиться биграммами и униграммами, чтобы не раздувать словарь, размер словаря для модели Fast Text будет рассмотрен ниже, а также, чтобы не увеличивать вычислительную сложность, мы хотим получить хороший базовый подход с низкими затратами на подготовку данных и низкой вычислительной мощностью. Наш подход основан на статистических данных. Таким образом, мы можем определить вероятность появления текущего символа на основе его соседей:

$$P(s_i | s_{i-3}, s_{i-2}, s_{i-1}, s_{i+1}, s_{i+2}, s_{i+3}), \quad (1)$$

где s_i – символ (или пара символов в случае биграмм).

Основной принцип, лежащий в основе Fast Text, заключается в том, что морфологическая структура слова несет важную информацию о значении этого слова. Таким образом, Fast Text позволяет получить векторное представление для каждого токена, токен в данном случае либо биграмма, либо униграмма. Эта структура не учитывается традиционными методами кодирования слов, такими как Word2Vec [9], которые предоставляют уникальное кодирование слов для каждого отдельного слова. Это особенно важно

для морфологически богатых языков (русский, немецкий, турецкий), в которых одно слово может иметь большое количество морфологических форм, каждая из которых может встречаться редко, что затрудняет обучение правильному кодированию слов.

Таким образом, мы получаем векторное представление для каждого символа или пар символов, усредняем полученные представления и получаем вектор всего слова:

$$V = \sum_i^{\text{length}(\text{word})} v_i, \quad (2)$$

где V – вектор слова, v_i – вектор символа (пары символов) внутри слова.

В результате мы получаем два вектора для одного слова: один из модели, обученной на униграммах, второй из модели, обученной на биграмах. Теперь, чтобы получить результирующий вектор для слова, можем найти среднее значение этих векторов. Также можно подобрать коэффициенты для каждого из векторов для конкретных данных (набора ошибок) таким образом, чтобы обеспечить лучший поиск правильной формы в базе знаний. Для поиска по векторной базе знаний лучше всего использовать классический метод k соседей.

Наша задача – реализовать быстрый и оптимальный, с точки зрения объема памяти и времени выполнения, подход к поиску слов из словаря с неправильным написанием, поэтому в своей работе мы используем только униграммы и биграмы, для триграмм и четырехграмм размер словаря модели FastText станет очень большим. Например, предположим, что в каком-то языке всего 20 букв, тогда размер словаря для биграмм будет равен

$$S = 20^2, \quad (3)$$

где S – размер словаря.

Для четырехграмм размер словаря будет соответственно:

$$S = 20^4. \quad (4)$$

Векторный поиск. Векторный поиск – неотъемлемая часть нашей системы, которая позволит искать правильное написание слова в базе знаний. Обычно поиск выполняется с использованием классического алгоритма k соседей, где в качестве расстояния берется либо косинусная мера близости, либо евклидово расстояние. В нашей работе мы использовали евклидово расстояние.

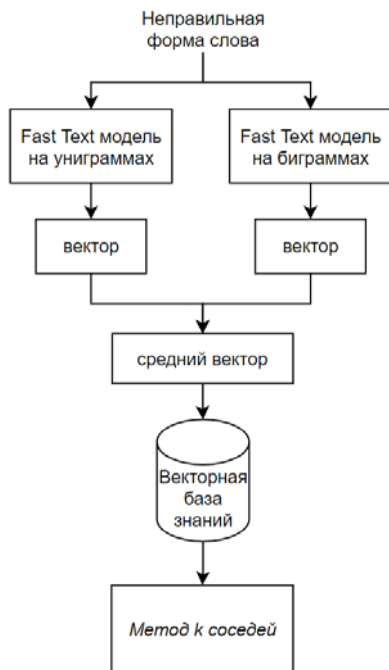


Рис. 1. Архитектура

Для векторного поиска было построено бинарное дерево, которое позволяет быстро выполнять поиск по нему. Каждый узел содержит вектор и ссылки на его дочерние узлы. В корне дерева хранится вектор, который является средним значением всех векторов в наборе. Затем дерево делится на две части: левую и правую, каждая из которых содержит векторы, расположенные ближе к среднему вектору в левом или правом направлении соответственно. Этот процесс продолжается рекурсивно для всех узлов до тех пор, пока не будут достигнуты листья дерева, содержащие сами векторы. Поиск ближайших соседей по индексу на основе бинарного дерева выполняется следующим образом:

1. Поиск в корне дерева: начинается с корня дерева. Вектор, сохраненный в корне, является средним значением всех векторов в наборе. Это позволяет быстро оценить, в каком направлении искать ближайшего соседа.
2. Выбор ветви: в зависимости от того, находится ли искомый вектор ближе к среднему вектору слева или справа, для дальнейшего поиска выбирается соответствующая ветвь дерева.

3. Рекурсивный поиск: процесс повторяется для выбранной ветви до тех пор, пока не будет достигнут лист дерева, содержащий сам вектор. Если вектор находится на листе, его расстояние до нужного вектора можно вычислить напрямую. Если вектор находится в узле, а не в листе, то поиск продолжается в соответствующих дочерних узлах.

4. Сбор результатов: после того, как найден ближайший сосед, процесс может быть завершён необходимым количеством раз, чтобы найти дополнительных соседей.

Эксперимент

В рамках исследования мы использовали открытый набор данных из Kaggle [<https://www.kaggle.com/datasets/harshwardhanpj/smartphones>], посвященный моделям мобильных телефонов.

Мы обучили две модели Fasttext на униграммах и биграмах [10], длина выходного вектора – 300 элементов, с контекстным окном из 3 униграмм (биграмм) в обоих направлениях, количество эпох обучения равно 40. В качестве результирующего вектора мы взяли среднее арифметическое между двумя векторами (один для униграмм, второй вектор для биграмм).

Чтобы протестировать работоспособность системы, мы сгенерировали набор нечетких запросов с ошибками в названиях мобильных телефонов на основе статистики ошибок при печати на английской раскладке клавиатуры [2]. Необходимо сгенерировать наборы данных для разного количества ошибок в словах, чтобы понять, при скольких орфографических ошибках системе не удастся выдать правильную форму слова в k ближайших соседях. Чтобы оценить нашу систему, мы рассмотрим точность нахождения нужного слова в k соседних словах. Для быстрого векторного поиска мы построили бинарное дерево.

В табл. 1 показана точность определения правильной формы слова для соответствующего количества соседей.

Таблица 1

Таблица точности для k соседей

Количество ошибок	1 сосед	5 соседей	10 соседей	15 соседей
1 ошибка	0,939	0,951	0,984	0,992
2 ошибки	0,84	0,843	0,924	0,948
3 ошибки	0,742	0,742	0,85	0,886
4 ошибки	0,574	0,574	0,70	0,753

Для наглядности построим график, показывающий полученные результаты. Можно проследить зависимость: по мере увеличения количества соседей повышается точность нахождения правильной формы слова в базе знаний среди k соседей. Также наблюдается тенденция к ухудшению качества поиска правильной формы с увеличением количества ошибок в слове.

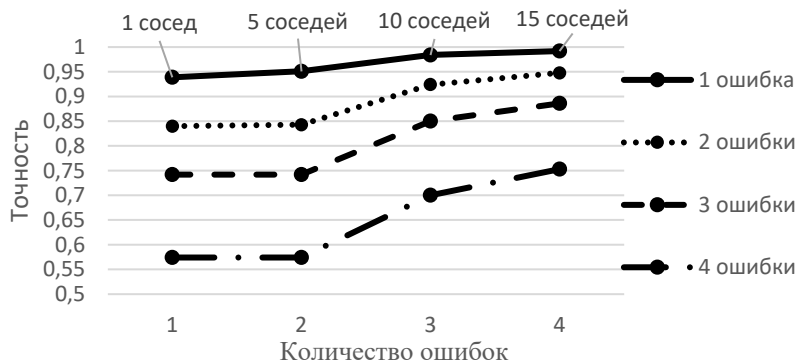


Рис. 2. Точность нахождения правильной формы слова среди k соседей для различного количества ошибок

Обсуждение

Проанализировав результаты тестирования, можно отметить следующее: для случаев, когда в слове допущено 1–2 ошибки, мы почти всегда сможем найти правильную форму. Для случаев, когда допущено 3–4 опечатки, ансамбль моделей справляется гораздо хуже.

К сильным сторонам этого подхода можно отнести следующее:

1. Для обучения модели необходим только исходный набор данных: для построения этого базового подхода вам не нужно размечать данные, все, что вам нужно, – правильные формы всех слов.

2. По сравнению с расстоянием Левенштейна мы получили заметное увеличение скорости поиска наиболее релевантного варианта коррекции.

Среди слабых сторон можно выделить высокую чувствительность к цифрам: обнаружить опечатки в цифрах чрезвычайно трудно.

В результате были достигнуты следующие цели:

1. Эффективность кодирования слов. Использование моделей, обученных на n -граммах, показывает потенциал в кодировании слов, которые являются не словарными, особенно если они содержат орфографические

ошибки. Это подтверждается анализом точности рекомендаций при различных уровнях орфографических ошибок.

2. Реализована векторная база знаний для быстрого поиска правильной формы.

В целом анализ показывает, что использование моделей, обученных на n -граммах, является многообещающим подходом к решению проблемы кодирования и исправления не словарных слов, написанных с ошибкой, особенно при отсутствии контекста. Однако для того, чтобы в полной мере оценить эффективность и применимость этого подхода, необходимы дополнительные исследования и тестирование в различных реальных условиях использования.

Выводы

В ходе наших исследований нам удалось разработать систему, которая позволяет гораздо быстрее искать правильную форму слова, в нашем случае это названия смартфонов, с учетом ошибок. Способность нашей системы находить правильную форму слова при 1–2 опечатках в словах составляет 90%, что является хорошим показателем для базового подхода. Говоря о скорости поиска правильной формы в базе знаний, мы в 706 раз быстрее расстояния Левенштейна, эти результаты получены при размере базы знаний в 980 элементов. Например, для того чтобы найти ближайшее слово по расстоянию Левенштейна, нам нужно вычислить его для каждого примера, мы получим сложность $O(n)$, в случае векторного поиска нам достаточно один раз получить вектор и найти ближайшее слово в бинарном дереве, мы получаем сложность $O(\log_2(n))$.

Список литературы

1. Hannah B. and Marjan C., Albert Ludwigs University “Efficient Fuzzy Search in Large Text Collections” // ACM Transactions on Information Systems. 2013. P. 1–2.
2. Yuzuru H., Yoshihiko O., Yamada-Hisao “AN ANALYSIS OF THE STANDARD ENGLISH KEYBOARD” // Dept. of Information Science Faculty of Science, University of Tokyo. P. 1–7.
3. Pennington, J., Socher, R., Manning, C.: GloVe: Global vectors for word representation. // Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP). 2014. P. 1532–1543.
4. Mikolov, T., Chen, K., Corrado, G., Dean, J.: Efficient estimation of word representations in vector space // 1st International Conference on Learning Representations, ICLR 2013. 2013. P. 1–5.
5. Nguyen, H., Dang, T.B., Nguyen, L.M.: Deep learning approach for vietnamese consonant misspell correction. In: Nguyen, L., Phan, X., Hasida, K., Tojo, S. (eds.) Computational Linguistics – 16th International Conference of the Pacific Association for Computational Linguistics, PACLING 2019, Hanoi, Vietnam, October 11–13, 2019,

- Revised Selected Papers. Communications in Computer and Information Science, V. 1215, P. 497–504. Springer (2019) Khan, Sardar Mudassar, “Clean Architecture Used in Software Development”, 2023.
6. Hieu Ngo Trung, Duong Tran Ham, Tin Huynh, Kiem Hoang “A Combination of BERT and Transformer for Vietnamese Spelling Correction” // The Saigon International University, Vietnam 2024. P. 1–7.
 7. Kanev A. I. and Terekhov V. I. “Evaluation issues of query result ranking for semantic search” // Computer Science and Control systems department, Bauman Moscow State Technical University, Moscow.
 8. Terekhov V. I., Kanev A. I. “Semantic Search System with Metagraph Knowledge Base and Natural Language Processing” // Proceeding of the 28th Conference of Fruct Association 2021. P. 652–658.
 9. Armand Joulin Edouard Grave Piotr Bojanowski Tomas Mikolov Facebook AI Research “Bag of Tricks for Efficient Text Classification” // arXiv 2016. P. 1–5.
 10. Nguyen Thi Xuan Huong , Tran-Thai Dang , The-Tung Nguyen , and Anh-Cuong Le “Using Large N-gram for Vietnamese Spell Checking” // Advances in Intelligent Systems and Computing. 2015. P. 1–5.

А.А. ЯКОВЕНКО, М.Р. ЛОЙКОНЕН, М.М. ЗАСЛАВСКИЙ

Санкт-Петербургский государственный электротехнический университет «ЛЭТИ»
yakovenant@mail.ru

МЕТОДЫ ВЫДЕЛЕНИЯ ПРИЗНАКОВ В ЗАДАЧЕ РАСПОЗНАВАНИЯ ЭМОЦИЙ ПО ГОЛОСУ

Рассматривается задача анализа эмоционально окрашенной русскоязычной устной речи с целью выделения характеристических признаков для задач классификации эмоций по голосу. Исследован подход к выделению признаков с позиции методов уменьшения размерности. Наилучшие результаты продемонстрировал метод S-UMAP, позволивший получить компактные и информативные описания, решить проблему мультиколлинеарности, ускорить процесс обучения и улучшить точность бинарной классификации базовых эмоций, выражающих радость и печаль.

Ключевые слова: *распознавание эмоций, классификация, выделение признаков, уменьшение размерности, анализ устной речи.*

Введение

Задача оценки эмоционального состояния диктора является существенным аспектом при анализе устной речи. Данная задача относится к аффек-

тивным вычислениям в рамках исследований искусственного интеллекта [1]. Технологии, реализующие аффективные вычисления, призваны улучшить пользовательский опыт в процессе человеко-машинного взаимодействия, используя информацию об эмоциональном состоянии. Эмоции в речи относятся к паралингвистической информации, использование которой позволит более точно интерпретировать смысл высказываний, верно расставлять акценты и знаки препинаний.

Известно, что смысловое содержание речевого сообщения в процессе устной передачи выражается не только лингвистически, на долю паралингвистической составляющей приходится не менее трети всей информации [2]. Поскольку эмоциональный окрас и интонация устной речи могут частично или полностью перекрывать смысловое содержание, полученное на основе анализа лингвистической составляющей [3], совместное использование паралингвистического анализа речи и языковых моделей позволяет улучшить точность распознавания эмоций [4]. Однако под воздействием эмоциональных состояний характер устной речи существенно изменяется, что затрудняет применение данного подхода на практике. Указанные особенности стимулируют растущий интерес к распознаванию эмоциональных состояний субъекта, как к самостоятельной задаче.

На текущем этапе существует ряд актуальных проблем, среди которых прежде всего – проблемы выбора корпуса эмоциональной речи, формализации эмоциональных состояний и выделения признаков эмоциональных состояний из речевого сигнала. Проблема выбора корпуса эмоциональной речи связана со спецификой данных: сбор и разметка эмоциональной речи чрезвычайно трудоемкая и плохо формализуемая задача. Формальные модели выделяют различное число базовых эмоций, а прикладные задачи требуют различной детализации в распознавании эмоциональных состояний [1; 5]. В представленной работе для проверки гипотез рассматривается простейший эмоциональный окрас устной речи: позитивный (радость) и негативный (печаль).

Для формализации и выделения признаков эмоциональных состояний на речевом сигнале может применяться ручной отбор признаков, автоматическое выделение признаков с использованием моделей искусственных нейронных сетей или выделение признаков путем вычисления вторичных значений исходных данных для уменьшения избыточности, сложности моделей и лучшей интерпретируемости данных [5]. В представленной работе рассматриваются методы выделения признаков эмоциональных состояний на речевом сигнале и их влияние на точность различных алгоритмов классификации.

Выбор корпуса эмоциональной речи

Речевой корпус имеет решающее значение для распознавания эмоциональных состояний, поскольку выбор исходных данных взаимосвязан с техническими характеристиками системы распознавания, целевым числом категорий эмоций и результирующей точностью классификации. Корпусы эмоциональной речи можно разделить на три типа по способу выражения эмоций: спонтанные, вызванные и искусственные [5].

Корпусы спонтанной эмоциональной речи являются наиболее востребованными, поскольку передают максимально естественные и достоверные эмоции. Для их сбора используются скрытые механизмы регистрации, позволяющие субъекту реагировать естественно. Однако процесс сбора сложный и результат не всегда предсказуемый, поскольку эмоциональные проявления могут быть неожиданными.

Для сбора корпусов вызванной эмоциональной речи создается ситуация, провоцирующая диктора на определенную эмоциональную реакцию. Недостаток данного типа речевых корпусов в том, что эмоциональная реакция зачастую оказывается недостаточно выразительной.

Корпусы искусственной эмоциональной речи, т.н. синтетические, собираются путем имитации эмоций актерами. При записи диктор произносит набор фраз с требуемой интонацией. Запись производится на протяжении некоторого периода времени с перерывами и в различное время суток, чтобы учесть голосовую вариативность диктора. Такой подход позволяет собрать необходимый объем фиксированных видов эмоциональных проявлений. Но актерские эмоции более выразительны и утрированы, в сравнении с естественными, что необходимо учитывать.

В качестве исходного набора данных для исследования выбран речевой корпус Dusha [6], так как это один из самых больших и доступных наборов данных эмоциональной речи на русском языке. Частота дискретизации аудио составляет 16 кГц, а общая продолжительность 350 часов. Набор данных содержит 4 типа эмоций: радость, печаль, гнев и нейтральная речь, и состоит из двух частей: Crowd – устная речь с имитированными эмоциями, Podcast – разговорная речь с естественными эмоциями, полученная из записей различных подкастов. В работе использовалась выборка Crowd, так как в ней более сбалансированное распределение записей по эмоциям и интенсивность выражения достаточно высокая для проверки гипотез. На рис. 1 показана гистограмма классов эмоций для выборки Crowd.

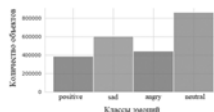


Рис. 1. Распределение эмоций в выборке Crowd речевого корпуса Dusha

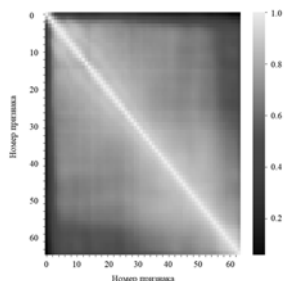


Рис. 2. Корреляционная матрица

Предварительный анализ и обработка данных

В представленном исследовании для проверки гипотез рассматривается задача бинарной классификации по двум категориям эмоций: радость и печаль. Была сделана подвыборка по 200 000 экземпляров класса каждой эмоции из исходного набора. Аудиозаписи в наборе данных представлены предварительно рассчитанными мел-частотными коэффициентами, выступающими в роли исходных признаков. Таким образом, каждая аудиозапись образует множество векторов, содержащих 64 коэффициента, посчитанных через фиксированный промежуток времени. Также каждая аудиозапись сопровождается меткой эмоции, которая была выражена.

На рис. 2 представлена корреляционная матрица признаков. По результатам предварительного анализа видно, что коэффициенты соседних частотных полос довольно сильно коррелируют, что приводит к проблеме мультиколлинеарности. Избыточные коэффициенты усложняют модель и процесс сходимости, приводят к неустойчивости и переобучению. Чтобы избавиться от мультиколлинеарности, необходимо исключить из модели порождающие признаки. Далее с этой целью будут рассмотрены методы выделения признаков, преобразующие и понижающие размерность пространства признаков.

Значения коэффициентов лежат в диапазоне от -80 до 0 дБ. На рис. 3 представлены примеры гистограмм распределения значений коэффициентов. Из приведенных графиков видно, что коэффициенты, значения которых лежат в области -80 дБ, выделяются на фоне остальных. Особенно это заметно для старших коэффициентов. Это говорит о наличии пропущенных значений в данных, которые соответствуют речевым паузам и не несут полезной информации о характеристиках голоса, но вносят определенные искажения. Для снижения влияния коэффициентов со значениями -80 дБ

применялась пороговая фильтрация к объектам выборки. Пример полученных в результате выполнения соответствующего преобразования гистограмм продемонстрирован на рис. 4.

Поскольку неизвестно, какие коэффициенты обладают наибольшей значимостью для распознавания эмоций, данные предварительно обрабатывались путем применения стандартизации и масштабирования, чтобы уравновесить вклад каждого коэффициента.

Постановка задачи и сравнение алгоритмов классификации

Запишем формальное описание задачи классификации эмоциональных состояний. Пусть дана обучающая выборка $X^L = (x_i, y_i)_{i=1}^L$, где $X^L \subset X$.

Данные представлены множеством объектов X и множеством ответов $Y = \{0, 1\}$, характеризующих позитивные (радость) и негативные (печаль) эмоциональные проявления. Существует целевая функция $z: X \rightarrow Y$, значения которой $y_i = z(x_i)$ известны на конечном подмножестве объектов обучающей выборки X^L . Необходимо восстановить зависимость z по X^L . Для этого нужно построить алгоритм a , приближающий целевую функцию $z(x_i)$ на множестве X , которое разбивается на классы эмоций $C_y = \{x_i \in X: z(x_i) = y_i\}$:

$$a: X \rightarrow \{C_y | y \in Y\}. \quad (1)$$

Восстановление зависимости z между X и Y по частным эмпирическим данным X^L характеризует задачу классификации, как задачу обучения по прецедентам. Для построения алгоритма a фиксируется аппроксимирующая модель восстанавливаемой зависимости $A: a \in A$, которая представляет параметрическое семейство решающих функций вида

$$A_w(X) = \{a(X, w) | w \in W\}, \quad (2)$$

$$a(X, w) = \text{sgn } g(X, w), \quad (3)$$

где $a: X \times W \rightarrow Y$, W – множество параметров модели, $g(X, w)$ – дискриминантная функция. Задача решается путем подбора параметров модели.

Как известно, выделяют линейные, байесовские, метрические, ансамблевые и нейросетевые модели A классификации. Используя исходный набор данных, рассмотрим основные методы для определения базового алгоритма a , относительно которого будем сравнивать методы выделения признаков эмоциональной речи, на примере следующих представителей: логистическая регрессия (англ. LR), наивный байесовский классификатор (англ. NB), алгоритм k ближайших соседей (англ. k-NN), метод опорных векторов (англ. SVC), дерево решений (англ. DT), случайный лес (англ. RF), градиентный бустинг (англ. GB) и многослойный персептрон (англ. MLP) [4; 5].

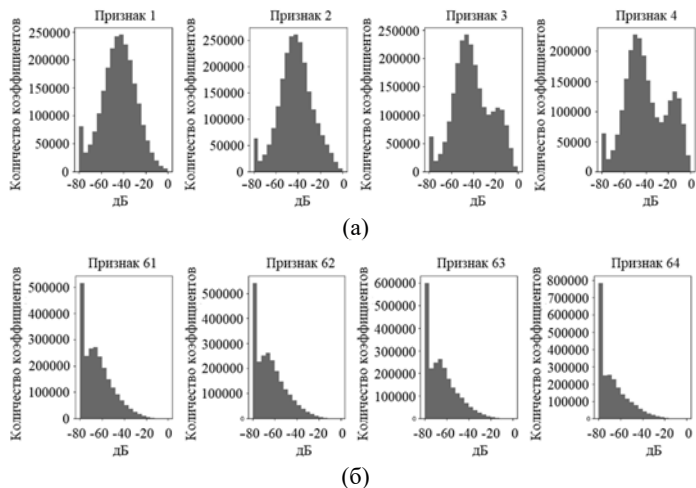


Рис. 3. Гистограммы распределения значений
а) первых четырех коэффициентов б) последних четырех коэффициентов

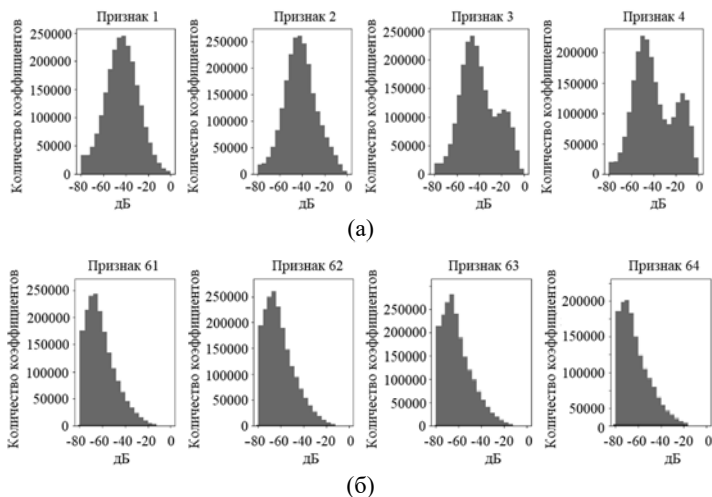


Рис. 4. Гистограммы распределения значений после восстановления пропусков
а) первых четырех коэффициентов б) последних четырех коэффициентов

Оценка полученных моделей на независимых данных выполнялась с использованием метода перекрестной проверки. Для этого исходные данные

случайным образом разбивались на 10 подвыборок. Обучение моделей выполнялось на 9 подвыборках, а на оставшейся производилось тестирование. Указанным образом процедура повторялась 10 раз, что позволило получить оценку эффективности классификации с наиболее равномерным использованием всей выборки данных.

Результаты обучения базовых моделей на задаче распознавания эмоций по голосу, с использованием исходного набора данных, представлены в таблице 1. В строке Train таблицы показана точность классификации на обучающей выборке, Test – на тестовой, а CV (от англ. Cross-Validation) демонстрирует результаты валидации в результате перекрестной проверки. Гиперпараметры каждой модели подбирались с помощью перекрестной проверки во избежание переобучения и для контроля производительности.

Из полученных результатов видно, что наибольшую эффективность продемонстрировал метрический алгоритм k -NN, точность на тестовом наборе и кросс-валидации при этом составила 92%. Также высокие результаты продемонстрировала нейросетевая модель MLP и ансамблевый классификатор RF с точностью 87 и 88% соответственно. Полученные оценки на исходном наборе данных будем использовать для анализа эффективности методов выделения признаков.

Таблица 1

Результаты распознавания эмоций на исходном наборе данных

Метод	<i>LR</i>	<i>NB</i>	<i>k-NN</i>	<i>SVC</i>	<i>DT</i>	<i>RF</i>	<i>GB</i>	<i>MLP</i>
Train	0.66	0.62	0.94	0.84	0.81	0.90	0.81	0.92
Test	0.65	0.62	0.92	0.83	0.79	0.88	0.80	0.87
CV	0.65	0.62	0.92	0.83	0.79	0.88	0.80	0.87

Методы преобразования пространства признаков

Методы преобразования пространства признаков можно разделить на линейные (сингулярное разложение, анализ главных или независимых компонент), позволяющие сохранить исходные закономерности и линейные взаимосвязи в пространстве меньшей размерности, и нелинейные (многомерное шкалирование, автокодеры), сохраняющие более сложные нелинейные взаимосвязи. Рассмотрим популярные методы функционального преобразования исходного пространства для выделения линейно независимых признаков и понижения размерности.

Метод анализа главных компонент (англ. PCA) [7] представляет вспомогательное преобразование, позволяющее определить эффективную размерность пространства признаков. Это классический алгоритм, применяемый для выделения признаков, понижения размерности и визуализации

данных. В PCA строится новое признаковое описание с минимальной размерностью, по которому исходные признаки восстанавливаются линейным преобразованием с минимальными погрешностями, максимизируя при этом общую дисперсию исходных данных. Основными преимуществами алгоритма на практике являются его детерминированность и скорость работы.

В отличие от PCA, идея алгоритма стохастического вложения соседей с t -распределением (англ. t -SNE) [7; 8] состоит не в максимизации дисперсии напрямую, а в нахождении такого пространства, в котором расстояние между объектами будет сохраняться. В процессе строится распределение вероятностей по всем парам объектов выборки данных. Каждая условная вероятность $P(i | j)$ означает, насколько близки объекты x_i и x_j при гауссовом распределении вокруг x_i . В качестве метрики близости при этом используется евклидово расстояние. Для получения отображения в пространство меньшей размерности также вводится распределение, описывающее насколько объекты в новом пространстве близки, используя t -распределение Стюдента с одной степенью свободы. В качестве критерия схожести распределений используется дивергенция Кульбака – Лейблера. Для минимизации используется алгоритм градиентного спуска, что позволяет найти компоненты в пространстве новой размерности. Особенность алгоритма в том, что он позволяет сохранить локальную структуру в данных.

Другой стохастический метод, развивающий идею t -SNE, – алгоритм унифицированной аппроксимации и проекции многообразия (англ. UMAP) [7; 9]. Алгоритм в процессе решает две подзадачи. На первом этапе выполняется построение графа, в котором ребрами соединены k ближайших объектов. Объекты рассматриваются как вершины графа. Веса ребер определяются расстоянием между вершинами – чем они ближе, тем вес больше. Далее задача состоит в том, чтобы построить аналогичный граф, но в пространстве меньшей размерности. Для этого оптимизируется низкоразмерное представление с использованием алгоритма градиентного спуска.

UMAP выгодно отличается скоростью работы и возможностью проецировать новые объекты на уже построенное представление. Кроме того, UMAP позволяет использовать разметку набора данных и получать представление, оптимизированное под имеющуюся информацию о классах объектов. Эта особенность работает как для полностью размеченных, так и для частично размеченных данных. Такая модификация называется Supervised UMAP (S-UMAP). Указанное свойство имеет важное значение для процедуры выделения признаков.

Оценка результатов классификации на выделенных признаках

Рассмотрим задачу проекции признаков в двумерное рабочее пространство. Такая проекция является наиболее удобной для визуального анализа и интерпретации полученных распределений. В таблице 2 представлены результаты распознавания эмоций по голосу на обучающей выборке, с применением методов выделения признаков. Аналогичным образом таблица 3 демонстрирует результаты перекрестной проверки. В целом, точность на обучающей выборке незначительно превосходит точность на валидации, что является ожидаемым результатом.

Применение метода PCA для выделения признаков показало прирост в точности классификации для алгоритмов DT и RF, в остальных случаях наблюдается снижение точности. Похожая ситуация с остальными алгоритмами. Метод t-SNE позволил сохранить на прежнем уровне точность классификатора k-NN. Алгоритм UMAP показал в среднем более качественное выделение признаков, сравнительно с предыдущими методами, существенно потеряли в точности только SVC и MLP.

Наилучшие результаты были получены с использованием метода S-UMAP. Благодаря возможности использования разметки, данный алгоритм построил наиболее качественное отображение в двумерное пространство. Точность классификации была увеличена для всех рассматриваемых алгоритмов как на обучающей выборке, так и в ходе перекрестной проверки. Учитывая, что S-UMAP является вычислительно эффективным и позволяет проецировать новые объекты на уже имеющееся представление, можно утверждать, что данный метод хорошо подходит для практики.

Таблица 2

Результаты классификации эмоций на обучающей выборке с выделением признаков

Метод	<i>LR</i>	<i>NB</i>	<i>k-NN</i>	<i>SVC</i>	<i>DT</i>	<i>RF</i>	<i>GB</i>	<i>MLP</i>
<i>PCA</i>	0.59	0.60	0.77	0.64	0.91	0.97	0.69	0.65
<i>t-SNE</i>	0.59	0.59	0.93	0.69	0.93	0.98	0.77	0.73
<i>UMAP</i>	0.56	0.58	0.88	0.68	0.94	0.99	0.76	0.70
<i>S-UMAP</i>	0.99	0.97	0.99	0.98	0.99	0.99	0.99	0.99

Таблица 3

Результаты перекрестной проверки с выделением признаков

Метод	<i>LR</i>	<i>NB</i>	<i>k-NN</i>	<i>SVC</i>	<i>DT</i>	<i>RF</i>	<i>GB</i>	<i>MLP</i>
<i>PCA</i>	0.58	0.59	0.76	0.62	0.90	0.97	0.68	0.63
<i>t-SNE</i>	0.58	0.57	0.91	0.68	0.93	0.98	0.76	0.72

<i>UMAP</i>	0.56	0.57	0.87	0.67	0.94	0.98	0.74	0.70
<i>S-UMAP</i>	0.98	0.96	0.99	0.97	0.98	0.99	0.98	0.98

Заключение

В работе решалась задача распознавания эмоциональных состояний дикторов по голосу на примере бинарной классификации базовых эмоций, выражающих радость и печаль. Исходные данные представлены мел-частотными коэффициентами из корпуса русскоязычной эмоциональной речи Dusha. Поскольку исходные данные являются достаточно сложными и избыточными, изучалось применение методов выделения признаков в пространстве меньшей размерности.

Метод S-UMAP позволил существенно улучшить качество классификации для всех рассмотренных алгоритмов, в сравнении с показателями, полученными на исходных данных. Так, наибольшую эффективность классификации исходных данных продемонстрировал алгоритм k-NN, точность которого на перекрестной проверке составила 92%. Использование метода S-UMAP позволило увеличить это значение до 99%, что говорит о решении проблемы мультиколлинеарности в данных.

Список литературы

1. Яковенко А.А., Мубаракшина Р.Т. Обзор подходов к проблеме распознавания эмоций по параметрам устной речи // Системный анализ в проектировании и управлении. Сборник научных трудов Международной научно-практической конференции. – Санкт-Петербург : Изд-во Политехнического университета. – 2019. – Т. 1. – С. 392–397.
2. Карпов А.А., Кайа Х., Салах А.А. Актуальные задачи и достижения систем паралингвистического анализа речи // Научно-технический вестник информационных технологий, механики и оптики. 2016. Т. 16. № 4. С. 581–592.
3. Верхоляк О.В. Автоматическое распознавание эмоциональных состояний дикторов по голосовым характеристикам и тональности текста высказывания: дисс. канд. техн. наук / ФГАОУ ВО «ИТМО» и Университет Ульма, Санкт-Петербург, 2021.
4. Schuller B. Voice and speech analysis in search of states and traits / Computer Analysis of Human Behavior. Eds. A.A. Salah, T. Gevers. Springer, 2011. P. 227–253.
5. Singh Y.B. A systematic literature review of speech emotion recognition approaches // Neurocomputing. 2022. V. 492. P. 245–263.
6. Kondratenko V., Sokolov A., Karpov N., Kutuzov O., Savushkin N., Minkin F. Large Raw Emotional Dataset with Aggregation Mechanism // arXiv. – 2022.
7. Wang Y., Huang H., Rudin C., Shaposhnik Y. Understanding how dimension reduction tools work: an empirical approach to deciphering t-SNE, UMAP, TriMAP, and PaCMAP for data visualization // Journal of Machine Learning Research. – 2021. – V. 22. – N 201. – P. 1–73.

8. Van der Maaten L., Hinton G. Visualizing data using t-SNE // Journal of machine learning research. – 2008. – V. 9. – N 11. – P. 1–73.
9. McInnes L., Healy J., Melville J. Umap: Uniform manifold approximation and projection for dimension reduction // arXiv preprint arXiv:1802.03426. – 2018.

**М.Ю. МАЛЬСАГОВ¹, Е.В. МИХАЛЬЧЕНКО¹,
Я.М. КАРАНДАШЕВ^{1,2}, Л.И. СТАМОВ¹**

¹Федеральный научный центр Научно-исследовательский институт системных исследований Российской академии наук, Москва

²Российский университет дружбы народов (РУДН), Москва

malsagov@niisi.ras.ru, mikhailchenkolena@yandex.ru,

karandashev@niisi.ras.ru, lyubens@mail.ru

АППРОКСИМАЦИЯ ПРОЦЕССА ДЕТОНАЦИИ ВОДОРОДА В ВОЗДУШНОЙ СРЕДЕ ИСКУССТВЕННОЙ НЕЙРОННОЙ СЕТЬЮ

Моделирование горения является ключевым аспектом полномасштабного 3D-моделирования современных и перспективных двигателей аэрокосмической промышленности. В настоящей работе исследуется возможность аппроксимации химической кинетики нейронными сетями. Используется более современный и точный, по сравнению с предыдущими работами, кинетический механизм, который применяется для получения обучающих данных. Слегка модифицировав некоторые параметры нейронной сети, нам удалось снять все ограничения на исходные данные, описанные в предыдущих работах, например, связь давления с временным шагом. Теперь нейронная сеть может предсказывать состояние системы через любой промежуток времени, а средняя квадратичная ошибка осталась на прежнем уровне – 10^{-4} .

Ключевые слова: численное моделирование химических процессов, горение, детонация, нейронные сети, глубокое обучение.

Введение

Во многих случаях задача расчета газодинамических процессов в двигателях, в энергетических устройствах и в других условиях включает моделирование физических и химических взаимодействий. Полномасштабное трехмерное моделирование современных и перспективных ракетных двигателей часто использует топливо, которое описывается детальными химическими механизмами, включающими сотни, а иногда и тысячи эле-

ментарных реакций между десятками и сотнями видов. Объединение вычислительной гидродинамики и моделирования горения является вычислительно дорогостоящей задачей из-за необходимости объединения нелинейных уравнений химической кинетики с гидродинамикой и явлениями переноса. Узким местом всего подхода является расчет химической кинетики, который требует интеграции этапа решения жесткой системы обыкновенных дифференциальных уравнений в численный процесс с пошаговым управлением по времени. Прямая интеграция детальных химических механизмов со многими видами и реакциями в переходные трехмерные потоки в настоящее время ограничена доступными вычислительными ресурсами [1–3].

Для сокращения времени вычислений во многих программах используются сокращенные кинетические механизмы [4–6]. Но даже при значительном сокращении кинетический шаг занимает значительную часть общего времени вычислений.

Другими способами сокращения времени вычислений являются методы, в которых при моделировании рассчитывается только термохимическое состояние пространства и сохраняется в таблицах для дальнейшего повторного использования [7–8].

В настоящее время подход нейронных сетей обещает ускорить такого рода вычисления. Методы машинного обучения получили широкое распространение благодаря многочисленным достижениям в алгоритмах и доступности вычислительных мощностей. Имеется множество работ по применению этого метода к задачам химической кинетики [9–12].

Процесс детонации происходит при определенных условиях и за очень малый промежуток времени. Поэтому большинство работ проводятся при определенных ограничениях: фиксированное давление, заданное соотношение между ключевыми компонентами топлива, фиксированный временной шаг наблюдений и т.д. В нашей работе не было задано каких-то ограничений, что создало множество сложностей как и в создании обучающей выборки, так и с процессом обучения. При случайной генерации начальных условий (давления, температуры, молярных плотностей и т.д.) процесс крайне редко включал в себя детонацию. Чтобы как-то исправить эту ситуацию, мы ввели эмпирическую зависимость между давлением в системе и временным интервалом между соседними наблюдениями. Данный подход подробно был описан в [13]. Однако, несмотря на хорошие результаты, подобная фиксация временного шага от давления оказалась не всегда применима при моделировании газовой динамики. Поэтому в этой работе мы отказались от прошлого результата и решили получить более гибкую сеть,

сохранив, по возможности, достигнутые преимущества. В настоящей работе удалось получить аналогичный результат по точности работы сети, сохранив архитектуру и процесс обучения сети, но учтя вариативности временного шага наблюдений. Так же был заменен кинетический механизм на более современный и точный.

Постановка задачи

В настоящей работе моделируется детонация водородно-кислородной смеси в воздухе (азот и аргон). Задав начальные давление и температуру, через равные промежутки времени производится расчёт процесса горения смеси. Для расчета химических преобразований используется сокращенный кинетический механизм (таблица 1), предложенный в работе [14]. В процессе работы механизма образуется 9 водородно-кислородных соединений (H_2 , O_2 , H_2O , OH , H , O , HO_2 , H_2O_2 , OH^*), а также азот и аргон, которые влияют на динамику газов и скорости реакций, но соединений не образуют.

Химическая кинетика описывает реакции распада и объединения химических элементов в каждой точке пространства, описываемых системой дифференциальных уравнений:

$$\frac{\partial \mathbf{X}}{\partial t} = f(p, T, \mathbf{X}), \quad (1)$$

где p – давление, T – температура, $\mathbf{X}$ – вектор молярных плотностей элементов.

Именно вычисление системы (1) численными методами занимает около 90% временных ресурсов. Поэтому в данной работе предлагается использовать искусственные нейронные сети для аппроксимации химической кинетики.

Данные для обучения и тестирования нейронной сети

Задача, поставленная в (1), имеет единственное решение, полученное численно с использованием метода Новикова решения жестких систем обыкновенных дифференциальных уравнений [15]. Ее решение моделирует самовоспламенение горючей смеси водорода с воздухом, предварительно нагретым настолько, чтобы обеспечить быстрый процесс воспламенения в адиабатических и объемно ограниченных условиях. Важной частью численной реализации является полуаналитический метод вычисления якобиана для правой части системы уравнений.

Вычисление развития газовой смеси во всей камере во времени является сложной вычислительной задачей, даже в небольшом двумерном случае.

Поэтому все пространство камеры разбивается на маленькие одинаковые ячейки, в которых независимо вычисляется химическая кинетика.

Таблица 1

Сокращенный кинетический механизм [14].

№	Формула	№	Формула	№	Формула
1	$H_2 + O_2 = HO_2 + H$	11	$H_2 + OH = H_2O + H$	21	$O + H = OH$
2	$H_2 + O_2 = OH + OH$	12	$OH + OH = H_2O_2$	22	$H + O = OH^*$
3	$H + O_2 = O + OH$	13	$H_2O_2 + O_2 = HO_2 + HO_2$	23	$OH^* + H_2O = OH + H_2O$
4	$H + O_2 = HO_2$	14	$H_2O_2 + OH = H_2O + HO_2$	24	$OH^* + H_2 = OH + H_2$
5	$HO_2 + O = OH + O_2$	15	$H_2O_2 + H = HO_2 + H_2$	25	$OH^* + N_2 = OH + N_2$
6	$HO_2 + H = OH + OH$	16	$H_2O_2 + H = H_2O + OH$	26	$OH^* + OH = OH + OH$
7	$HO_2 + H = H_2O + O$	17	$H_2O_2 + O = HO_2 + OH$	27	$OH^* + H = OH + H$
8	$HO_2 + OH = O_2 + H_2O$	18	$H_2O = H + OH$	28	$OH^* + Ar = OH + Ar$
9	$H_2 + O = H + OH$	19	$O_2 = O + O$	29	$OH^* = OH$
10	$OH + OH = H_2O + O$	20	$H_2 = H + H$	30	$OH^* + O_2 = OH + O_2$

Задавая начальные условия (давление, температура, молярные плотности веществ) в одной ячейке, можно, используя метод Новикова, вычислять изменения химической системы во времени.

В данной работе было сгенерировано 60 000 случайных начальных состояний системы. С помощью численного метода определялись молярные плотности веществ через равные интервалы времени (dt). Эти таблицы состояний системы во времени и использовались в дальнейшем для обучения и тестирования нейронной сети. Температура выбиралась от 250 К до 5000 К, давление от 10^4 Па до $2 \cdot 10^7$ Па. Размер временного шага выбирался случайно в диапазоне от 10^{-9} с до 10^{-5} с. Эти 60 000 примеров разбили на три набора данных: 5000 – тестовый набор, 5000 – валидационный набор, 50 000 – обучающий набор. На рис. 1 показан пример полученных данных.

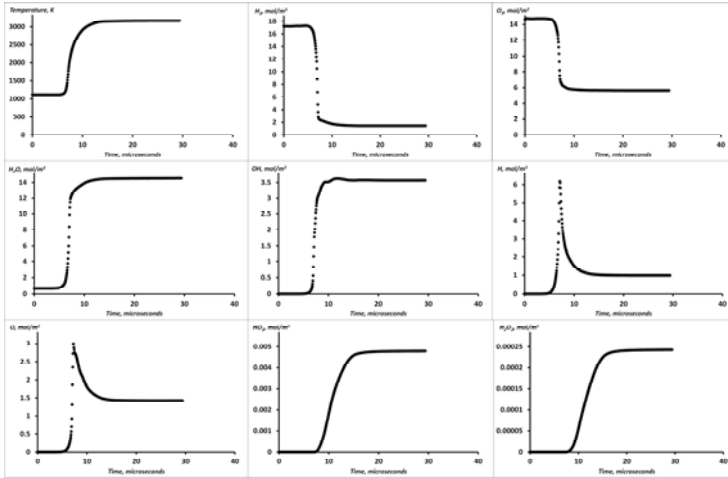


Рис. 1. Один пример из обучающего набора данных

Давление системы легко определяется из температуры T и молярных плотностей веществ χ , поэтому оно не использовалось в качестве компоненты входного вектора нейронной сети. Помимо температуры на вход подавался временной шаг dt , через который мы хотим получить ответ. Таким образом, входом и выходом сети были 13-мерные векторы $\tilde{X} = (dt, T, \chi)$.

Важно заметить, что каждая компонента этого вектора не только подчиняется своему закону во времени, но и меняется каждая в своем диапазоне величин. Тысячи Кельвин для температуры, от 10^{-9} с до 10^{-5} с для временного шага, от 10^{-15} до десятков и сотен для молярных плотностей. Поэтому ко всем данным применялась логарифмическая нормировка данных:

$$\mathbf{X}^* = \ln(\varepsilon + \tilde{\mathbf{X}}), \quad \mathbf{X} = \frac{\mathbf{X}^* - M(\mathbf{X}^*)}{\sigma(\mathbf{X}^*)}. \quad (2)$$

Здесь $M(\mathbf{X}^*)$ и $\sigma(\mathbf{X}^*)$ – среднее значение и стандартное отклонение (по компонентным), вычисленные на тренировочном наборе после логарифмирования, а $\varepsilon = 10^{-15}$ – коэффициент, чтобы избежать нулевых значений молярных плотностей. Подобная нормировка позволила избавиться от экспоненциальных величин, а все компоненты вектора выровнять около нулевого среднего значения с единичной дисперсией.

Архитектура нейронной сети

Минимализм – основной критерий для нашей нейронной сети. Можно построить нейронную сеть с сотней слоёв, последовательных или параллельных, и тысячами нейронов в слоях. Такая сеть, вероятно, будет работать с большой точностью, но вычислительно окажется медленнее большинства численных методов. Смысла в таком подходе нет совсем. Поэтому поиск баланса между объемом вычислений, точностью и устойчивостью решения является основной задачей.

На рисунке 2 представлена наилучшая найденная архитектура нейронной сети. Сеть состоит из 5 слоев. Первый и последний слой (в виде трапеции) переводят соответственно входной 13-мерный вектор в 100-мерное пространство и обратно. Средние три – полносвязные слои, по 100 или 150 нейронов (в отличие от [13], где все слои были одинаковые), с активационной функцией LeakyReLU($a=0.01$) после каждого. Присутствуют также прямые связи, передающие входной вектор без изменений на выход, как в остаточных нейронных сетях ResNet или Unet.

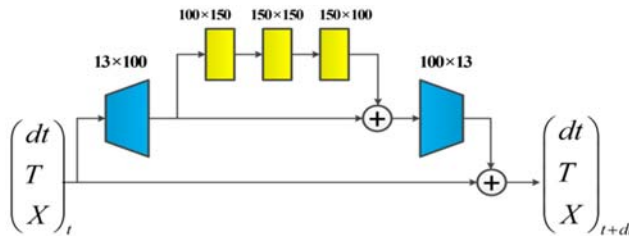


Рис 2. Архитектура нейронной сети

Сеть получилась простая и компактная, всего 55 613 обучаемых параметров. Каждый слой – это скалярное произведение матрицы весовых коэффициентов на входной вектор слоя:

$$\mathbf{h}^l = \mathbf{W}^l \mathbf{h}^{l-1} + \mathbf{b}^l. \quad (3)$$

Активационная функция:

$$\text{LeakyReLU} = \max \{ a \mathbf{h}^l, \mathbf{h}^l \}. \quad (4)$$

Здесь l – номер слоя, $\mathbf{h}^l$ – выход слоя l , $a=0.01$, $\mathbf{W}$ и $\mathbf{b}$ – обучаемые параметры нейронной сети.

Вычислительная сложность такой сети легко считается и составляет $\square 7 \square 10^6$ операций сложения и умножения. Сравнить нейронную сеть по количеству операций с численным методом невозможно. В отличие от сети количество операций в использовавшемся численном методе меняется

(например, шаг интегрирования подбирается в процессе и т.д.). Так же современные фреймворки и графические ускорители позволяют параллельно запускать одну нейронную сеть на сотнях и тысячах входных векторах. Например, тестовая выборка из 5000 примеров на 500 шагов каждый считается менее чем за секунду на видеокарте NVIDIA GeForce RTX 3090.

Модификация алгоритма обучения

Такая сеть может предсказывать только на один шаг по времени. Т.е. на вход сети подаётся начальная конфигурация $\mathbf{X}_t$, после чего сеть предсказывает изменённые значения через один шаг времени $\mathbf{X}_{t+dt}$, что сравнивается с желаемым значением $\hat{\mathbf{X}}_{t+dt}$, полученным численно. При этом сеть учится минимизировать разницу между $\mathbf{X}_{t+dt}$ и $\hat{\mathbf{X}}_{t+dt}$. С другой стороны, в процессе функционирования нам будет необходимо предсказание на несколько шагов вперёд, в идеале на несколько сотен или даже тысяч шагов. Такое предсказание можно получить, если обученную сеть запускать в рекуррентном режиме, т.е. выход сети подавать на вход и т.д. При этом очень важно, чтобы ошибка предсказания на один шаг была бы как можно меньше, поскольку при многократном запуске сети ошибка быстро накапливается и растёт как минимум линейно.

Как и в предыдущей работе [13], мы включили предсказание на несколько шагов в рекуррентном режиме в сам процесс обучения. Обучение сети производилось сразу на несколько шагов, т.е. выход сети подавался на вход некоторое количество раз n_{step} , а функция потерь вычислялась по накопленным данным:

$$Loss = \sum_{k=1}^{n_{step}} \frac{1}{k} MSE(\mathbf{X}_{t+dt}^k, \hat{\mathbf{X}}_{t+dt}^k), \quad (5)$$

где $\hat{\mathbf{X}}_{t+dt}$ и $\mathbf{X}_{t+dt}$ – эталонный и предсказанный сетью выходы.

Благодаря суммированию ошибок предсказания на несколько шагов вперёд, сеть автоматически учится учитывать рекуррентное накопление ошибки и минимизировать его. В экспериментах мы использовали $n_{step} = 50$ шагов.

Результаты

На рисунке 3 представлен результат работы обученной нейронной сети на одном примере из тестового набора. Из рисунка видно, что хоть и не идеально, но сеть обучилась аппроксимировать необходимые зависимости.

Среднее значение среднеквадратичной ошибки на тестовом наборе составило $\langle MSE \rangle \approx 10^{-4}$.

На рисунке 4 показан результат другого эксперимента с обученной сетью. Суть эксперимента в том, чтобы определить, как нейронная сеть научилась учитывать изменение временного шага наблюдений. Для этого с помощью численного метода определялось развитие одного и того же процесса через разные промежутки времени (рис. 4, верхний ряд). Так как количество шагов наблюдения не менялось, то при разных временных шагах реальный интервал наблюдения был различен. На рисунках видно, что при шаге $3 \cdot 10^{-7}$ сек между стартовым состоянием и следующей точкой уместятся почти все 500 точек наблюдения, полученные с шагом 10^{-9} сек и т.д. На рисунке 4, нижний ряд, показано, как нейронная сеть справилась с данной задачей. В целом, сеть учитывает изменение временного шага и с некоторой ошибкой дает подходящий результат. Определить, допустима ли такая погрешность, предстоит в дальнейшей работе, когда химическая кинетика будет объединена с газовой динамикой.

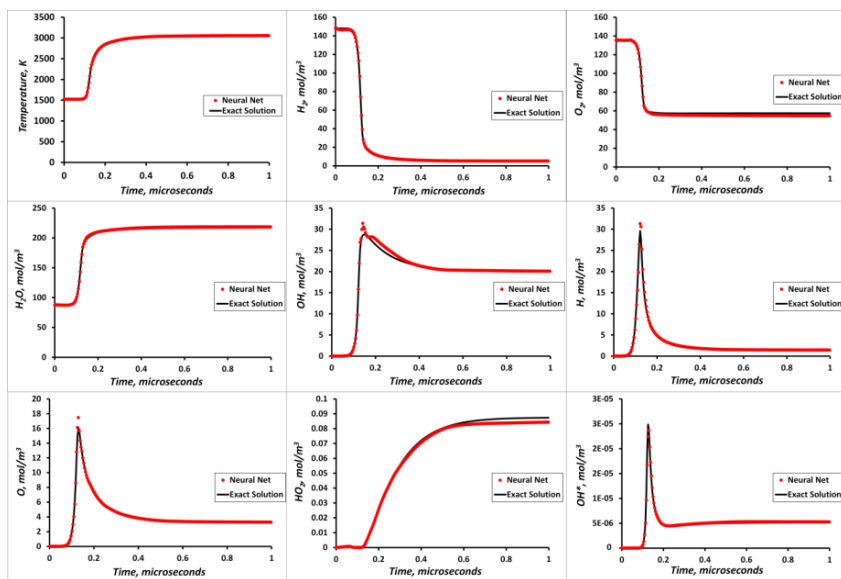


Рис 3. Тестовый пример. Аппроксимация нейронной сетью ($dt = 4 \cdot 10^{-9}$)

Заключение

В настоящей работе рассмотрена задача аппроксимации химической кинетики, детонация водорода в воздушной среде, с помощью нейронных сетей. Основной проблемой было учесть временной шаг, интервал времени, через который проводятся наблюдения развития процесса детонации газовой смеси.

Хотя нейронная сеть и показала свою способность решать поставленную задачу, на данный момент остается множество проблем, с которыми предстоит справиться для достижения наилучшего результата.

Во-первых, правильное формирование обучающего набора. До сих пор мы генерировали стартовые состояния случайным образом, что означало множество комбинаций как нереальных, так и не рабочих. Рисунки 1 и 4 хорошо демонстрируют недостатки такого подхода. Неправильный выбор временного шага может привести к тому, что в наблюдаемый диапазон не попадет процесс детонации вообще или из стартовой точки сразу попадем в конечное равновесное состояние. Все это означает горизонтальные линии без какой-либо информации, которой могла бы учиться нейронная сеть. Сейчас большая часть обучающих данных состоит из таких непрезентабельных примеров, что ухудшает весь результат.

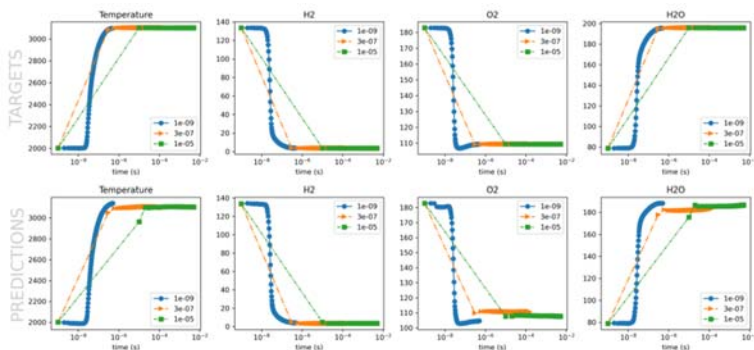


Рис 4. Изменение временного шага. Результаты получены из одного стартового состояния на 500 шагов с тремя различными шагами по времени dt : 10^{-5} с, $3 \cdot 10^{-7}$ с, 10^{-9} с. Вверху – полученные численные значения. Внизу – предсказанные нейронной сетью

Во-вторых, поиск оптимальной архитектуры. Несмотря на то, что предложенная нейронная сеть очень проста, всего пять небольших полносвязных слоев, в ней присутствует много параметров, которые можно варьировать в широких диапазонах и различных комбинациях. Количество слоев и

количество нейронов в слоях, типы активационных функций и их параметры, количество шагов прогнозирования – это лишь часть переменных, меняя которые получаются различные результаты [16]. Все это трудно решаемые проблемы, требующие больших временных и вычислительных затрат.

Решение этих проблем позволит качественно улучшить полученный результат.

Работа выполнена в рамках государственного задания ФГУ ФНЦ НИИСИ РАН по теме № FNEF-2024-0001 "Создание и реализация доверенных систем искусственного интеллекта, основанных на новых математических и алгоритмических методах, моделях быстрых вычислений, реализуемых на отечественных вычислительных системах" (1023032100070-3-1.2.1).

Список литературы

1. Pantano C. Direct simulation of non-premixed flame extinction in a methane–air jet with reduced chemistry // *J. Fluid Mech.* 2004. 514. 231–270.
2. Smirnov N. N., Betelin V. B., Nikitin V. F., Stamov L. I., Altoukhov D. I. Accumulation of errors in numerical simulations of chemically reacting gas dynamics. *Acta Astronautica*, 2015. 117. 338–355.
3. Smirnov N.N., Penyazkov O.G., Sevrout K.L., Nikitin V.F., Stamov L.I., Tyurenkova V.V. Onset of detonation in hydrogen-air mixtures due to shock wave reflection inside a combustion chamber // *Acta Astronaut.* 2018. V. 149. 77–92.
4. Lovas T. Automatic generation of skeletal mechanisms for ignition combustion based on level of importance analysis // *Combust. Flame.* 2009. 156 (7). 1348–1358.
5. Koniavitis P., S. Rigopoulos, W.P. Jones, A methodology for derivation of RCCE-reduced mechanisms via CSP // *Combust. Flame.* 2016. 183. 126–143.
6. Lu T., Law C.K. Toward accommodating realistic fuel chemistry in large-scale computations // *Prog. Energy Combust. Sci.* 2009. 35 (2). 192–215.
7. Mikharchenko E. V., Nikitin V. F., and Goryachev V. D. Simulation of the operation of a detonation engine // *Lecture notes in mechanical engineering*, 2020. P. 98–107. DOI:10.1007/978-3-030-92144-6_7
8. Pope S.B., *Combust. Theory Modell.* 1997. 1 (1). 41–63.
9. Tonse S.R., Moriarty N.W., Brown N.J., Frencklach M. // *Israel J. Chem.* 1999. 39 (1). 97–106.
10. Peng, W.Y. and Pinkowski, N.H. «Efficient and Accurate Time-Integration of Combustion Chemical Kinetics Using Artificial Neural Networks», – 2017, – <https://cs229.stanford.edu/proj2017/final-reports/5241836.pdf>.
11. Sharma, A.J. Johnson, R.F. Kessler, D.A. Moses, A. Deep Learning for Scalable Chemical Kinetics // *Proceedings of the AIAA Scitech 2020 Forum*, Orlando, FL, USA, – 6–10 January 2020. P.0181.
12. Weiqi Ji, Sili Deng. KiNet: A Deep Neural Network Representation of Chemical Kinetics // <https://arxiv.org/abs/2108.00455>. – 2021.
13. Мальсагов М.Ю., Карандашев Я.М., Михальченко Е.В., Никитин В.Ф.. Аппроксимация химической кинетики искусственной нейронной сетью // XXIV

Международная научно-техническая конференция "НЕЙРО-ИНФОРМАТИКА-2022": сборник научных трудов. – Москва : МФТИ, 2022. ISBN 978-5-7417-0823-1. С. 37–45.

14. Tereza A.M., Medvedev S.P., Smirnov V.N. Experimental study and numerical simulation of chemiluminescence emission during the self-ignition of hydrocarbon fuels // Acta Astronautica. – 2019. – V. 163, Part A. – P. 18–24.

15. E.A. Novikov. Investigation of (m,2)-methods of stiff systems calculation // Vychislitel'nye tekhnologii (Calculation technologies). 2007. V. 12, 5. 103.

16. Malsagov, M.Y., Mikhilchenko, E.V., Karandashev, I.M., Nikitin V.F. Simulation of Hydrogen Combustion at Different Pressures Using a Neural Network // Combustion, Explosion, and Shock Waves. 2023. V. 59. P. 145–150. <https://doi.org/10.1134/S0010508223020041>.

А.В. КУЗЬМЕНКО, В.С. КИРЕЕВ

Национальный исследовательский ядерный университет «МИФИ», Москва
andrey_kuzmenko2907@mail.ru

МЕТОДОЛОГИЯ РЕКОНСТРУКЦИИ ГРАФА ЗНАНИЙ В МАЛОЧИСЛЕННЫХ КОЛЛЕКЦИЯХ ТЕКСТОВЫХ ДОКУМЕНТОВ

В работе формализуется задача реконструкции графа знаний текстовых документов. Производится анализ существующих подходов решения подзадач реконструкции графа, таких как извлечения реляционных троек, нормализация именованных сущностей и сопоставление сущностей. Предлагается трехступенчатая методология реконструкции графа в условиях малочисленных коллекций текстовых документов. В основе методологии лежит взаимодействие с большой языковой моделью посредством промпт-инженерии. В работе проводятся эксперименты, иллюстрирующие способность большой языковой модели извлекать компоненты графа знаний без переобучения модели под целевую задачу. Совокупность малого объема требуемых данных и использование высокоуровневого API делает данную технологию извлечения знаний доступной в малоизученных предметных областях.

Ключевые слова: *нейронные сети, граф знаний, большие языковые модели, реляционная тройка, нормализация именованных сущностей, сопоставление сущностей.*

Введение

Извлечение знаний – это фундаментальная задача обработки текстов на естественном языке. Наиболее сжатым средством представления знаний является граф. Узлы графа – это сущности, ребра отношения между ними. Граф знаний обладает мощной выразительной способностью, позволяет моделировать сущности, понятия и отношения [1–3]. Применение графа знаний нашлось в поиске знаний [4], ответов на вопросы [5], рекомендации знаний и визуализации знаний [6].

Задача реконструкции графа представима в виде комбинации задачи извлечения реляционных троек, сопоставления текстовых сущностей и нормализации сущностей. Обычно для решения каждой из подзадач используют отдельную модель. И при этом для достижения наилучших результатов в своей задаче каждая модель обычно реализуется в двухэтапной парадигме: 1) обучение модели структуре языка, как правило, через задачу псевдоязыкового моделирования на большом объеме не размеченных текстовых данных; 2) обучение под целевую задачу на существенно меньшем объеме размеченных текстовых документов. Первый этап позволяет существенно сократить затраты на дорогостоящую разметку. Однако для построения подсистем высокого уровня точности все равно требуются тысячи размеченных текстовых документов, что в большинстве практических приложений невозможно.

В данной работе предлагается методология реконструкции графа знаний на наборах данных порядка сотен и даже десятков с использованием больших языковых моделей. Описываются проблемные стороны валидации такого рода подходов, предлагаются способы их решения.

Реконструкция графа знаний текстовых документов

Реконструкция графа знаний довольно широкого плана. Таксономия задач и методов их решения представлена в [7]. В данной работе задача реконструкции графа трактуется как идентификация всех требуемых сущностей и определение связей между ними. Данная задача может считаться решенной, если решены следующие три подзадачи: извлечение реляционных троек, сопоставление сущностей из разных источников, отображение сущностей из источников в реальные.

Извлечение реляционных троек

Цель извлечения реляционных троек RTE (англ. Relational triplet extraction) – поиск набора тройных множеств $Y = \{(s_1, r_1, o_1), (s_2, r_2, o_2), \dots, (s_k, r_k, o_k)\}$, извлеченного из текстового документа $w = \{w_1, w_2, \dots, w_m\}$, в результате решения задачи:

$$p(Y|w, \theta) = p(n|w) \prod_{i=1}^k p(Y_i|w, Y_{i \neq j}, \theta), \quad (1)$$

где $S_i = [w_{t_{start}}^i, \dots, w_{t_{end}}^i]$ – сущность, определенная как субъект;

t_{start}, t_{end} – индексы токенов начала и конца сущности соответственно;

$o_i = [w_{t_{start}}^i, \dots, w_{t_{end}}^i]$ – сущность, определенная как объект;

$p(n|w)$ – моделирует размер набора реляционных троек;

$p(Y_i|w, Y_{i \neq j}, \theta)$ – моделирует тройку Y_i , при условии, что она связана с другими тройками $Y_{i \neq j}$,

θ – параметры модели.

Систематизации существующих подходов предложены в [8–9]. У [7] она включает четыре направления: (1) конвейерные методы, которые состоят из четко разграниченных двух шагов, на первом из текста извлекаются сущности, на втором идентифицируется отношение между ними [9–12]; (2) табличные методы, ключевая технология которых опирается на заполнение таблицы сущностей [13–16]; (3) методы, использующие дополнительную маркировку текстовых последовательностей [17–19]; (4) методы преобразования последовательностей [20–27]. Наиболее интересными в условиях малых наборов размеченных данных являются подходы, использующие большие языковые модели [24–28].

Сопоставление сущностей

Задача сопоставления сущностей ЕМ (англ. Entity Matching) сформулирована в [29–30]. Пусть A и B – два источника данных. A имеет атрибуты $(A_1, A_2, \dots, A_n)$, и мы обозначаем записи как $a = (a_1, a_2, \dots, a_n) \in A$. Аналогично, B имеет атрибуты $B = (B_1, B_2, \dots, B_m)$, и мы обозначаем записи как $b = (b_1, b_2, \dots, b_m) \in B$. Источник данных — это набор записей, а запись — это кортеж, соответствующий определенной схеме атрибутов. Атрибут определяется предполагаемой семантикой его значений. Таким образом, $A_i = B_j$ тогда и только тогда, когда значения a_i из A_i предназначены для передачи той же информации, что и значения b_j из B_j , и конкретный синтаксис значений атрибута не имеет значения. Атрибуты также могут иметь связанные с ними метаданные (например, имя), но это не влияет на равенство между ними. Назовем кортеж атрибутов $(A_1, A_2, \dots, A_n)$, схемой источника данных A и соответственно для B . Цель сопоставления сущностей — найти максимально возможное бинарное отношение $M \subseteq A \times B$ МАВ такое, что a и b относятся к одному и тому же объекту для всех $(a, b) \in M$.

Традиционные подходы сопоставления сущностей обычно включают 4 стадии [30–35]: 1 – предобработку данных, необходимую для приведения

данных из разных источников к единому формату; 2 – разработку схемы сопоставления, предназначенную для определения набора атрибутов для сравнения сущностей; 3 – фильтрацию кандидатов, предназначенную для снижения вычислительных затрат, поскольку процедура сопоставления имеет сложность $O(n * m)$; 4 – финальное парного сравнение кандидатов. Наибольшую эффективность, с точки зрения точности, показывают подходы на базе псевдоязыковых моделей [36]. Однако для их синтеза требуются тысячи обучающих примеров. Добиться хорошей производительности в условиях малых наборов данных, позволяют подходы на базе больших языковых моделей [37–38].

Нормализация сущностей

Заключительным этапом в реконструкции графа знаний является отображение текстовых сущностей в реальные NEN (англ. Named Entity Normalization). Мы выделили три группы текстовых сущностей: 1 – явные, наиболее простые для извлечения из текста, представленные в начальной форме; 2 – неявные – представленные в других склонениях слова; 3 – косвенные, наиболее сложные для извлечения, представленные в виде упоминания. Рассмотрим фрагмент, взятый из набора данных NEREL [39].

«Пресс-служба американской военной базы, расположенной в международном аэропорту "Манас" столицы Киргизии, отказалась комментировать данный инцидент с участием американских военных».

В данном текстовом фрагменте явная текстовая сущность – «Манас», неявная – «Киргизии», косвенная «американской», им соответствуют реальные сущности «Манас», «Киргизия», «Соединенные Штаты Америки». Нормализация сущностей, на наш взгляд, является наиболее трудоемкой задачей, поскольку требует от системы не только понимания текстового фрагмента, но и внешние знания о природе сущности. Поэтому стандартные инструменты лемматизации [40] недостаточны, особенно для нормализации косвенных сущностей. Такой природой знаний, на наш взгляд, обладают большие языковые модели.

Большие языковые модели

Современные большие языковые модели являются представителями архитектуры трансформер [41]. При этом, как правило, наибольшим успехом обладают системы – декодеры [42]. Современные подходы обучения больших языковых моделей обычно включают три стадии.

Первая стадия – обучение задаче классического языкового моделирования. Это – одна из наиболее популярных техник обучения с самоконтролем SSL (англ. self-supervised learning). Данный этап является самым вычисли-

тельно затратным и вместе с тем наименее требовательным к качеству исходных данных. Однако количество этих данных огромно, количество очищенных, обработанных текстовых документов исчисляется терабайтами. Это позволяет модели изучить структуру языка.

Вторая стадия SFT (с англ. supervised fine-tuning) производится на размеченных данных вида пара: инструкция–ответ. Количество таких пар уже исчисляется сотнями тысяч штук, но существенно повышается требование к их качеству.

Заключительная стадия – выравнивание (англ. alignment). Используемый формат данных схож со второй стадией. Но теперь модель учится не предсказывать ответ, а ранжировать имеющиеся. Используя элементы обучения с подкреплением, модель уточняет знания, полученные на второй стадии, учится обращать внимание на детали.

Большое количество знаний, которое впитала в себя большая языковая модель, позволяет ей быстро перенастраиваться под разные задачи буквально на нескольких примерах (англ. few-shot learning). И даже решать задачи для класса примеров, которые не присутствовали в обучении (англ. zero-shot learning). Еще одно интересное свойство больших языковых моделей – решать задачи посредством передаваемого запроса (контекста) (англ. prompt-engineering). Исследованию последнего в задаче реконструкции графа посвящена данная работа.

Методология

Описание исходных данных. В своих экспериментах мы используем, наиболее крупный, публично доступный, русскоязычный набор данных NEREL [39]. Он состоит из 746/94/93 документов Russian Wikinews, разделенных на обучающее, валидационное и тестовое множества соответственно. В нем представлено 29 типов сущностей и 49 типов отношений.

Модель. В исследовании мы используем GigaChat-lite. Данная версия модели поддерживает длину контекста 8192 токена, чего достаточно для проведения наших экспериментов. Доступ осуществлялся через API (англ. application programming interface).

Предлагаемая методология реконструкции графа знаний состоит из трех шагов: 1) поиск оптимального запроса-контекста; 2) обработка выходов большой языковой модели; 3) замер финального качества после обработки. Жизненный цикл нашего подхода иллюстрирован на рисунке 1. Наши эксперименты проводились над конкретным типом реляционных троек: объект – наименование населенного пункта, отношение – расположение города в стране, субъект – страна, представленные в NEREL как {CITY, COUNTRY, LOCATED_IN}. Всего таких троек в исследуемом

наборе данных 139, 13, 27 для обучающей, валидационной и тестовых частях соответственно. Для получения наиболее статистически значимых результатов мы используем валидационную часть для подбора контекста, тестовую – для постобработки и наиболее представленную тренировочную часть – для оценки финального качества.

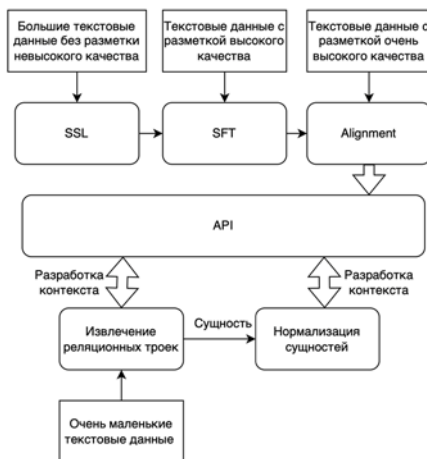


Рис. 1. Жизненный цикл реконструкции графа знаний

Поиск оптимального контекста. Задачей этого этапа является подбор контекста, посредством которого модель способна извлечь требуемые сущности и отношение. Проведенный анализ набора данных NEREL показал, что каждая из реляционных троек находится в рамках одного предложения. Поскольку мы не можем достоверно проверить, насколько наша методология позволяет обрабатывать документы с несколькими предложениями, мы упрощаем задачу, обрабатывая каждое из предложений отдельно. В наших экспериментах тестировались наиболее общие формулировки контекста, наиболее удачной оказалась конкатенация префикса:

«В следующем предложении необходимо выделить наименование населенного пункта и страну, в которую он входит»

и предложения – кандидата. Ниже показаны примеры ответов большой языковой модели для ряда предложений.

Исходное предложение: *«После окончания срока работы на посту президента в январе супруги арендовали особняк в стиле эпохи Тюдоров в престижном районе столицы США – Калорама, где долгое время селились дипломаты, лоббисты и политики».*

Ответ модели: «Наименование населенного пункта: Калорама. Страна: США».

Исходное предложение: «Полиция задержала 5 человек по делу о перестрелке в армянском Гаваре, в результате которой 28 апреля 2 человека были убиты, еще четверо – госпитализированы».

Ответ модели: «Наименование населенного пункта: Гавар. Страна, в которую входит населенный пункт: Армения».

Исходное предложение: «Семья президента Обамы также владеет домом в Чикаго — третьем по величине городе в США, где Обама приобрёл значительную политическую поддержку перед тем, как дважды стал главой Белого дома.»

Ответ модели: «Наименование населенного пункта: Чикаго. Страна, в которую входит населенный пункт: США».

Обработка выходов языковой модели. Анализируя результаты, полученные на первом этапе, мы столкнулись с двумя проблемами:

- ответы модели не имеют единого формата, даже в условиях низкой температуры;
- размеченные сущности, как правило, находятся не в начальной форме.

Для минимизации первого фактора мы обогащаем префикс требуемым форматом вывода: «В следующем предложении необходимо выделить наименование населенного пункта и страну, в которую он входит, вывести в формате JSON». Ниже показаны ответы большой языковой модели, для ранее представленных предложений:

`{"населенный_пункт": "Калорама", "страна": "США"}`

`{"населенный_пункт": "Гавар", "страна": "Армения"}`

`{"населенный_пункт": "Чикаго", "страна": "США"}`

Преобразование исходной разметки к некоторой форме для расчета метрик качества оказалось наиболее трудоемкой задачей. Поскольку в большом количестве документов нет явных формулировок принадлежности населенного пункта к какой-либо стране. Например, предложение:

«Полиция задержала 5 человек по делу о перестрелке в армянском Гаваре, в результате которой 28 апреля 2 человека были убиты, еще четверо – госпитализированы».

В исходной разметке сущности представлены словами: «Гаваре» и «армянском». И если в случае слова «Гаваре» осуществить преобразование к валидной начальной форме названия города возможно стандартными методами морфологического анализа, как, например, `ru morphology2` [42], то для «армянском» осуществить такое преобразование невозможно. Кроме того,

зачастую сущности – это именованные объекты, что еще сильнее усложняет применение данного инструмента. Мы также экспериментировали с инструментами, основанными на морфологическом анализе с использованием контекста [43], которые довольно неплохо решали задачу в случае прямых и неявных сущностей, однако на косвенных совсем не работали. Поэтому в этапе постобработки мы также используем большую языковую модель. Мы отправляем в модель запросы следующего вида:

«Преобразуй слово “х” в название страны, с которой оно связано. Верни только название страны. Если название населенного пункта “у” не в начальной форме, преобразуй в начальную, иначе оставь как есть. Верни только название населенного пункта».

для получения начальных форм слов - наименований. Такое преобразование исходной разметки может породить случаи, в которых модель в силу неспособности генерировать корректно название города будет подгонять под свои ответы, однако такое поведение крайне маловероятно и в наших экспериментах не наблюдалось.

Мы также проводили эксперименты с синтезом одноэтапного подхода с более подробным, детальным контекстом. Наши результаты показали, что такая методология с текущим уровнем развития LLM менее производительна как с точки зрения времени подбора контекста, так и с точки зрения качества извлечения.

Оценка результатов. Анализ ошибок. По результатам наших экспериментов, полученная система обнаруживала более 45% реляционных троек с отображением в реальные сущности. Мы выделили следующие группы ошибок:

- Блокировка модели, при запросах не удовлетворяющих требованиям запросов, связанных с политикой компании, предоставляющей API.
- Смежные конвертации. Например, модель определяет название страны как США, при этом конвертирование исходной разметки выдало Соединенные Штаты Америки. Ливия – Ливийская народная республика.
- Ошибки определения страны или населенного пункта.

Выводы

В работе сформирована задача реконструкции графа знаний текстовых документов. Проанализированы подходы в решении смежных задач: извлечение реляционных троек, сопоставление сущностей, нормализации именованных сущностей. Предложена методология реконструкции графа зна-

ний, не требующая больших наборов размеченных текстовых данных. Проведены эксперименты, оценивающие способность большой языковой модели извлекать знания из текстов на естественном языке посредством контекста.

Список литературы

1. Xue Zheng, Bing Wang, Yunmeng Zhao, Shuai Mao, Yang Tang. A knowledge graph method for hazardous chemical management: Ontology design and entity identification. *Neurocomputing*, 2017, V. 77. P. 48–52.
2. A knowledge graph method for hazardous chemical management: Ontology design and entity identification. *Neurocomputing*, 2021. V. 430. P. 104–111.
3. Weidong Li, Rong Peng, Yaqian Wang, Zhihuan Yan. Knowledge graph based natural language generation with adapted pointer-generator networks. *Neurocomputing*, 2020. V. 382. P. 174–187.
4. Xutang Zhang, Xin Hou, Xiaofeng Chen, Ting Zhuang. Ontology-based semantic retrieval for engineering domain knowledge. *Neurocomputing*, 2013. V. 116. P. 382–391.
5. Shuguang Zhu, Xiang Cheng, Sen Su. Knowledge-based question answering by tree-to-sequence learning. *Neurocomputing*, 2020. V. 372. P. 64–72.
6. Tong Yu, Jinghua Li, Qi Yu, Ye Tian, Xiaofeng Shun, Lili Xu, Ling Zhu, Hongjie Gao. Knowledge graph for TCM health preservation: Design, construction, and applications. *Artificial Intelligence in Medicine*, 2017. V. 77. P. 48–52.
7. Hongbin Ye, Ningyu Zhang, Hui Chen, Huajun Chen. Generative Knowledge Graph Construction: A Review // *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, 2022. P. 1–17.
8. Кузьменко А.В., Киреев В.С. Классификация методов извлечения реляционных троек из текстов на естественном языке // *Материалы XXV международной научно технической конференции Нейроинформатика-2023*, 2023. С. 302–311.
9. Shang, et al. Relational Triple Extraction: One Step is Enough // *Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence (IJCAI-22)*, 2022. P. 4360–4366.
10. Zenelnko, D., Aone, C., Richardella, A., et al. Kernel methods for relation extraction. *Journal of Machine Learning Research*, 2003. V. 3. P 1083–1106.
11. Chan, S., Y., and Roth, D. Exploiting syntactico-semantic structures for relation extraction. *Proceedings of the 49th Annual Meeting of the Association for Computational Linguistics: Human Language Technologies*, 2011. P. 551–560.
12. Zhong, Z., and Chen, D., A. A frustratingly easy approach for entity and relation extraction. *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, 2021. doi: 10.18653/v1/2021.naacl-main-5.

13. Miwa, M., and Sasaki, Y. Modeling Joint Entity and Relation Extraction with Table Representation. *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*. 2014. P. 1858–1869.
14. Gupta, P., Schutze, H. and Andrassy, B. A Table Filling Multi-Task Recurrent Neural Network for Joint Entity and Relation Extraction. *Proceedings of COLING 2016, the 26th International Conference on Computational Linguistics: Technical Papers*. 2016. P. 2537–2547.
15. Zhang, M., Zhang, Y., Fu, G. End-to-End Neural Relation Extraction with Global Optimization. *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*. 2017. P. 1730–1740.
16. Miwa, M., Bansal, M. End-to-end Relation Extraction using LSTMs ON Sequences and Tree Structures. *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics*. 2016. V.1: Long Papers. P. 1105–1116.
17. Zheng, S., et al. Joint Extraction of Entities and Relations Based on a Novel Tagging Scheme. *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics*. 2017. V. 1: Long Papers. P.1227–1236.
18. Dai, D., et al. Joint Extraction of Entities and Overlapping Relations Using Position-Attentive Sequence Labeling. *Proceedings of the AAAI Conference on Artificial Intelligence*. 2019. P. 6300–6308.
19. Wei, Z., et al. A Novel Cascade Binary Tagging Framework for Relational Triple Extraction. *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*. 2020. P. 1476–1488.
20. Zeng, X., et al. Extracting Relational Facts by an End-to-End Neural Model with Copy Mechanism. *Proceedings of the 56th Annual Meetings of the Association for Computational Linguistic*. 2018. V. 1: Long Papers. P 506–514.
21. Zeng D., Zhang, H. and Liu, Q. CopyMTL: Copy Mechanism for Joint Extraction of Entities and Relations with Multi-Task Learning. *Proceedings of the AAAI Conference on Artificial Intelligence*. 2020. P. 9507–9514.
22. Sui, D., Zeng, X., et al. Joint entity and relation extraction with set prediction networks. *IEEE Transactions on Neural Networks and Learning Systems*. 2023, doi: 10.1109/TNNLS.2023.3264735.
23. Wen, X., et al. End-to-end entity detection with proposer and regressor. *Neural Proceedings Letters*. 2023, doi: 10.1007/s111063-023-11201-8.
24. Hu, Y., Ameer, I. and Zou, X. Zero-shot Clinical Entity Recognition using ChatGPT. 2024, doi: arxiv-2303.16416.
25. Yuan, C., Xie, Q. and Ananiadou, S. Zero-shot Temporal Relation Extraction with CharGPT. 2023, doi: 10.48550/arxiv.2304.05454.
26. Feng, P., Wu, H. and Yang, Z. Leveraging Prompt and Top-K Predictions with ChatGPT Data Augmentation for Improved Relation Extraction // *Applied Sciences* (ISSN 2076-3417). 2023. V. 13(23).
27. Dagdelen, J., Dunn, A., Lee, S., et al. Structured information extraction from scientific text with large language models // *Nature communications*. 2024. V. 15, N 1418.

28. Liang Xu, Changxia Gao, Xuetao Tian. Domain-control prompt-driven zero-shot relational triplet extraction // *Neurocomputing*. 2024. V. 574. Doi: 10.1016/j.neucom.2024.127270.
29. Hanna Köpcke, Erhard Rahm. Frameworks for entity matching: A comparison. *Data & Knowledge Engineering*. 2010. V. 69, issue 2. P. 197–210.
30. Nils Barlaug and Jon Atle Gulla. 2021. Neural Networks for Entity Matching: A Survey. *ACM Trans. Knowl. Discov. Data*. 15, 3 (April 2021). P. 1–37.
31. Peter Christen. 2012. *Data Matching: Concepts and Techniques for Record Linkage, Entity Resolution, and Duplicate Detection*. Springer Science & Business Media.
32. Vassilis Christophides, Vasilis Efthymiou, Themis Palpanas, George Papadakis, and Kostas Stefanidis. 2019. End-to-End Entity Resolution for Big Data: A Survey. (May 2019). *arXiv:cs.DB/1905.06397*
33. Muhammad Ebraheem, Saravanan Thirumuruganathan, Shafiq Joty, Mourad Ouzani, and Nan Tang. 2018. Distributed representations of tuples for entity resolution. *Proceedings VLDB Endowment* 11, 11 (July 2018). 1454–1467.
34. Sairam Gurajada, Lucian Popa, Kun Qian, and Prithviraj Sen. 2019. Learning-Based Methods with Human-in-the-Loop for Entity Resolution // *Proceedings of the 28th ACM International Conference on Information and Knowledge Management (CIKM '19)*. ACM, New York, NY, USA, 2969–2970.
35. Felix Naumann and Melanie Herschel. 2010. *An Introduction to Duplicate Detection*. Morgan and Claypool Publishers.
36. Yuliang Li, Jinfeng Li, Yoshihiko Suhara, Anhui Doan, and Wang-Chiew Tan. 2020. Deep Entity Matching with Pre-Trained Language Models. (April 2020). *arXiv:cs.DB/2004.00584*.
37. Ralph Peeters, Christian Bizer. Using ChatGPT for Entity Matching. *New Trends in Database and Information Systems*, 2023. V. 1850. P. 221–230.
38. Ralph Peeters, Christian Bizer. Entity Matching using Large Language Models. <https://arxiv.org/pdf/2310.11244>.
39. Artemova, E., et al. NEREL: A Russian Dataset with Nested Named Entities, Relations and Events // *International Conference Recent Advances in Natural Language Processing*. 2021, doi: 10.26615/978-954-452-072-4_100.
40. Korobov M. Morphological Analyzer and Generator for Russian and Ukrainian Languages. <https://arxiv.org/abs/1503.07283>.
41. Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. Attention is all you Need. // *Advances in Neural Information Processing Systems* 30, I Guyon, U V.
42. Hugo Touvron, Thibaut Lavril, Gautier Izacard. LLaMA: Open and Efficient Foundation Language Models. <https://arxiv.org/abs/2302.13971>.
43. URL: <https://github.com/natasha/natasha>.

В.И. КУЗНЕЦОВ, Д.А. ЮДИН

Московский физико-технический институт
(национальный исследовательский университет)
vd.kuznetcov@yandex.ru

РАСПОЗНАВАНИЕ ТРЁХМЕРНЫХ ОБЪЕКТОВ НА ПОСЛЕДОВАТЕЛЬНОСТИ ИЗОБРАЖЕНИЙ С ИСПОЛЬЗОВАНИЕМ НЕЙРОННЫХ ПОЛЕЙ ЯРКОСТИ

Распознавание трёхмерных объектов является важной и актуальной задачей в компьютерном зрении. Однако малая часть методов может учитывать геометрию сцены при детектировании объектов. Поэтому в данной работе рассматривается задача распознавания трёхмерных объектов на основе нескольких изображений с разных ракурсов и с использованием нейронных полей яркости. Одним из таких методов является NeRF-Det, который позволяет учитывать геометрию сцены при детектировании объектов. Данная работа исследует влияние различных экстракторов признаков на качество детекции модели. Эксперименты проводились на наборе данных ScanNetV2 по показателям mAP@.25 и mAP@.50.

Ключевые слова: *распознавание трёхмерных объектов, дифференцируемое представления сцены, нейронные сети.*

Введение

Трёхмерное распознавание объектов является одним из ключевых направлений в области компьютерного зрения и имеет широкий спектр применений, включая робототехнику, дополненную реальность и автономные транспортные средства. Большинство стандартных методов зависят от специализированных сенсоров, что накладывает дополнительные ограничения на их использование.

В последние годы была разработана технология нейронных полей яркости (Neural Radiance Fields, NeRF) [1], которая привлекла значительное внимание благодаря своей способности генерировать высококачественные трёхмерные представления объектов и сцен на основе набора двумерных изображений. Это делает её перспективным инструментом для решения задач трёхмерного компьютерного зрения, таких как реконструкция сцен, распознавание и трекинг объектов, а также для приложений, связанных с виртуальной и дополненной реальностью.

Применение нейронных полей яркости для распознавания трёхмерных объектов на основе изображений открывает новые перспективы благодаря независимости от специализированных сенсоров, которые накладывают

дополнительные ограничения и увеличивают стоимость системы. Примерами таких работ являются [2, 3].

Данное исследование продолжает развитие идей, представленных в работе [3], и направлено на анализ влияния различных экстракторов признаков на качество детекции модели. Было установлено, что современная архитектура сверточной нейронной сети способна демонстрировать более высокие показатели качества при обучении на меньших объемах данных, чем исходная модель. Эксперименты проводились на наборе данных ScanNetV2 по показателям mAP@.25 и mAP@.50.

Литературный обзор

Традиционные методы распознавания трёхмерных объектов можно разделить на несколько категорий: на основе ключевых точек, плотных трёхмерных облаков и глубоких нейронных сетях.

Методы, основанные на ключевых точках, включают такие подходы, как SIFT и SURF, которые извлекают ключевые особенности из изображений и сопоставляют их с трёхмерными моделями. Однако эти методы часто оказываются чувствительными к изменениям условий освещения и углов обзора [4]. Методы, использующие плотные трёхмерные облака, применяют данные от сенсоров, таких как LIDAR, для создания облаков точек и их последующей обработки. Эти методы демонстрируют высокую точность, но требуют значительных вычислительных ресурсов и наличия специализированных сенсоров [5]. Глубокие нейронные сети, такие как PointNet [6], показали значительный прогресс в распознавании трёхмерных объектов, однако они требуют больших объёмов данных для обучения и могут быть чувствительны к шуму.

NeRF представляет собой новый подход к моделированию трёхмерных объектов, основанный на использовании объёмного рендеринга для описания сцены. Основная идея заключается в представлении сцены в виде непрерывной функции, моделирующей плотность и цвет для каждой точки пространства в каждом направлении. В оригинальной работе в качестве функции использовались два многослойных перцептрона, а последующие модификации исследовали сочетание этой методики с такими архитектурами, как сверточные и трансформерные сети [7–9].

К преимуществам модели NeRF относят получение детализированной трёхмерной реконструкции сцены с фотореалистичным качеством, устойчивость к изменениям освещения и ракурсов, что делает его особенно полезным для приложений, где важна высокая визуальная достоверность. Кроме того, данный метод может быть адаптирован для решения различных задач, таких как генерация новых ракурсов сцены или синтез новых объектов [10].

На начальных этапах развития приложений NeRF большинство исследований использовали его косвенно для распознавания объектов: обучали детекторы на основе карт глубины [11, 12], применяли предварительную сегментацию [13] совместно с детекцией [14], либо увеличивали количество обучающих данных [15]. Со временем появились работы, предлагающие распознавание трёхмерных объектов с явным применением нейронных полей яркости [2, 3].

Сверточные нейронные сети в исследованиях по распознаванию трёхмерных объектов в основном применяются в качестве экстракторов признаков. Информативность извлечённых признаков на ранних этапах существенно влияет на итоговую точность детекции. В последних работах [3, 16] экстракторам признаков уделяется сравнительно меньше внимания, чем другим блокам в модели, и используется классическая архитектура ResNet [17], в частности ResNet-50 и ResNet-101. Однако с тех пор были разработаны модификации этих моделей [18–20], которые сохраняют преимущества оригинальных архитектур и при этом улучшают их производительность.

Постановка задачи

На вход нейросети G_Φ поступает 5D вектор-функция нейронных полей яркости F_Θ и поза наблюдения $P \in SE(3)$, на выходе ограничивающий параллелепипед $\hat{B}$ и его класс $\hat{p}$ из множества объектов J в текущем представлении:

$$\{(\hat{B}^j, \hat{p}^j)\}_{j=1}^J = G_\Phi(P; F_\Theta), \quad (1)$$

где $\hat{B} = \{x_{i=1}, y_{i=1}, z_{i=1}\}_{i=1}^{N_c}$ – координаты для каждого ограничивающего параллелепипеда, а N_c – количество угловых точек, в рассматриваемой задаче $N_c = 8$. В данном исследовании используется постановка задачи, предложенная в работе [21], с переопределением переменной G_Φ , теперь она может быть представлена как последовательный [2] или параллельный [3] ансамбль нейронных сетей.

Метод

Предложенный метод, основанный на NeRF-Det, предполагает применение более эффективного экстрактора признаков, выбор которого будет определён результатами экспериментов. Однако, прежде чем принимать решение, важно рассмотреть, почему архитектура ResNet была выбрана в качестве стандарта. Её ключевым преимуществом является использование остаточных связей (skip connections), которые значительно упрощают процесс обучения глубоких моделей, позволяя преодолевать проблему затухания градиента.

Следующие две работы представляют собой развитие архитектуры ResNet. В ResNeSt [20] был предложен механизм разделённого внимания (split-attention) и различные стратегии обучения, которые позволили более эффективно учитывать пространственные и каналные зависимости в изображениях. В свою очередь, ResNeXt [19] использует декомпозицию остаточного блока, что повышает эффективность извлечения признаков, позволяя модели обучаться на более детальных признаках изображений без значительного увеличения вычислительной сложности. RegNet [18] представляет собой семейство моделей, ориентированных на оптимизацию производительности нейронных сетей посредством более эффективной параметризации и архитектурных решений.

Экспериментальные результаты

Для проведения исследования влияния экстракторов признаков был выбран общепринятый набор данных ScanNetV2, предназначенный для различных задач компьютерного зрения. Все рассматриваемые архитектуры использовали одинаковые настройки модели NeRF-Det, чтобы обеспечить объективное сравнение. Обучение всех моделей проводилось с зафиксированным значением параметра под названием “seed”. Каждая модель была обучена на определённом количестве изображений, максимально возможном для данной архитектуры. Несмотря на то, что некоторые архитектуры являются предшественниками, они способны обработать большее количество данных. Это позволило дополнительно рассмотреть в исследовании вопрос о том, улучшает ли увеличение объёма обучающего набора качество модели. Также исследуется, требуется ли такое же количество данных для обучения более современных моделей, или они могут превосходить ранее разработанные архитектуры, обучаясь на меньшем объёме данных. Соотношение количества изображений к архитектуре экстракторов признаков распределилось следующим образом: RegNetX-3.2GF — 35, ResNet-50 — 37, ResNeXt-50 — 32, ResNeSt-50 — 32, предобученная модель NeRF-Det с ResNet-50, представленная авторами [3] — 50.

Основными показателями качества для оценки работы моделей были выбраны следующие: средняя точность по всем классам (mean average precision, mAP) с порогами пересечения над объединением (intersection over union, IoU), равными 0.25 и 0.5, обозначенные как mAP@.25 и mAP@.50 соответственно и размер общей модели.

В таблице 1 представлены результаты по проведенным экспериментам. Из неё видно, что NeRF-Det с RegNetX-3.2GF является эффективной моделью по размеру, но требует более детальной настройки или подбора определенной архитектуры RegNet. В таблице приведены две модели с использованием ResNet-50: первая была обучена в рамках данного исследования

на меньшем объёме данных, а вторая представлена авторами [3]. Экстрактор ResNeXt-50 является промежуточным вариантом между ResNet-50 и ResNeSt-50 по показателям качества и более оптимизированным по размеру общей модели. Из анализа таблицы следует, что NeRF-Det с ResNeSt-50 демонстрирует наивысшее качество среди рассмотренных моделей при относительно небольшом размере по сравнению с ResNet-50.

Таблица 1

Предложенные методы были обучены на меньшем объёме данных по сравнению с исходной моделью NeRF-Det

Метод	mAP@.25	mAP@.50	Размер, МБ
NeRF-Det (RegNetX-3.2GF)	21.93	8.03	1091
NeRF-Det (ResNet-50)	50.14	25.33	1201
NeRF-Det (ResNeXt-50)	50.66	25.40	1196
NeRF-Det (ResNeSt-50)	53.34	27.02	1224
NeRF-Det (ResNet-50, предобученные веса от авторов)	52.68	26.96	1201

Результаты сравнения предложенной модели NeRF-Det с ResNeSt-50 и предобученной с ResNest-50 от авторов [3] показаны в таблице 2, где классы ScanNetV2 были переведены следующим образом: “кабинет” – “cabinet”, “диван” – “sofa”, “книжная полка” – “bookshelf”, “стул” – “chair”, “занавеска” – “curtain”, “стол” – “table”, “письменный стол” – “desk”, “окно” – “window”, “картина” – “picture”, “мусорное ведро” – “garbagebin”, “кровать” – “bed”, “дверь” – “door”, “счётчик” – “counter”, “раковина” – “sink”, “холодильник” – “refrigerator”, “занавеска для душа” – “showercurtain”, “туалет” – “toilet”, “ванна” – “bathtub”. Для AP@.25 предложенное решение опережает оригинал на 4 класса, а в AP@.50 отстаёт на 2 класса, однако при детальном рассмотрении можно увидеть, что отставание в двух классах, таких как “раковина” и “ванна”, является незначительным.

Таблица 2

Сравнение предложенной модели NeRF-Det с архитектурой ResNeSt-50 и предобученной с ResNet-50, представленной авторами, на наборе данных ScanNetV2

Показатели качества	AP@.25		AP@.50	
Классы	NeRF-Det	Предложенная модель	NeRF-Det	Предложенная модель
Кабинет	39.77	39.21	15.11	14.11

Показатели качества	AP@.25		AP@.50	
	NeRF-Det	Предложенная модель	NeRF-Det	Предложенная модель
Диван	76.41	75.16	39.24	42.24
Книжная полка	48.57	41.66	14.16	12.58
Стул	75.70	76.59	44.74	46.32
Занавеска	21.34	29.63	2.68	6.07
Стол	55.78	59.73	39.16	35.71
Письменный стол	72.90	77.31	45.68	41.71
Окно	22.58	21.70	2.17	2.52
Картина	3.36	3.74	0.66	1.74
Мусорное ведро	31.02	35.78	11.99	15.81
Кровать	82.31	83.29	71.91	65.27
Дверь	32.14	32.94	7.54	5.91
Счётчик	52.94	49.29	11.19	16.29
Раковина	52.05	52.27	27.85	27.43
Холодильник	48.81	56.65	25.89	21.85
Занавеска для душа	63.94	53.36	7.65	4.25
Туалет	89.56	93.67	60.89	70.59
Ванна	79.01	78.12	56.69	56.00
Среднее	52.68	53.34	26.96	27.02

Заключение

В данной работе исследовано влияние четырёх архитектур экстракторов признаков на качество детекции модели в задаче распознавания трёхмерных объектов на основе изображений с использованием нейронных полей яркости. В результате исследования была предложена и реализована модель NeRF-Det на основе экстрактора признаков ResNeSt-50, которая, несмотря на обучение на меньшем объёме данных по сравнению с предобученным методом NeRF-Det на базе ResNet-50, продемонстрировала превосходство по ключевым показателям качества.

Дальнейшие исследования в этом направлении включают в себя рассмотрение экстракторов признаков, которые комбинируют преимущества нескольких архитектур и оптимизацию других блоков модели NeRF-Det.

Список литературы

1. Mildenhall B., Srinivasan P. P., Tancik M., Barron J. T., Ramamoorthi R., Ng R. Nerf: Representing scenes as neural radiance fields for view synthesis // Communications of the ACM. 2021. Т. 65. № 1. С. 99–106.

2. Hu B., Huang J., Liu Y., Tai Y. W., Tang C. K. Nerf-rpn: A general framework for object detection in nerfs // *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 2023. C. 23528–23538.
3. Xu C., Wu B., Hou J., Tsai S., Li R., Wang J., Zhan W., He Z., Vajda P., Keutzer K., Tomizuka M. Nerf-det: Learning geometry-aware volumetric representation for multi-view 3d object detection // *Proceedings of the IEEE/CVF International Conference on Computer Vision*. 2023. C. 23320–23330.
4. Lowe D. G. Object recognition from local scale-invariant features // *Proceedings of the seventh IEEE international conference on computer vision*. Ieee, 1999. T. 2. C. 1150–1157.
5. Rusu R. B., Cousins S. 3d is here: Point cloud library (pcl) // *2011 IEEE international conference on robotics and automation*. IEEE, 2011. P. 1–4.
6. Qi C. R., Su H., Mo K., Guibas L. J. Pointnet: Deep learning on point sets for 3d classification and segmentation // *Proceedings of the IEEE conference on computer vision and pattern recognition*. 2017. P. 652–660.
7. Wang Q., Wang Z., Genova K., Srinivasan P. P., Zhou H., Barron J. T., Martin-Brualla R., Snavely N., Funkhouser T. Ibrnet: Learning multi-view image-based rendering // *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. 2021. P. 4690–4699.
8. Varma M., Wang P., Chen X., Chen T., Venugopalan S., Wang Z. Is attention all that neRF needs? // *The Eleventh International Conference on Learning Representations*. 2022.
9. Cong W., Liang H., Wang P., Fan Z., Chen T., Varma M., Wang Y., Wang Z. Enhancing nerf akin to enhancing llms: Generalizable nerf transformer with mixture-of-view-experts // *Proceedings of the IEEE/CVF International Conference on Computer Vision*. 2023. P. 3193–3204.
10. Zhang K., Riegler G., Snavely N., Koltun V. Nerf++: Analyzing and improving neural radiance fields // *arXiv preprint arXiv:2010.07492*. 2020.
11. Ichnowski J., Avigal Y., Kerr J., Goldberg K. Dex-NeRF: Using a neural radiance field to grasp transparent objects // *arXiv preprint arXiv:2110.14217*. 2021.
12. Yen-Chen L., Florence P., Barron J. T., Lin T. Y., Rodriguez A., Isola P. Nerf-super-vision: Learning dense object descriptors from neural radiance fields // *2022 international conference on robotics and automation (ICRA)*. IEEE, 2022. P. 6496–6503.
13. Kundu A., Genova K., Yin X., Fathi A., Pantofaru C., Guibas L. J., Tagliasacchi A., Dellaert F., Funkhouser T. Panoptic neural fields: A semantic object-aware neural scene representation // *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 2022. P. 12871–12881.
14. Müller N., Simonelli A., Porzi L., Bulò S. R., Nießner M., Kotschieder P. Autorf: Learning 3d object radiance fields from single view observations // *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 2022. P. 3971–3980.
15. Ge Y., Behl H., Xu J., Gunasekar S., Joshi N., Song Y., Wang X., Itti L., Vineet V. Neural-sim: Learning to generate training data with nerf // *European Conference on Computer Vision*. Cham: Springer Nature Switzerland, 2022. P. 477–493.

16. Rukhovich D., Vorontsova A., Konushin A. Imvoxelnet: Image to voxels projection for monocular and multi-view general-purpose 3d object detection // Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision. 2022. P. 2397–2406.
17. He K., Zhang X., Ren S., Sun J. Deep residual learning for image recognition // Proceedings of the IEEE conference on computer vision and pattern recognition. – 2016. P. 770–778.
18. Radosavovic I., Kosaraju R. P., Girshick R., He K., Dollár P. Designing network design spaces // Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. 2020. P. 10428–10436.
19. Xie S., Girshick R., Dollár P., Tu Z., He K. Aggregated residual transformations for deep neural networks // Proceedings of the IEEE conference on computer vision and pattern recognition. 2017. P. 1492–1500.
20. Zhang H., Wu C., Zhang Z., Zhu Y., Lin H., Zhang Z., Sun Y., He T., Mueller J., Manmatha R., Li M., Smola A. Resnest: Split-attention networks // Proceedings of the IEEE/CVF conference on computer vision and pattern recognition. 2022. P. 2736–2746.
21. Sun J., Xu Y., Ding M., Yi H., Wang C., Wang J., Zhang L., Schwager M. NeRF-Loc: Transformer-based object localization within neural radiance fields // IEEE Robotics and Automation Letters. 2023.

L. SAQQOUR, I. KARABULATOVA, YU. GAPANYUK

Bauman Moscow State Technical University
1830@bmstu.ru, radogost2000@mail.ru, gapyu@bmstu.ru

EXPLORING THE SHORTCOMINGS OF OPEN TRANSLATION SERVICES FOR ARABIC AND RUSSIAN LANGUAGES

This article discusses the shortcomings of existing open translation services for the Arabic language, such as Google Translate, Yandex.Translate, DeepL. Examples of errors that services make during translation are shown. As an alternative, a hybrid system builds on Yandex.Translate is developed.

Keywords: *machine translation, neural machine translation, hybrid machine translation.*

Introduction

Machine translation is the process of translating texts from one natural language to another using a special software. Machine translation systems are available everywhere any time and almost for all the languages and that is an undeniable advantage. However, despite all the progress and successes

of machine translation, its simplicity and ease of use, regular users and specialists note existing errors in translations obtained using such systems. Nowadays translational software can only be used as an auxiliary tool for manual translation.

According to [1], there are four periods in the history of machine translation:

1. Rule-Based Machine Translation or RBMT (the main period of development until the early 1980s). RBMT systems rely on countless linguistic rules that cover semantic, morphological, and syntactic patterns of required languages, as well as built-in dictionaries for the languages pairs.
2. Example-Based Machine Translation or EBMT (the main period of development is from the early 1980s to the early 1990s). The knowledge-base of EBMT system is a bilingual corpus with parallel texts that is used during translation process.
3. Statistical Machine Translation or SMT (the main period of development is from the early 1990s to the year 2014). SMT systems were based on various statistical methods and machine learning methods (such as SVM, decision trees, boosting methods) and used vectorized bilingual corpus as knowledge-base.
4. Neural Machine Translation or NMT (the main period of development is from the year 2014 to present). They are built using neural networks and deep learning techniques and typically consider translation as a sequence prediction problem [2].

Currently, the main efforts in the development of NMT systems are related to Large Language Models (LLM). LLMs are used both for machine translation [3] and for machine translation quality evaluation [4]. The results show that in some cases, LLM models are inferior in translation quality to specialized NMT neural network architectures. To a large extent, this depends on the text corpus on which the model was trained.

Hybrid Machine Translation systems or HMT can also be considered as a separate class of software, which can combine individual elements of RBMT, EBMT, SMT, and NMT approaches. Also, these approaches can be used as components of an ensemble model, because SMT systems typically give better translations in short sentences, while long sentences are better analyzed with RBMT or NMT systems.

Currently, there are a large number of open machine translation services that work with varying degrees of accuracy for the Arabic language. Most of these services use the NMT or HMT approach.

The Existing Problems of Translation from Russian Language into Arabic Language

To conduct the experiments, the authors used the most well-known open machine translation services:

- **Google Translate**¹ is an HMT system developed by Google, that uses both SMT and NMT approaches. By analyzing vast amounts of text data, Google Translate continuously improves its translation quality and adapts to the nuances of different languages.
- **Yandex.Translate**² is also an HMT system developed by Yandex. Yandex.Translate uses the CatBoost technology to provide users with highly accurate and nuanced translations. When it comes to Russian language, Yandex.Translate gives better translations than Google Translate, and it ensures seamless and natural language conversion.
- **DeepL**³ use NMT approach, based on deep convolutional neural networks, which results in a more natural-sounding translations across various languages. Typically, both the grammar and the sentence structure of the produced translations are correct.

Most problems arise because all translations are first done into English and then into another language. It is shown in the first row of Table 1.

SMT systems give high-quality translations when the direction of the translations is very common. The problems start to appear when the direction is low-coverage, like Russian-Arabic, or when the subject area is not common like medical topics for example, because there are much fewer parallel texts than in everyday and conversational topics. In such situations English is used as a mediator, and the translations are made from Russian, in our situation, to English, then from English to Arabic.

Another problem that arises from using the English language as an intermediary is that English does not have gender for adjectives, while adjectives have gender in both Arabic and Russian languages. It is shown in the second and third rows of Table 1.

Arabic has a peculiarity in terms of the order of the verb and the noun in a sentence. The verb comes first, and then the noun, which is the ideal case. There is nothing wrong with mentioning the noun first, but it is preferable to maintain this feature to give the language more fluency. It is shown in the fourth row of Table 1.

¹ <https://translate.google.com>

² <https://translate.yandex.ru>

³ <https://www.deepl.com>

Some texts used in creating translators are considered to infringe on women’s rights, as any text about cooking or the kitchen is translated in the second person feminine form. It is shown in the fifth row of Table 1.

Table 1

Examples of the translations’ errors

	<i>RU</i>	<i>AR (Yandex)</i>	<i>AR (Google)</i>	<i>AR (DeepL)</i>	<i>AR (True)</i>
1	штраф	بخير	بخير	العقوبة	غرامة
2	Я женат	أنا متزوج	أنا متزوج	أنا متزوج	أنا متزوج
3	Я замужем	أنا متزوج	أنا متزوج	أنا متزوج	أنا متزوجة
4	Учитель задает вопрос	المعلم يسأل سؤالاً	المعلم يسأل سؤالاً	يطرح المعلم سؤالاً	يسأل المعلم سؤالاً
5	Выложите пшено в кастрюлю	ضعي الدخن في قدر	ضعي الدخن في قدر	ضعي الدخن في قدر	ضع الدخن في قدر

The last problem noticed is that machine translation systems process the past tense verb in Arabic as if it were an imperative mood, knowing that at the beginning of the imperative mood there is the letter ʾ. It is shown in the Table 2.

Table 2

Example of the error translating imperative

<i>AR</i>	<i>RU (MT)</i>	<i>RU (True)</i>
اقرأ الدرس	Он прочитал урок	Прочтите урок

It was also noticed that DeepL is giving better results, but it adds the definition tool.

The Proposed Hybrid System

The system architecture is designed to meet the needs of users who require accurate and efficient translations between Arabic and Russian. The architecture includes three blocks, each of which performs a specific function. The first block

is responsible for translation from Arabic into Russian (activity diagram of this block is shown in Figure 3), the second block is responsible for translation from Russian into Arabic (activity diagram of this block is shown in Figure 2), and the third block is for determining the language in case the source and target languages are not specified. The following diagram in Figure 1 shows how the components of the system will work to achieve the goal of the system.

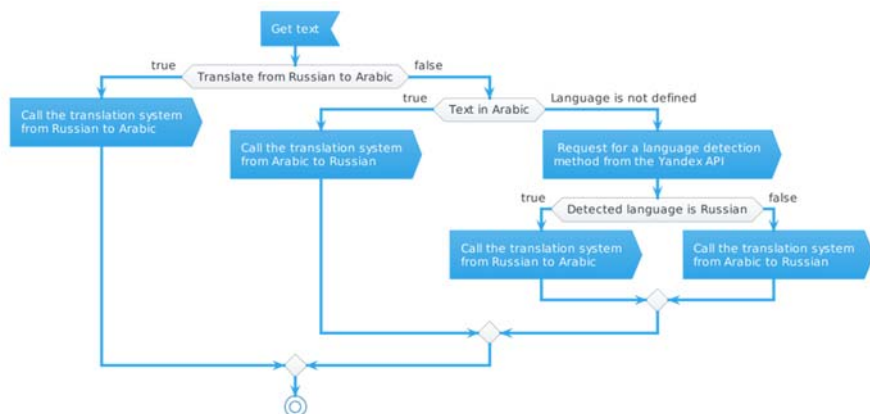


Fig.1. The structure of the developed system

To carry out these translations, the system uses the Yandex API, which provides a reliable and efficient translation service. The API is integrated into each block, allowing the system to receive translations quickly and accurately. However, to ensure high quality translations, the system also uses a rule-based translation system. This system uses a set of predefined rules to provide more accurate translations in specific situations where the Yandex API cannot provide a better translation.

Although rule-based translation systems can be complex to develop and expand, they offer significant advantages in terms of accuracy and quality. By combining the Yandex API with a rules-based translation system, the proposed architecture provides users with the best of both worlds - reliable and efficient translations, enhanced by more accurate translations in specific situations.

To further improve the accuracy of the system, new rules can be added as needed. This approach allows for continuous improvement and refinement of the translation process, ensuring that users receive the most accurate and reliable translations. Overall, this system architecture is designed to provide users with a seamless and efficient translation experience between Arabic and Russian.

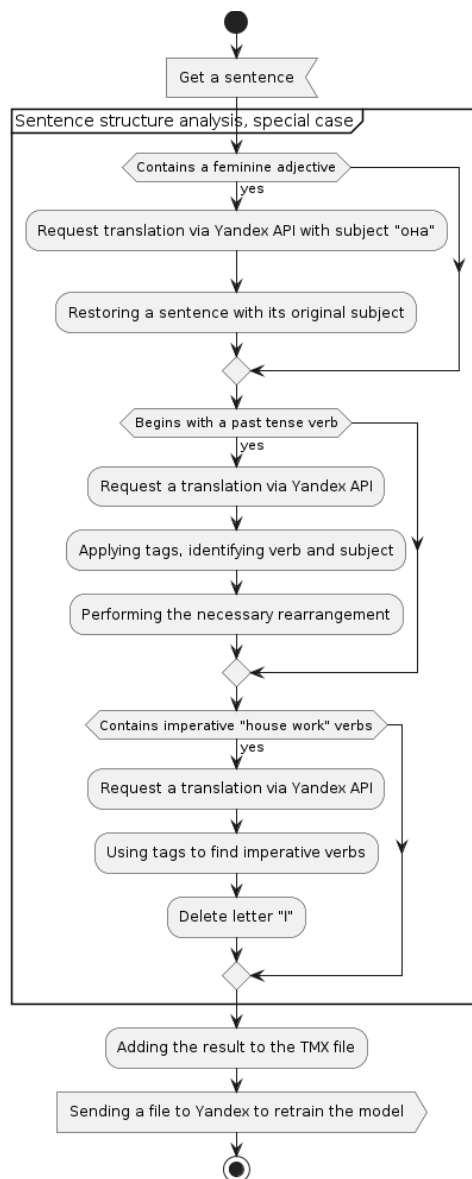


Fig.2. The activity diagram of the Russian to Arabic part

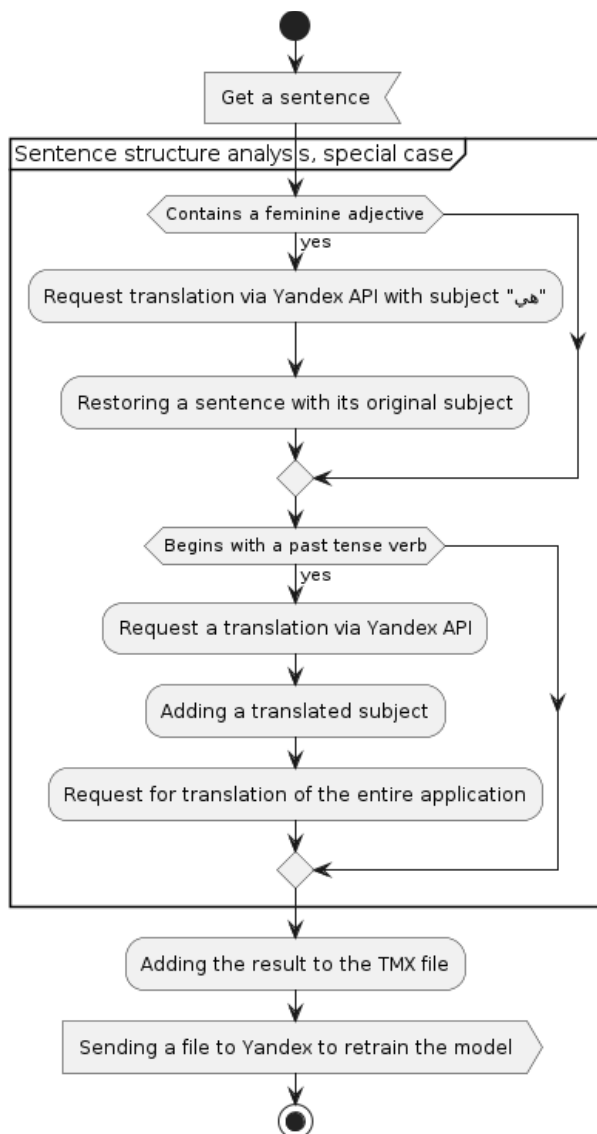


Fig. 3. The activity diagram of the Arabic to Russian part

The diagram shown in Figure 1 illustrates that the system uses a multi-step approach to ensure the highest level of accuracy. If the input is a single word, the system first checks its translation database, which contains entries for special cases where the Yandex API may not provide the correct translation due to its Russian/Arabic to English translation process with subsequent translation into Arabic/Russian, respectively. If a matching entry is found, the system returns the translation; otherwise, it requests a translation using the Yandex API.

For sentences with more than one word, the system uses the Yandex API as the main translation source. However, it also analyzes the words of the sentence and their part-of-speech tags to identify any special cases where the API may not provide the most accurate translation. In such cases, the system applies a set of predefined rules to improve translation accuracy. These rules will differ between the translation subsystems from Russian to Arabic and from Arabic to Russian due to the fact that translation from Russian to Arabic and vice versa is not always identical.

To further improve the accuracy of the system, it undergoes continuous training and refinement. The system compares translations obtained from the Yandex API and those obtained by applying the rules to evaluate their accuracy in sentences containing special cases and those that do not. This process allows the system to formulate more precise new rules, improve existing ones and thereby improve the overall quality of the translation.

Using a combination of a translation database, Yandex API and a rules-based translation engine, this architecture provides users with much more accurate and efficient Russian to Arabic translations.

A good practice when using the Yandex open API is to save all the results of these rules in TMX files for further use and submission to Yandex technical support for further improvement. This approach allows for continuous improvement and refinement of the translation process, ensuring that users receive the most accurate and reliable translations. On the other hand, such a multi-step approach may lead to increased system response time to the request, and, as a result, slower processing, but this approach will ensure the highest level of accuracy and quality of translations.

Conclusion

Statistical and Neural machine translation currently are the most common approaches.

Many modern open translation services for Arabic use English as an intermediate language, which cause to the problems with the quality of translation.

The hybrid translation system has shown promising results in bridging the gap between machine and human translation. However, to further enhance

the quality of translation, our new hypothesis suggests that incorporating ontologies and thesaurus into the system could provide a more accurate and contextually relevant output. This approach aims to leverage semantic knowledge and linguistic resources to refine the translation process and deliver more precise and nuanced results.

References

1. Sharma, Sonali and Diwakar, Manoj and Singh, Prabhishek and Singh, Vijendra and Kadry, Seifedine and Kim, Jungeun: Machine Translation Systems Based on Classical-Statistical-Deep-Learning Approaches, Electronics, (2023)
2. Zhixing Tan, Shuo Wang, Zonghan Yang, Gang Chen, Xuancheng Huang, Maosong Sun, Yang Liu: Neural Machine Translation: {A} Review of Methods, Resources, and Tools, CoRR abs/2012.15515, <https://arxiv.org/abs/2012.15515> (2012).
3. Wenhao Zhu and Hongyi Liu and Qingxiu Dong and Jingjing Xu and Shujian Huang and Lingpeng Kong and Jiajun Chen and Lei Li: Multilingual Machine Translation with Large Language Models: Empirical Results and Analysis. (2023).
4. Xu Huang and Zhirui Zhang and Xiang Geng and Yichao Du and Jiajun Chen and Shujian Huang: Lost in the Source Language: How Large Language Models Evaluate the Quality of Machine Translation. (2024)

НЕЙРОИНФОРМАТИКА И ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ

А.Р. ИСХАКОВ, М.Р. БОГДАНОВ

Башкирский государственный педагогический университет им. М. Акмуллы,
Россия, Уфа
Уфимский университет науки и технологий, Россия, Уфа
intellab@mail.ru

ПАРАМЕТРИЧЕСКАЯ ИДЕНТИФИКАЦИЯ СИСТЕМ МАШИННОГО ЗРЕНИЯ КАК ЗАДАЧА ОБУЧЕНИЯ

Статья описывает элементы теории модифицированных дескриптивных алгебр изображений. Математическая модель системы машинного зрения получается в ходе структурного синтеза. Параметрическая идентификация формулируется как задача нелинейной оптимизации. Математические модели систем машинного зрения в теории МДАИ обладают гетерогенной архитектурой с малым числом вариационных параметров. Задача оптимизации математической модели системы машинного зрения аналогична обучению нейронной сети.

Ключевые слова: *система машинного зрения, математическая модель, модифицированные дескриптивные алгебры изображений, структурный синтез, параметрическая идентификация, измерение площади.*

Введение

Теория модифицированных дескриптивных алгебр изображений (МДАИ) [1–4] является одной из направлений исследований в области математического моделирования систем машинного зрения, как и теория дескриптивных алгебр изображений (ДАИ) И. Б. Гуревича и дескриптивных алгебр изображений с одним кольцом (ДАИ1К) В. В. Яшиной. Основные результаты исследований по темам ДАИ и ДАИ1К представлены в цикле работ [5–11]. Отличие теории МДАИ от теории ДАИ и ДАИ1К заключается в использовании структурного синтеза и параметрической идентификации математических моделей систем машинного зрения методами численной оптимизации [1–4].

Алгебраический подход используется в ДАИ и ДАИ1К как основной способ решения задач с применением методологии этой теории [5–11]. МДАИ тоже допускает применение алгебраических методов для описания процедурных и параметрических преобразований, применяемых в разработке моделей систем машинного зрения. В МДАИ введены следующие

понятия, относящиеся к алгебраическому подходу: пространство состояний изображений (ПСИ) и воронка со срезами; комбинаторная оценка числа состояний изображения в ПСИ и формула вероятности попадания в воронку при измерении; универсальные алгебры как алгебраические модели срезов воронки. Процедурные преобразования МДАИ, формально описывающие методы цифровой обработки изображений, замкнуты в ПСИ. Алгебраический подход направлен на разработку моделей систем машинного зрения через структурный синтез с применением алгебраических элементов и операций над ними. В теорию МДАИ введено понятие модели сцены наблюдения в двух состояниях – модели априори и апостериори. Модель сцены априори состоит из базовых и производных элементов, формирующих общее описание объекта наблюдения без уточнения значений характеристик объекта наблюдения. В результате анализа изображения системой машинного зрения первоначальная модель (модель априори) превращается в модель объекта апостериори [1–4].

В статьях [12, 13] предлагается системно-объектный детерминантный анализ, заключающийся в построении партитивной классификации с применением формально-семантической нормативной системы. В работе [14] авторы решают задачу формализации структурного синтеза технических систем. Результатом этих исследований стала квазиаксиоматическая теория, формализующая процедуры порождения смысла для естественно-языкового описания процесса создания нового технического решения. Основным инструментом создания описания процедуры синтеза технической системы стала порождающая грамматика над нечеткими структурами. Методология МДАИ допускает применение данного инструмента в структурном синтезе математических моделей систем машинного зрения. В работе [15] рассматривается проблема декомпозиции образов, что приводит к одной из форм лингвистических моделей.

Используя методологию теории МДАИ, можно разрабатывать математические модели подсистем обработки и анализа изображений систем машинного зрения (СМЗ). СМЗ, с точки зрения моделирования, представляет совокупность функций измерения характеристик и признаков различных модальностей у наблюдаемого объекта на изображении. Математические модели функций СМЗ могут быть получены в ходе последовательного решения следующих задач [1–4]:

1. структурного синтеза моделируемой функции СМЗ с помощью процедурных и параметрических преобразований в форме дескриптивных алгебраических схем преобразований изображений (ДАСПИ) в виде Т-ДАСПИ или Р-ДАСПИ;

2. параметрической идентификации суперпозиции Т-ДАСПИ и Р-ДАСПИ, представляющей математическую модель функции СМЗ.

Для решения задачи структурного синтеза используются методы цифровой обработки и анализа изображений в терминах теории МДАИ. Т-ДАСПИ представляет последовательность процедурных преобразований, направленных на преобразование начальных изображений. Р-ДАСПИ используются для анализа преобразованных изображений, и представляет совокупность взаимосвязанных параметрических преобразований. Таким образом, Т-ДАСПИ – модель подсистемы преобразования изображения СМЗ, а Р-ДАСПИ – модель ее подсистемы для анализа изображений. Определения этих понятий в методологии МДАИ приведено ниже по тексту статьи. Совокупность Т-ДАСПИ и Р-ДАСПИ является математической моделью СМЗ, получаемой в ходе ее структурного синтеза. В составе Т-ДАСПИ и Р-ДАСПИ используются методы обработки и анализа изображений, которые имеют свои собственные вариационные параметры. В данной работе приведены необходимые теоретические конструкции методологии МДАИ для формального описания модели СМЗ, строгая постановка задачи идентификации данной модели и генетический алгоритм для ее решения [1–4].

Теоретические основы методологии МДАИ

Пусть в одиночном изображении или в совокупности изображений наблюдается один объект. Такие изображения называются *сценой наблюдения* или просто *сценой*. Каждое изображение может быть задано в виде совокупности его *реализаций* $I \in \{I_{bin}, I_{gray}, I_{color}\}$ в соответствии с определением 1 в формате бинарного, полутонового или цветного изображения.

Определение 1. Реализациями изображения I называются конечно-мерные матрицы размерности $n \times m$ [1–4]:

1. $I_{bin} \stackrel{def}{=} \|x_{ij}\|, \quad x_{ij} \in \{0, 1\}$
 2. $I_{gray} \stackrel{def}{=} \|x_{ij}\|, \quad x_{ij} \in \{0, \dots, 255\}$
 3. $I_{color} \stackrel{def}{=} \|x_{ij}\|, \quad x_{ij} = \langle r_{ij}, g_{ij}, b_{ij} \rangle, \quad r_{ij}, g_{ij}, b_{ij} \in \{0, \dots, 255\}$
- $i = \overline{1, n}, \quad j = \overline{1, m},$

где $x_{ij}, r_{ij}, g_{ij}, b_{ij}$ – яркости пикселей, i, j – индексы пикселей в матрице реализации изображения.

Системы машинного зрения решают задачи преобразования и анализа изображения с использованием операций [1–4]:

1. процедурные преобразования (2) позволяют изменить содержание изображения (1);
2. параметрические преобразования (3) позволяют вычислить признаки на изображении (1).

Определение 2. Процедурным преобразованием $O_T(\cdot, \bar{p})$ изображения I_f , $f \in \{bin, gray, color\}$ с вектором вариационных параметров $\bar{p} = \langle p_1, p_2, \dots, p_k \rangle$ называется отображение

$$O_T(\cdot, \bar{p}): I_f \rightarrow I_g, \quad (2)$$

где $f, g \in \{bin, gray, color\}$ – реализации изображений (1) [1–4].

Определение 3. Параметрическим преобразованием $O_p(\cdot, \bar{p})$ изображения I_f , $f \in \{bin, gray, color\}$ с вектором вариационных параметров $\bar{p} = \langle p_1, p_2, \dots, p_k \rangle$ называется отображение

$$O_p(\cdot, \bar{p}): I_f \rightarrow q, \quad (3)$$

где $f, g \in \{bin, gray, color\}$ – реализации изображений (1), q – значение выбранного признака объекта наблюдения на изображении [1–4].

Необходимо разработать функцию измерения детерминированного количественного признака p объекта наблюдения на изображении (1) в виде суперпозиции математических моделей первого и второго уровней СМЗ. Математические модели первого и второго уровней СМЗ по отдельности представляют последовательности процедурных (2) и параметрических (3) преобразований, называемых дескриптивными алгебраическими схемами преобразования изображений (Т-ДАСПИ и Р-ДАСПИ соответственно). Т-ДАСПИ представляет суперпозицию («вложенную» друг в друга последовательность) процедурных преобразований, каждое из которых преобразовывает одно изображение в другое изображение согласно определению 4 [1–4].

Определение 4. Т-ДАСПИ называется суперпозиция процедурных преобразований (4):

$$\begin{aligned} \mathfrak{R}_T(I_{color}, \langle \alpha_1, \alpha_2, \dots, \alpha_n \rangle) &= \times \dots \\ \dots \times &= O_T^n \left(O_T^{n-1} \left(\mathfrak{R}_T(I_{color}, \langle \alpha_1, \dots, \alpha_{n-2} \rangle), \alpha_{n-1} \right), \alpha_n \right), \end{aligned} \quad (4)$$

где $O_T^n(I_{bin}, \alpha_n)$ и $O_T^1(I_{color}, \alpha_1)$ – последнее и первое процедурные преобразования изображений I_{bin} и I_{color} с вариационными параметрами α_n и α_1 соответственно в суперпозиции (4) [1–4].

Определение 5. Р-ДАСПИ называется совокупность процедурных преобразований (10):

$$p = \mathfrak{R}_p \left(I_{bin}, \langle \beta_1, \beta_2, \dots, \beta_m \rangle \right) = \bigcirc_{i=1}^m O_p^i \left(I_{bin}, \beta_i \right), \quad (5)$$

где $O_p^i \left(I_{bin}, \beta_i \right)$ – параметрические преобразования, представляющие методы измерения выбранного признака на бинарном изображении I_{bin} , с вариационными параметрами β_i ; p – измеряемый признак [1–4].

Определение 6. Математической моделью системы машинного зрения называется суперпозиция Т-ДАСПИ и Р-ДАСПИ согласно (6):

$$p = \mathfrak{R}_p \left(\mathfrak{R}_T \left(I_{color}, \langle \alpha_1, \alpha_2, \dots, \alpha_n \rangle \right), \langle \beta_1, \beta_2, \dots, \beta_m \rangle \right), \quad (6)$$

где I_{color} – начальное цветное изображение сцены наблюдения, представляемое в виде матрицы согласно п. 3 из определения 1; $I_{bin} = \mathfrak{R}_T \left(I_{color}, \langle \alpha_1, \alpha_2, \dots, \alpha_n \rangle \right) = O_T^n \left(O_T^{n-1} \left(\dots O_T^1 \left(I_{color}, \alpha_1 \right), \alpha_{n-1} \right), \alpha_n \right)$ – последовательность из n процедурных преобразований с вариационными параметрами $\bar{\alpha} = \langle \alpha_1, \alpha_2, \dots, \alpha_n \rangle$; $p = \mathfrak{R}_p \left(I_{bin}, \langle \beta_1, \beta_2, \dots, \beta_m \rangle \right)$ – совокупность из m параметрических преобразований с вариационными параметрами $\bar{\beta} = \langle \beta_1, \beta_2, \dots, \beta_m \rangle$ [1–4].

Задача идентификации математической модели системы машинного зрения как задача обучения

Параметрическая идентификация математической модели СМЗ предполагает вычисление всех допустимых наборов значений вариационных параметров $\bar{\alpha} = \langle \alpha_1, \alpha_2, \dots, \alpha_n \rangle$ и $\bar{\beta} = \langle \beta_1, \beta_2, \dots, \beta_m \rangle$ функции измерения признака p для СМЗ. Функция измерения признака p является математической моделью СМЗ, если решается задача измерения только одного признака. Если же перед СМЗ стоит задача измерения комплекса признаков $p_1, p_2, \dots, p_k$, то под математической моделью СМЗ нужно понимать совокупность k функций для измерения этих признаков. Параметрическая идентификация формулируется как задача численной нелинейной оптимизации (7) [1–4].

$$F(\bar{\alpha}, \bar{\beta}) = \left| p^* - \mathfrak{R}_p \left(\mathfrak{R}_T \left(I_{color}, \langle \alpha_1, \alpha_2, \dots, \alpha_n \rangle, \langle \beta_1, \beta_2, \dots, \beta_m \rangle \right) \right) \right| \xrightarrow{G} \min$$

$$G: \begin{cases} \alpha_i^{\min} \leq \alpha_i \leq \alpha_i^{\max}, & i = \overline{1, n} \\ \beta_j^{\min} \leq \beta_j \leq \beta_j^{\max}, & j = \overline{1, m} \end{cases}, \quad (7)$$

где $F(\bar{\alpha}, \bar{\beta})$ – целевая функция с вариационными параметрами $\bar{\alpha} = \langle \alpha_1, \alpha_2, \dots, \alpha_n \rangle$ и $\bar{\beta} = \langle \beta_1, \beta_2, \dots, \beta_m \rangle$ в преобразованиях Т-ДАСПИ и Р-ДАСПИ соответственно; $\alpha_i^{\min} \leq \alpha_i \leq \alpha_i^{\max}$, $i = \overline{1, n}$ – ограничения на вариационные параметры α_i ; $\beta_j^{\min} \leq \beta_j \leq \beta_j^{\max}$, $j = \overline{1, m}$ – ограничения на вариационные параметры β_j преобразований Р-ДАСПИ [1–4].

Для численного решения задачи нелинейной оптимизации необходима предварительная оценка значения выбранного признака $p = p^*$ для одного из объектов наблюдения на изображении. Множество допустимых наборов значений вариационных параметров $\{\langle \alpha_1^*, \alpha_2^*, \dots, \alpha_n^* \rangle\}$ и $\{\langle \beta_1^*, \beta_2^*, \dots, \beta_m^* \rangle\}$ целевой функции $F(\bar{\alpha}, \bar{\beta})$ позволяет измерить выбранный признак p с заданной точностью $\varepsilon > 0$, что означает выполнение условия (8) [1–4].

$$\left| p^* - \mathfrak{R}_p \left(\mathfrak{R}_T \left(I_{color}, \langle \alpha_1^*, \alpha_2^*, \dots, \alpha_n^* \rangle, \langle \beta_1^*, \beta_2^*, \dots, \beta_m^* \rangle \right) \right) \right| < \varepsilon. \quad (8)$$

Допустим, что фиксация изображения велась при неизвестных внешних факторах (например, освещение сцены наблюдения), которые несущественно изменяются в течение короткого срока времени. Следовательно, используя те же самые наборы значений вариационных параметров, можно с заданной точностью $\varepsilon > 0$ измерить признак p и у остальных объектов наблюдения на анализируемом изображении [1–4].

Предлагаемый подход был использован в работах [16–18]. Для решения в них поставленной задачи (7) с условием (8) были использованы различные модификации генетических алгоритмов. Ясно, что генетический алгоритм не может гарантировать нахождение глобального экстремума. Однако, как показывает практика, для выполнения условия (8) достаточно использование некоторых наборов значений вариационных параметров $\{\langle \alpha_1^*, \alpha_2^*, \dots, \alpha_n^* \rangle\}$ и $\{\langle \beta_1^*, \beta_2^*, \dots, \beta_m^* \rangle\}$ математической модели системы машинного зрения (6), при которых она достигает локального экстремума. Несмотря на это, вычисление оптимальных наборов значений вариационных параметров остается не решенной задачей в теории МДАИ.

Аналогичный подход используется в теории классических нейронных сетей. Классические нейронные сети способны решать задачи классификации, кластеризации, аппроксимации и т.п. Нейронные сети подлежат обучению на обучающем множестве, которое представляет собой настройку весовых коэффициентов. В обучении классических нейронных сетей (сетей неглубокого обучения) используется широкий спектр методов минимизации функционала ошибок 0-го, 1-го и 2-го порядков. К градиентным алгоритмам обучения относят: алгоритм градиентного спуска (алгоритм GD), алгоритм градиентного спуска с возмущением (алгоритм GDM), алгоритм градиентного спуска с выбором параметра скорости настройки (алгоритм GDA), алгоритм обратного распространения ошибки (алгоритм Rprop). Алгоритмы обучения на основе сопряженных градиентов основаны на методах: Флетчера–Ривса (алгоритм CGF), Полака–Ребейры (алгоритм CGP), Биеле–Пауэлла (алгоритм CGB), Молера (алгоритм SCG). Для обучения классических нейронных сетей используются также квазиньютоновы алгоритмы, основанные на методах: Бroyдена–Флетчера–Гольдфарба–Шанно (алгоритм BFGS), секущих плоскостей Баттити (алгоритм OSS), Левенберга–Марквардта (алгоритм LM) [19].

Другие методы численной оптимизации можно найти в работе [20], но проблема решения задачи (7) с условием (8) в теории МДАИ остается неразрешенной из-за специфичной природы методов обработки (процедурные преобразования) и анализа (параметрические преобразования) изображений.

Выводы

Системы машинного зрения разрабатываются с применением различных подходов и технологий. В большинстве случаев они направлены на решение задач распознавания объектов различной природы. В статье предлагается решение задачи высокоточного измерения площади объекта наблюдения с применением методологии МДАИ. Площадь является важной геометрической характеристикой любого объекта наблюдения. Измерения площадей объектов применяются во многих отраслях экономики. В нефтяной и химической промышленности дистанционно измеряют на спутниковых снимках и аэрофотоснимках площади разливов нефтепродуктов. В лесной промышленности возникает необходимость автоматического измерения объема перерабатываемых лесоматериалов. В геодезии измерение площадей является одной из ключевых задач. В военном деле на спутниковых и аэрофотоснимках необходимы высокоточные измерения площадей гражданских архитектурных сооружений, военных объектов и различной военной техники с целью их классификации и распознавания. В медицине из-

мерение площадей позволяют определять стадии онкологических и легочных заболеваний. Таким образом, высокоточное измерение площади является актуальной задачей в цифровой обработке и анализе изображений.

Методология МДАИ позволяет решать задачи высокоточных измерений различных детерминированных числовых признаков путем структурного синтеза и параметрической идентификации математических моделей систем машинного зрения. Структурный синтез математической модели проводится исследователем вручную. Параметрическая идентификация математической модели формулируется в виде задачи нелинейной оптимизации. Даются рекомендации по статистической оценке значений вариационных параметров модели при обработке множества изображений. Достоинствами такого подхода являются: математический подход к формальному описанию математических моделей, малое количество вариационных параметров моделей, постановка задачи параметрической идентификации в виде оптимизационной задачи и статистический подход к уточнению значений вариационных параметров. Недостатком предлагаемого подхода являются отсутствие методов и технологий автоматизации структурного синтеза таких моделей. Другая проблема заключается в применении генетических алгоритмов для решения задач параметрической идентификации, которые не гарантируют нахождение глобального минимума целевой функции и сходимостъ вычислительного процесса. Параметрическая идентификация основывается на допущении влияния одних и тех же природных факторов на объекты наблюдения. Это предполагает использование вычисленных значений вариационных параметров в измерении тех же признаков у других объектов наблюдения в данной сцене. Несмотря на имеющиеся недостатки, теория МДАИ является перспективным математическим направлением исследований.

Список литературы

1. Исхаков А.Р. Моделирование систем технического зрения в модифицированных дескриптивных алгебрах изображений / А.Р. Исхаков, Р.Ф. Маликов. Уфа : Башкирский государственный педагогический университет им. М. Акмуллы, 2015. – 159 с.
2. Исхаков А.Р. Методы математического моделирования обработки и анализа изображений в модифицированных дескриптивных алгебрах изображений: специальность 05.13.18 "Математическое моделирование, численные методы и комплексы программ": диссертация на соискание ученой степени кандидата физико-математических наук / Исхаков Алмаз Раилович. – Челябинск, 2017. – 164 с.
3. Iskhakov A.R. Structural Synthesis of the Computer Vision and Its Parametric Identification with Statistical Estimation of Variational Parameters / A.R. Iskhakov, V.R.

- Akbashev // Journal of Computational and Engineering Mathematics. 2023. V. 10, N 1. P. 56–63. – DOI 10.14529/jcem230106.
4. Iskhakov A.R. Mathematical Methods of Modeling of Image Processing and Analysis in the Modified Descriptive Algebras of Images / A.R. Iskhakov // Journal of Computational and Engineering Mathematics. 2016. V. 3, N 1. P. 3–9. DOI 10.14529/jcem160101. – EDN VRELJR.
 5. Gurevich I.B. Descriptive Image Analysis: Genesis and Current Trends / I.B. Gurevich, V.V. Yashina // Pattern Recognition and Image Analysis. Advances in Mathematical Theory and Applications. 2017. V. 27, N 4. P. 653–674. DOI 10.1134/S1054661817040071.
 6. Gurevich I.B. Descriptive Image Analysis: Part II. Descriptive Image Models / I.B. Gurevich, V.V. Yashina // Pattern Recognition and Image Analysis. Advances in Mathematical Theory and Applications. 2019. V. 29, N 4. P. 598–612. DOI 10.1134/S1054661819040035.
 7. Gurevich I.B. Algebraic Interpretation of Image Analysis Operations / I.B. Gurevich, V.V. Yashina // Pattern Recognition and Image Analysis. Advances in Mathematical Theory and Applications. 2019. V. 29, N 3. P. 389–403. DOI 10.1134/S105466181903009X.
 8. Gurevich I.B. Descriptive Image Analysis: Part III. Multilevel Model for Algorithms and Initial Data Combining in Pattern Recognition / I. B. Gurevich, V. V. Yashina // Pattern Recognition and Image Analysis. Advances in Mathematical Theory and Applications. – 2020. V. 30, N 3. P. 328–341. DOI 10.1134/S1054661820030086.
 9. Gurevich I.B. Descriptive Image Analysis: Part IV. Information Structure for Generating Descriptive Algorithmic Schemes for Image Recognition / I.B. Gurevich, V.V. Yashina // Pattern Recognition and Image Analysis. Advances in Mathematical Theory and Applications. 2020. V. 30, N 4. – P. 638–654. – DOI 10.1134/S1054661820040161.
 10. Gurevich I.B. Descriptive Models of Information Transformation Processes in Image Analysis / I. B. Gurevich, V. V. Yashina // Pattern Recognition and Image Analysis. Advances in Mathematical Theory and Applications. – 2021. – V. 31, N 3. – P. 402–420. – DOI 10.1134/S105466182103010X.
 11. Gurevich I.B. On Modeling Descriptive Image Analysis Procedures on a Specialized Turing Machine / I. B. Gurevich, V. V. Yashina // Pattern Recognition and Image Analysis. Advances in Mathematical Theory and Applications. – 2022. – V. 32, N 3. – P. 469–476. – DOI 10.1134/S1054661822030142.
 12. Маторин С.И. Системно-объектный детерминантный анализ. Построение генетической и партитивной классификаций предметной области / С.И. Маторин, В.В. Михелев // Искусственный интеллект и принятие решений. – 2022. – № 1. – С. 26–34. – DOI 10.14357/20718594220103.
 13. Маторин С.И. Системно-объектный детерминантный анализ. Партитивная классификация с помощью формально-семантической нормативной системы / С.И. Маторин, В.В. Михелев // Искусственный интеллект и принятие решений. – 2022. – № 2. – С. 17–26. – DOI 10.14357/20718594220202.

14. Заболевса-Зотова А.В. Формализация структурного синтеза технических систем на начальном этапе проектирования / А.В. Заболевса-Зотова, А.Б. Петровский // Искусственный интеллект и принятие решений. – 2022. – № 4. – С. 44–54. – DOI 10.14357/20718594220405.
15. Поляков О.М. О связи модели знаний и задачи распознавания образов / О.М. Поляков, С. Рудницкий // Искусственный интеллект и принятие решений. – 2023. – № 3. – С. 16–22. – DOI 10.14357/20718594230302.
16. Iskhakov A.R. Calculation of Aircraft Area on Satellite Images by Genetic Algorithm / A.R. Iskhakov, R.F. Malikov // Bulletin of the South Ural State University. Series: Mathematical Modelling, Programming and Computer Software. – 2016. – V. 9, N 4. – P. 148–154. – DOI 10.14529/mmp160414.
17. Исследование динамики площади озера Аслыкуль (Южное Предуралье) методом обработки изображений космических снимков на основе алгебраического подхода / Б.И. Кочуров, Р.Ф. Маликов, А.Р. Исхаков [и др.] // Теоретическая и прикладная экология. – 2021. – № 1. – С. 58–64. – DOI 10.25750/1995-4301-2021-1-058-064.
18. Исхаков А.Р. Метод структурного синтеза системы технического зрения для задачи измерения площади / А.Р. Исхаков // Прикладная информатика. – 2022. – Т. 17, № 6(102). – С. 122–134. – DOI 10.37791/2687-0649-2022-17-6-122-134.
19. Медведев В.С., Потемкин В.Г. Нейронные сети. MATLAB 6./Под общ. Ред. В.Г. Потемкина. – Москва : ДИАЛОГ-МИФИ, 2001. – 630 с.
20. Пантелеев А.В. Методы оптимизации в примерах и задачах : учеб. пособие / А.В. Пантелеев, Т.А. Летова. – 2-е изд., исправл. – Москва : Высш. шк., 2005. – 544 с.

В.Н. ШАЦ

Независимый исследователь, Санкт-Петербург
vlnash@mail.ru

УПОРЯДОЧИВАНИЕ ПРИЗНАКОВ КАК СПОСОБ СНИЖЕНИЯ ЭНТРОПИИ ОБУЧАЮЩЕЙ ВЫБОРКИ И ОСНОВА ПРОСТЕЙШИХ АЛГОРИТМОВ КЛАССИФИКАЦИИ

Рассматриваются вопросы применения упорядоченных признаков, основанные на предложенной автором понятие близости для объектов конечного множества. Значения признаков двух объектов одного и того же класса полагаются близкими при достаточно малой разности их величины. В центре вычислительных процедур здесь находятся не объекты в целом, а значения их признаков. Благодаря этому основная масса расчетов выполняется для одномерных функций, что приводит к качественному упрощению алгоритмов классификации. Каждая строка матрицы данных описывает совокупность значений признаков некоторого объекта. Однако при

упорядочивании признаков это свойство строк матрицы нарушаются, так как значения признаков располагаются в порядке неубывания. Поэтому для каждого упорядоченного признака исходная матрица данных трансформируется в матрицу, которая отличается от нее своим порядком строк, но состоит из тех же элементов. Показано, что при снижении энтропии обучающей выборки, вызванном переходом к упорядоченным признакам, обнаруживаются детерминированные зависимости между признаками объектов и их классом. При применении упорядоченных признаков можно найти большее число вариантов разбиения множества на кластеры, близкие по составу объектов соответствующим классам. Предложены принципиально новые и простейшие алгоритмы решения задачи классификации.

Ключевые слова: *концепция сходства, упорядочивание признаков, снижение энтропии выборки, алгоритмы классификации.*

Введение

Мы знаем, что органы чувств животного воспринимают информацию из окружающей среды и передают в мозг, где она сопоставляется с уже накопленной там аналогичной информацией, и суммарная информация из всех источников получает оценку в виде реакции (классификации). Тем не менее в аналогичных задачах машинного обучения, как правило, используется другая схема оценки информации. Здесь объект (аналог конкретной ситуации в окружающей среде) описан вектором признаков (любое его значение – это аналог реакции определенного органа чувств животного) и на каждом этапе вычислений рассматривается аналог объекта, как правило, в виде единой функции элементов вектора.

Нами была предложена схема обработки информации для решения задач машинного обучения, близкая к той, которую использует животное [1]. Согласно этой схеме вычисления ведутся для каждого признака в отдельности, и только на заключительном этапе (в «мозге») происходит объединение полученных результатов. Заключительный этап сводится к применению простейшей формулы полной вероятности. Таким образом, используется новая технология решения задач машинного обучения, где в центре вычислительных процедур находится не объект как элемент многомерного пространства признаков, а значение признака как элемента одного из векторов признаков объекта. Поэтому основная масса расчетов выполняется для одномерных функций, что приводит к качественному упрощению алгоритма.

Эта технология исходит из нового понятие сходства объектов, которое является одним из основополагающих в машинном обучении, поскольку

позволяет сопоставлять наборы данных, которыми описаны объекты обучающей выборки (ОВ), чтобы распознавать объекты разных классов и применить эти знания при разделении объектов тестовой выборки (ТВ) на классы [2]. Обычно мерой сходства двух объектов служит оценка их близости по расстоянию между ними в метрическом пространстве. Новая концепция рассматривает сходство объектов конечного множества, которая измеряется согласно понятию близости объектов одного класса. Значения некоторого признака объектов одного класса полагаются близкими при достаточно малой разности их величины.

Новая технология была реализована разными методами, но все они преобразовывали структуру матрицы данных путем упорядочивания признаков (УП). Их роль выполняли порядковые номера интервалов, на которые разбивались значения каждого признака. Среди методов отметим метод списков [3], согласно которому множество значений каждого признака объектов объединенной выборки разбивается на равное число интервалов (оно служит параметром), в пределах которых разница между значениями признака не учитывается. Списки номеров объектов ОВ одного и того же класса, попадающих в эти интервалы, были названы информационными гранулами [4].

В соответствии с новым понятием близости все объекты гранулы рассматриваются как ближайшие соседи по признаку. Поэтому приближенно считается, что гранула в целом и образующие ее объекты обладают одинаковыми статистическими характеристиками. Частота гранулы находится как частота сложного события появления определенной упорядоченной пары (номер интервала, класс объекта), которая вычисляется непосредственно по соотношению длин соответствующих подмножеств. Тогда частота любого значения признака в определенном классе равна частоте соответствующей гранулы, частота объекта в каждом классе равна средней частоте всех его признаков в каждом классе, а класс объекта соответствует максимуму этих частот.

Здесь проявился эффект систематизации знаний на основе идеи холизма [5] при многократном разделении целого на взаимосвязанные части. Так множества значений каждого признака разбивается на равные интервалы, для каждого признака формируются списки объектов, попадающих в интервалы, а эти списки разбиваются по классам, образуя информационные гранулы. Согласно новой концепции, при указанном выше допущении, объекты одного класса, являются «сходными». Фактически это означает, что объекты гранулы обладают близкими свойствами и возможно принадлежат одному и тому же классу.

Эти методы получили развитие в настоящей статье. Основное внимание обращено на изучение свойств УП и использование полученных результатов при классификации объектов. Установлено, что существуют свойства УП, которые отличаются для объектов разных классов. Кроме того, было показано, что благодаря снижению энтропии ОБ при УП некоторые стохастические зависимости трансформируются в детерминированные, что естественно качественно упрощает алгоритмы классификации.

Свойства упорядоченных признаков

Уточним терминологию [6]. Мы будем рассматривать множества, состоящие из объектов ОБ и ТВ, которые содержат полный объем информации, характеризующий объекты отдельных классов. Существование «глобальной» выборки, генеральной совокупности, не предполагается.

Пусть $X = \{x_{sk}\}$ – это множество значений признака k объекта номер $s \in \{1, \dots, l\}$ выборки длиной l , которое может иметь равные элементы. Это множество называется упорядоченным, если его элементы были перенумерованы и получили такие новые номера $(s) = (1), (2), \dots, (l)$, что $x_{(1)k} \leq x_{(2)k}, \dots, \leq x_{(l)k}$.

Применение УП опирается на результаты анализа диаграмм УП. Для каждого признака k эти диаграммы представляют собой отображение $\{\{s\}|k\} \rightarrow \{\{(s)|k\}$ в виде множества точек, расположенных на отрезках прямых $i = \text{const}$, соответствующих классам $i \in \{1, \dots, C\}$. Пример диаграммы для базы данных Iris [7] приведен на рис. 1, где $r = (s)/l_i$, l_i – длина класса i . Анализ диаграмм нескольких баз данных показал, что объекты каждого класса отличаются распределениями УП, и привел к предположению о возможности классификации объектов на основе учета особенностей таких распределений.

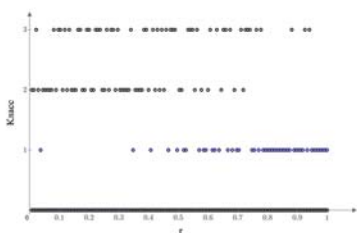


Рис. 1. Диаграмма признака $k = 1$ базы данных Iris

В ходе дальнейшего анализа была определена количественная характеристика взаимосвязи УП и классов множества. Ею является частота $f((s), i, k)$ появления между объектами класса одного, двух и т.д. значений признаков объектов любого другого класса. На рис. 2 представлены распределения частот $f((s), i, k)$ в точках оси $r = (s)/M$ для классов $i = 1, 2, 3$ баз данных Iris

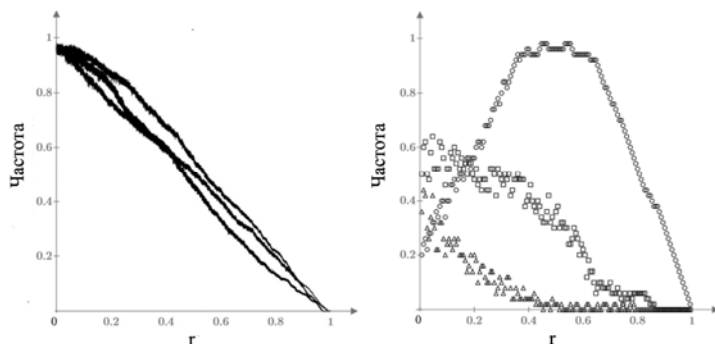


Рис. 2. Распределение частот $f(\tilde{s}, i, k)$ по длине объединенной выборки базы данных Iris (справа, $k=1$) и Letter Image Recognition (слева, $k=3$) при $i = 1, 2, 3$

(слева) и Letter Image Recognition (справа), вычисленные для объединённых выборок длиной M .

Рисунки показывают, что при каждом i точечные графики $f((s), i, k)$ могут быть аппроксимированы гладкой кривой с одним максимумом, на которую наложены случайные колебания, как правило, достаточно малой амплитуды. Наблюдается следующая закономерность: при любом k эти кривые для разных i отличаются почти на всем своем протяжении, а уровень их близости существенно отличаться для отдельных классов. Очевидно, что графики $f((s), i, k)$ являются еще одним свидетельством существования устойчивой взаимосвязи УП и классов и могут служить качестве идентификатора классов.

Использование УП влечет за собой изменение структуры данных. Объект с порядковым номером s в матрице данных $G = \|x_{sk}\|_{M \times N}$ описывает строка значений признаков $x_{s1}, x_{s2}, \dots, x_{sN}$. При упорядочивании любого признака объект и все его признаки изменяют свой номер s на (s) .

Для того, чтобы строки матрицы по-прежнему описывали объекты, все элементы строк должны изменить свою нумерацию. Поэтому при упорядочивании признака k мы получим новую матрицу данных G_k . Фактически все матрицы G_k и матрица G будут состоять из одних и тех же элементов, но отличаться порядком расположения строк.

Вопросы упорядочивания при классификации и кластеризации

УП позволяет найти инварианты классов, которые представляют собой особый вид кластеров, названных кластерами класса i , длина которых равна длине класса i .

Поиск таких инвариантов основан на следующих соображениях. Из постановки вопроса следует, что при любом k задано отображение $g: \{s\} \rightarrow \{i\}$ номера объекта s матрицы данных G на класс i . Класс объектов, представленных матрицей G_k , определяется зависимостью $i = g((s))$. Поэтому мы можем оценить уровень совпадения номеров объектов каждого класса матриц G и G_k , если найдем общее число объектов, для которых выполняется зависимость $g((s)) = g(s)$. Среднее число таких совпадений (по отношению к длине множества) названо индексом совпадения ψ_k признака k . Тогда мы получим N вариантов значений ψ для матрицы G .

Вычисления, выполненные для 10 баз данных Abalone, Adult Breast Cancer, Car evaluation, Glass, Haberman's Survival, Iris, Letter Image Recognition, Spect и Wine, показали [8]: для трех из 10 рассмотренных баз данных значение ψ превышает 0.9, 0.7 или 0.5, максимальное значение $\psi = 0.961$, только для одной базы значение $\psi \sim 0$. Из этих результатов следует, что для некоторых баз классы и кластеры i , соответствующие сопоставляемым матрицам, состоят из одних и тех же объектов, а в отдельных случаях списки их объектов могут практически полностью совпадать.

Как известно, классы – это подмножества объектов некоторого множества, которые отличаются своим свойствами от объектов других подмножеств. Поиск классов выполняют специалисты в определенной области знаний на основании анализа свойств объектов. Тем не менее, мы приближенно вычислили классы, не учитывая каких-либо сведений о свойствах объектов множества.

Этот феномен может быть объяснен, только если предположить существование взаимосвязи между классами и упорядоченными значениями его признаков. Но в таком случае объекты одного класса (иначе говоря, близкие по своим свойствам), являются также близкими по упорядоченным номерам значений своих признаков. В равной степени справедливо и обратное утверждение: упорядоченные номера объектов одного класса образуют непрерывные последовательности целых чисел.

Тогда разбив множество $\{s\}$ на C подмножеств, длина которых равна количеству объектов соответствующих классов, мы найдем списки объектов каждого класса, точнее кластера l_i класса i . Учитывая, что порядок расположения классов по длине выборки является произвольным, мы получим $C!$ вариантов разбиений на классы каждой матрицы G_k . В частности, для баз данных Iris и Wine, у которых $C = 3$, возможны $3! = 6$ вариантов расположения этих последовательностей по длине множества: (l_1, l_2, l_3) , (l_1, l_3, l_2) , (l_2, l_3, l_1) , (l_2, l_1, l_3) , (l_3, l_1, l_2) и (l_3, l_2, l_1) .

Всего существует $NC!$ вариантов приближенного разбиения множества, представленного матрицей данных G , на кластеры рассматриваемого типа.

Вычисления показали, то для многих вариантов баз данных значения $\psi > 0.70$. Таким образом, дополнительное упорядочивание в отношении уже упорядоченных номеров объектов УП (второй уровень упорядочивания признаков) позволяет приближенно найти состав объектов каждого класса.

Новые алгоритмы классификации

Продолжим изучение вариантов взаимосвязи между классами объектов и значениями их УП, которые являются отображениями всей совокупности данных задачи. Примем во внимание, что применение УП снижает неопределенность информации, содержащейся в ОБ, и ее энтропию. Будем разыскивать некоторую функцию, которая позволяет определить класс i для значения признака k , которое имеет упорядоченный номер объекта (s).

Пусть $\tilde{s} \in \{1, \dots, l_i\}$ – порядковый номер упорядоченных номеров (s) по признаку k в подмножестве объектов ОБ класса i . Очевидно, что при любых значениях (s), k и i мы получим единственное значение $\tilde{s}$. Поэтому существует функция φ , которая отображает множество $\{\tilde{s}\}$ на нормированное множество значений признаков Y для ОБ. Нормирование производится путем деления значений каждого признака на его максимальное значение в объединенной выборке.

Приращение функции $\Delta = \varphi(\tilde{s} + 1, k, i) - \varphi(\tilde{s}, k, i)$ будет достаточно малым, поскольку элементами множества $\{\tilde{s}\}$ являются ближайшие соседи по классу и признаку. Отметим, что распределение значений признаков по длине класса для порядковых номеров s имеет множество разрывов, в которых возникает скачок значений одного порядка с размахом значений признака. Поэтому функция $\varphi(s, k, i)$ не обладает содержательной информацией о взаимосвязи признаков и классов.

На рис. 3–4 представлены некоторые точечные графики $\varphi(\tilde{s}, k, i)$ для баз данных Iris и Wine. Они демонстрируют равномерный рост функции и отсутствие выбросов. Сам факт существования и свойства рассматриваемой функции в определенной степени соответствует идеям относительно роли порядка и хаоса в природе [9].

Следует отметить, что $\varphi(\tilde{s}, k, i)$ это детерминированная, а не случайная функция, которая устанавливает закономерности взаимосвязи значений признаков объектов и их классов. Для определения класса i объекта номер t , принадлежащего ТВ, по нормированному значению $u_k(t)$ его признака k нам нужно найти подмножество значений $\varphi(\tilde{s}, k, i)$, расстояние от которого до точки $(t, u_k(t))$ будет минимальным. Расчеты сводятся к вычислению функции:

$$i(t) = \arg \min_{1 \leq i \leq C} |\varphi(\tilde{s}, k, i) - u_k(t)| \quad (1)$$

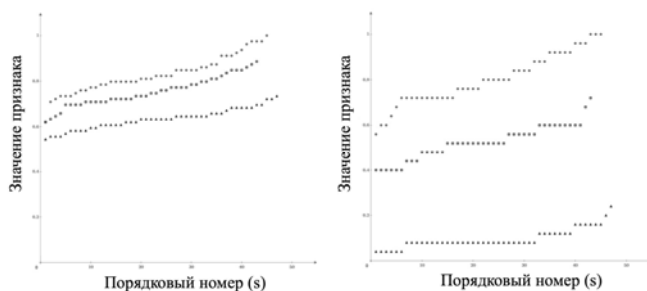


Рис.3. Точечные графики значений $\varphi(\tilde{s}, k, i)$ для базы данных Iris (слева $k=1$, справа $k=4$)

Практические вычисления по этой формуле требуют учета всех отношений между множеством упорядоченных значений (s) и множеством порядковых значений t . Для учета связей между этими множествами при каждом k точечные графики значений признаков $\varphi(\tilde{s})_{ki}$ рассматриваются как аналогии диаграммы Эйлера–Венна.

С этой целью разобьём прямыми линиями $y = a$ и $y = b$ диапазон значений $u_k(t)$ на несколько участков $a < u_k(t) \leq b$. Значения a и b должны удовлетворять требованиям: $a < \varphi(\tilde{s}, k, i) \leq b$ при единственном значении i . Они определяются с помощью вычисленных графиков. Тогда из формулы (1) получим, что $i(t) = i$.

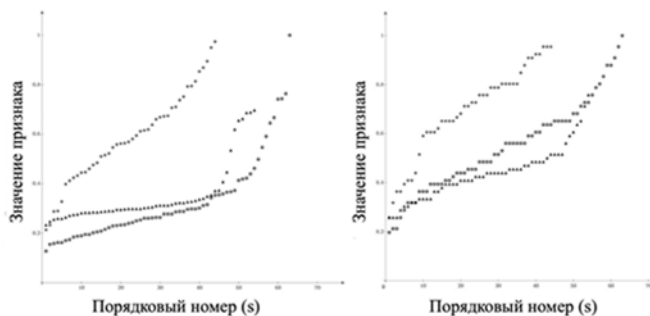


Рис. 4. Точечные графики значений $\varphi(\tilde{s}, k, i)$ для базы данных Wine (слева $k=2$, справа $k=8$).

Однако для многих признаков суммарная длина таких участков охватывает лишь небольшую часть множества $\{(s)\}$ или на длине этих участках расположены несколько значений $\varphi(\tilde{s}, k, i)$, относящихся к разным классам. В этих случаях нужно будет рассматривать кривые для большего числа признаков, и сопоставлять результаты для различных вариантов разбиения множества значений тестовой выборки на участки. Если окажется, что некоторый объект принадлежит интервалу, где расположены подмножества, например, для классов 2 и 3, а на другом участке для классов 1 и 2, то логично предположить, что объект имеет класс 2.

Зависимость (1) определяет класс любого объекта при произвольном значении k . Однако при некоторых k учет отмеченных особенностей взаимного расположения указанных кривых приводит к решению только для небольшого подмножества объектов тестовой выборки. Например, для базы Iris достаточно было разбить на указанные участки множество значений только одного признака, а для базы Wine, аналогичная процедура выполнялась уже для 11 признаков.

Существует также способ вычисления $i(t)$, применение которого не требует построения точечных графиков $\varphi(\tilde{s}, k, i)$, но он может давать менее точные результаты. Алгоритм вычисления состоит из нескольких этапов:

1. Найдем подмножества УП для всех признаков и классов ОБ: $\{\varphi(\tilde{s}, k, i) | k \in \{1, \dots, N\}, i \in \{1, \dots, C\}, (s) \in \{1, \dots, l_i\}\}$.
2. Для каждого значения признака $u_k(t)$ объекта тестовой выборки t вычислим множество значений $d = |u_k(t) - \varphi(\tilde{s})_{ki}|$, удовлетворяющих условию $d \leq \delta$, где δ – параметр. При всех k и i находим число случаев $q_{ki}(t, \delta)$, при которых это условие выполняется.
3. Вычислим суммарную по всем признакам величину $p_i(t, \delta) = \sum_k q_{ki}(t, \delta)$. Класс объекта t равен

$$i(t, \delta) = \arg \max_{1 \leq i \leq C} p_i(t, \delta) \quad (2)$$

Расчеты для баз данных Iris и Wine на основе формулы (1) показали, что количество ошибок классификации равно нулю. Из аналогичных расчетов по формуле (2) при 20 вариантах значений $\delta \in (0, 0.19)$ следует, что максимальное число ошибок соответствуют случаям $\delta \approx 0$. Они возникают в случае, если подмножества значений признаков ОБ и ТВ не пересекаются. С увеличением δ снижается вероятность такой ситуации, однако возрастает среднее число ближайших соседей i , соответственно, увеличивается риск ошибочной классификации.

Опыт применения обеих формул показал, что их не следует использовать при значительной несбалансированности классов.

Заключение

1. В статье рассмотрены вопросы о свойствах УП и их применении на основе новой концепции сходства и разработаны принципиально новые алгоритмы решения задач классификации.

2. Показано, что использование УП приводит к снижению энтропии обучающей выборки и позволяет выявить детерминированные зависимости между значениями признаков объектов и их классами. Эти зависимости служат основой простейших алгоритмов классификации.

3. Установлено, что благодаря упорядочиванию признаков можно найти большое число вариантов разбиения множества на кластеры, близких соответствующим классам по составу объектов.

4. Полученные результаты указывают на перспективность исследований по применению УП в машинном обучении и изучении свойств объектов множества данных.

Список литературы

1. Shats V. The classification of objects based on a model of perception // Advances in Neural Computation, Machine Learning, and Cognitive Research, Studies in Computational Intelligence, Springer, Moscow, P. 125–131. 2018. https://doi.org/10.1007/978-3-319-66604-4_1.
2. Luger G. Artificial intelligence: structures and strategies for complex problem solving. Addison-Wesley Springer, New York. 2016.
3. Shats V. Properties of the ordered feature values a classifier basis. Cybernetics and physics, 1(11), P. 25–29. 2022.
4. Yao J., Vasiliacos V., Pedrycz W. Granular Computing: Perspective and Challenges // IEEE Transactions on cybernetics. V. 43, N 6. P. 1977–1989. 2013.
5. Calabrese M. A., Amato, A., Lecce, V. D., and Piuri, V. (2010). Hierarchical-granularity holonic modelling // Journal of Ambient Intelligence and Humanized Computing. 1(3), P. 199–209. 2010.
6. David H., Nagaraja, H. Order Statistics. 3-rd ed. Wiley-Interscience, New Jersey. 2003.
7. Asuncion A., Newman, D.J. UCI Machine Learning Repository. Irvine. University of California, Irvine. 2007.
8. Shats V.N. Principle splitting of finite set in classification problem // Proceeding XXV International Conference Neuroinformatics, Moscow, P. 262–270. 2023.
9. Prigogine I., Stengers I. Order Out of Chaos: Man's New Dialogue with Nature. Flamingo Edition, London. Solso R. Cognitive psychology. 6-th ed. Allyn and Bacon. Boston. London. Toronto. P. 589. 2006.

НЕЙРОННЫЕ СЕТИ И КОГНИТИВНЫЕ НАУКИ. АДАПТИВНОЕ ПОВЕДЕНИЕ И ЭВОЛЮЦИОННОЕ МОДЕЛИРОВАНИЕ

В.Е. КУРЬЯН

ООО «АиС», ГОСНИИАС, Москва
kuryan@phystech.edu

ПОСТРОЕНИЕ ГРАФА МОДЕЛИ МИРА БЕЗ УЧИТЕЛЯ

Отмечена принципиальная неоднозначность соответствия реального мира и его описания на естественном языке. Описан подход к автоматическому построению модели мира без учителя. Модель мира строится в виде графа. Продемонстрирована зависимость структуры графа от выбора языка описания и сходство моделей, соответствующих различным языкам. Сделан вывод о том, что сходство моделей обусловлено отражением объективных закономерностей реального мира.

Ключевые слова: *граф модели мира, аксиомы обучения.*

Введение

Построение систем искусственного интеллекта подразумевает решение ряда разнообразных задач. Наиболее сложными из этого ряда задач являются автоматический перевод с одного естественного языка на другой и диалог на естественном языке. В последние годы произошел значительный прогресс в решении этих задач. Основным подходом при этом является использование парадигмы «последовательность в последовательность», при которой исходный текст рассматривается как упорядоченная последовательность слов (токенов). По этой входной последовательности строится выходная последовательность, которая считается переводом исходного текста или ответом системы. Реализация этой парадигмы при решении задачи автоматического перевода может осуществляться различными способами. Среди них стоит отметить использование синтаксического анализа текста и грамматических правил [1], статистический подход [2], использование нейронных сетей для получения перевода [3, 4], метод аналогий [5]. Наибольшего успеха удалось добиться системам на основе больших языковых моделей, среди которых наиболее известна GPT4 [6]. Принципиальная ограниченность систем на основе парадигмы «последовательность в последовательность» отмечалась ранее автором в работе [7]. Когда человек читает текст, у него в сознании возникает образ реального мира, с которым

связан текст и между текстом и возникшим образом устанавливается некоторое соответствие. В системах, реализующих парадигму «последовательность в последовательность» такой связи с реальным миром не образуется. Отсутствие такой связи является принципиальным ограничением этой парадигмы что нередко приводит к некорректным ответам.

При альтернативных подходах используется модель мира. В них система определяет какому фрагменту модели мира соответствует входная информация и на этой основе выдает ответ. Подход с использованием модели мира на основе онтологий описан в [8]. Недостатком этого метода является необходимость ручного построения онтологий. В работах автора [7, 9–13] развит подход, использующий модель мира в виде графа. При этом в процессе обучения системы модель мира строится автоматически. Построение модели мира (обучение системы) имитирует обучение человека. Построенная таким образом система обладает свойством объяснимости.

Идея автоматического перевода на основе модели мира на естественных языках состоит в том, что одна и та же ситуация во внешнем мире описывается на различных естественных языках. Эти описания называются переводами с одного языка на другой. Перевод выполняется по следующему алгоритму: по входному тексту, описывающему некоторую ситуацию в реальном мире на входном языке, выбирается часть модели мира, соответствующая этому тексту. В качестве перевода на выходной язык выбирается описание этой же части модели мира, но уже на выходном языке. Детально алгоритм обучения и перевода описан в предыдущих работах автора [7, 9–13]. В этих работах были сформулированы правила (аксиомы обучения) построения модели мира и проиллюстрировано их применение. Для обучения системе подавались описания ситуаций на входном и выходном языках. Этот метод соответствует обучению с учителем. В настоящей работе предложен метод обучения, соответствующий обучению без учителя.

Неоднозначность соответствия модели мира и реального мира

Отображение реального мира в его описание на каком-то языке может быть неоднозначным. Одна и та же ситуация в реальном мире может быть описана различными способами. Например, можно сказать: «Я хожу за продуктами на рынок» или «Я хожу за продуктами на базар». Оба эти описания ситуации эквивалентны.

Обратное отображение из описания в реальный мир тоже неоднозначно. В качестве примера неоднозначности можно привести фразу «Эти типы стали есть на складе». Приведенная последовательность слов может соответствовать совершенно различным ситуациям – утверждению о том, что

определенные марки металла имеются на складе либо утверждению о том, что какие-то люди начали принимать пищу в помещении склада.

Описания на различных языках могут соответствовать различным частям реального мира. Это означает, что существуют ситуации в реальном мире, которые хорошо описываются на одном языке, но не могут быть описаны на другом. Например, для описания состояния снега у народов севера существует довольно много различных слов и понятий, а в языках народов экваториального пояса понятия снег вообще нет. Модель мира, построенная с помощью языка народов экваториального пояса, не сможет описывать ситуации, связанные со снегом, и перевод описания сугроба с языка эвенков на языки экваториальных народов будет невозможен.

Построение модели мира без учителя

В предыдущих работах [7, 9–13] были сформулированы правила (аксиомы) построения модели мира в виде графа. При этом для построения модели мира подавались описания ситуации во внешнем мире на входном и выходном языках и отмечалось, что этот подход не зависит от выбора конкретных языков. Обучение на параллельных текстах соответствует обучению с учителем.

Для широко распространенных языков сравнительно просто найти необходимый набор обучающих параллельных текстов, но в тех случаях, когда требуется построить модель для редких языков, это часто невозможно. Например, для такой пары языков как язык саамских оленеводов и язык чукчей не получится найти достаточного набора параллельных текстов. Оба эти языка используют очень много слов для описания снегового покрова, поэтому попытка использования в качестве промежуточного языка одного из широко распространенных языков типа русского или английского для построения обучающего набора параллельных текстов не приведет к успеху.

Как было отмечено выше, правила построения графа модели мира не зависят от выбора входного и выходного языков, поэтому можно в качестве выходного языка выбрать входной язык, при этом проблема с подбором параллельных тестов автоматически исчезает. Для обучения (построения модели мира в виде графа) достаточно иметь набор текстов на одном языке. В этом случае будет построена модели мира, соответствующая одному выбранному языку. Такой подход реализует обучение без учителя.

Обучение без учителя можно применять для анализа ДНК, белков и прочих биомолекул (это линейные молекулы, состоящие из повторяющихся блоков: белки – из 20 аминокислот; а РНК и ДНК – из 4 нуклеотидов). Каждый тип аминокислоты или нуклеотида можно считать токеном или одной

из букв алфавита, при этом каждой биомолекуле соответствует упорядоченный набор букв – последовательность (текст на «белковом языке» или языке нуклеотидов). Аминокислотная последовательность – первичная структура белка; она определяет его структуру и функцию. Такие последовательности можно рассматривать как на тексты, описывающие функции белка.

Белки синтезирует рибосома, после «нанизывания» на мРНК начинающая синтезировать белковую цепочку. Прочитав три подряд идущих нуклеотида (кодон), рибосома прикрепляет новую аминокислоту к будущему белку. Таким образом она переводит текст с языка РНК на белковый. Правила перевода называются генетическим кодом, а процесс синтеза белка на основе РНК – трансляцией. Этот процесс аналогичен переводу с одного естественного языка на другой, поэтому подходы, развитые для естественных языков, применимы и при анализе биомолекул.

Индуктивное построение графа модели мира

Индуктивный подход подразумевает выделение общих закономерностей в обучающих текстах и проведение обобщений.

Основная идея автоматического построения модели мира состоит в том, что рассматриваются две ситуации во внешнем мире, имеющие нечто общее. Общей части этих двух ситуаций во внешнем мире соответствует общая часть в описании. Эти общие части выделяются и предполагается, что общей части описания на естественном языке соответствует общая часть графа модели мира, различающимся частям графа модели мира соответствуют различные части описаний на естественном языке. Выделение общих фрагментов во внешнем мире и их описаний на естественном языке и в виде фрагмента графа модели мира соответствует выделению общих понятий и закономерностей внешнего мира. Далее мы рассматриваем выделенные фрагменты как самостоятельные ситуации во внешнем мире и их описания на естественном языке. Это соответствует выделению обобщенных понятий и построению иерархической модели.

Считаем, что в языке имеется конечное число слов и словоформ, описание ситуации во внешнем мире представляется упорядоченным набором слов (токенов). Такое описание далее будем называть выражением. Любое выражение может состоять из упорядоченного набора других выражений более низкого уровня иерархии. Выражения удобно представлять в виде графа, в котором каждому выражению соответствует вершина, несущая информацию о порядке следования подвыражений. Такие вершины будем называть структурными вершинами, они несут информацию об общей структуре соответствующего естественного языка. Вершины, описывающие множество одинаково употребляемых выражений, будем называть

вершинами конкретизации. Выбор и подстановка определенного выражения из всего множества одинаково употребляемых выражений фиксирует конкретную ситуацию из множества похожих. Вершины более низкого уровня соответствуют элементам исходного выражения. Эти вершины соединены ребрами с вершиной исходного выражения. Детальное описание структуры графа модели мира приведено в работе [9]. В работах [7, 9–13] приведены правила построения графа в процессе обучения, при котором подаются пары предложений на входном и выходном языке.

Зависимость графа модели мира от языка описания

Построим графы модели мира для случая, когда при обучении входной и выходной языки совпадают. Выполним это построение для различных языков и сравним получившиеся графы моделей. Можно ожидать, что графы получатся различными для разных естественных языков. Выберем для построения графа часть реального мира, описываемого следующим набором предложений. 1. *Петя ходит в магазин.* 2. *Коля ходит в магазин.* 3. *Коля ходит в кинотеатр.* 4. *Петя ходит в зоопарк.* 5. *Коля ходит в зоопарк.* 6. *Петя ходит в кинотеатр.* 7. *Петя ходит в ближайший магазин.* 8. *Коля ходит в ближайший магазин.* 9. *Петя ходит в ближайший зоопарк.* 10. *Коля ходит в ближайший магазин.* 11. *Коля ходит в ближайший кинотеатр.* 12. *Коля ходит в магазин и кинотеатр.* 13. *Петя ходит в магазин и кинотеатр.* 14. *Коля ходит в магазин и зоопарк.* 15. *Коля ходит в магазин или зоопарк.* 16. *Коля ходит в магазин или кинотеатр.* 17. *Коля ходит в ближайший магазин и кинотеатр.* 18. *Коля ходит в ближайший магазин и зоопарк.* 19. *Коля ходит в ближайший магазин или кинотеатр.* Для построения графа модели мира используем аксиомы *stak2*, *stak31*, *stak32* [7, 13]. В результате получим граф модели мира, представленный на рис. 1.

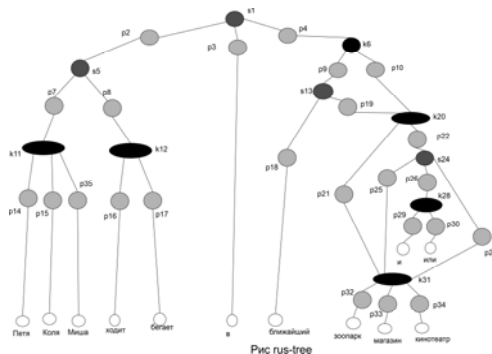


Рис. 1. Граф модели мира, соответствующий русскому языку

Эту же часть мира можно описать на английском языке следующим набором предложений. 1. *Petya goes to the shop.* 2. *Kolya goes to the shop.* 3. *Kolya goes to the cinema.* 4. *Petya goes to the zoo.* 5. *Kolya goes to the zoo.* 6. *Petya goes to the cinema.* 7. *Petya goes to the nearest shop.* 8. *Nick walks to the nearest shop.* 9. *Petya goes to the nearest zoo.* 10. *Nick walks to the nearest shop.* 11. *Kolya goes to the nearest cinema.* 12. *Nick walks into the shop and a cinema.* 13. *Petya goes to the shop and the cinema.* 14. *Nick walks into the shop and the zoo.* 15. *Nick goes to the shop or the zoo.* 16. *Nick walks into the shop or theater.* 17. *Kolya goes to the nearest shop and a cinema.* 18. *Nick walks to the nearest shop and the zoo.* 19. *Kolya goes to the nearest shop or movie theater.* Соответствующий ей граф изображен на рис. 2.

Аналогичное описание на немецком языке и граф приведен на рисунке 3. 1. *Petya geht in den laden.* 2. *Kohl geht in den laden.* 3. *Kohl geht ins KiN* 4. *Petya geht in den Zoo.* 5. *Kohl geht in den Zoo.* 6. *Petya geht ins KiN* 7. *Petya geht in den nächsten laden.* 8. *Kohl geht in den nächsten laden.* 9. *Petya geht in den nächsten Zoo.* 10. *Kohl geht in den nächsten laden.* 11. *Kohl geht ins nächste KiN* 12. *Kohl geht in den laden und ins KiN* 13. *Petya geht in einen laden und ein KiN* 14. *Kohl geht in den laden und den Zoo.* 15. *Kohl geht in einen Laden oder Zoo.* 16. *Kohl geht in ein Geschäft oder ein KiN* 17. *Kohl geht zum nächsten Geschäft und KiN* 18. *Kohl geht in den nächsten Laden und Zoo.* 19. *Kohl geht zum nächsten Geschäft oder KiN*

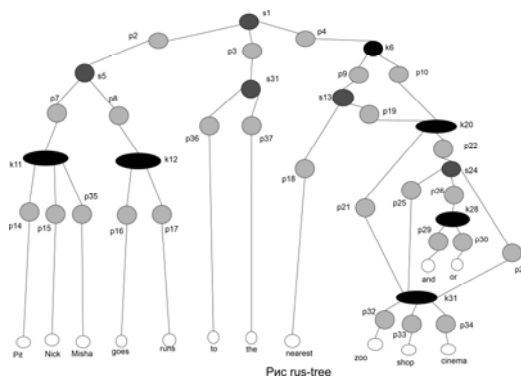


Рис. 2. Граф модели мира, соответствующий английскому языку

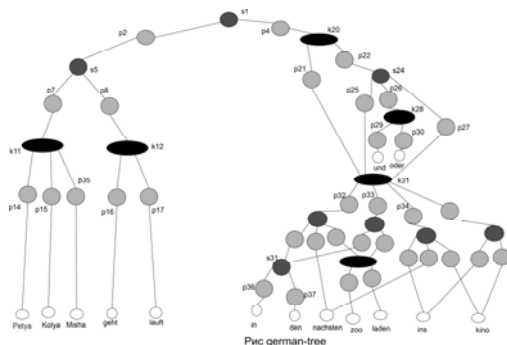


Рис. 3. Граф модели мира, соответствующий немецкому языку

Для французского языка получится результат, представленный на рис. 4. 1. *Petya va au magasin*. 2. *Kolya va au magasin*. 3. *Kolya va au cinéma*. 4. *Petya va au zoo*. 5. *Kolya va au zoo*. 6. *Kolya va au cinéma*. 7. *Petya va au magasin le plus proche*. 8. *Kolya va au magasin le plus proche*. 9. *Petya va au zoo le plus proche*. 10. *Kolya va au magasin le plus proche*. 11. *Kolya va au cinéma le plus proche*. 12. *Kolya va au magasin et au cinéma*. 13. *Petya va au magasin et au cinéma*. 14. *Kolya va au magasin et au zoo*. 15. *Kolya va au magasin ou au zoo*. 16. *Kolya va au magasin ou au cinéma*. 17. *Kolya va au magasin le plus proche et au cinéma*. 18. *Kolya va au magasin le plus proche et au zoo*. 19. *Kolya va au magasin ou au cinéma le plus proche*.

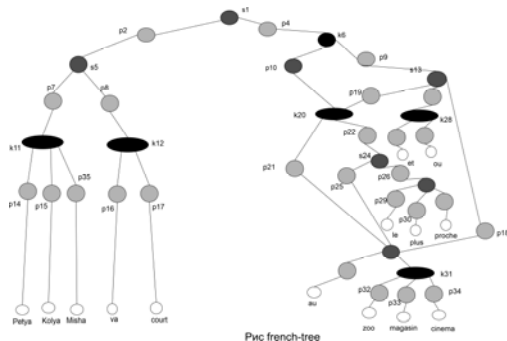


Рис. 4. Граф модели мира, соответствующий французскому языку

При описании на испанском языке получается следующий результат (рис. 5). 1. *Petya va a la tienda*. 2. *Kolya va a la tienda*. 3. *Kolya va al cine*. 4. *Petya va al parque zoológico*. 5. *Kolya va al parque zoológico*. 6. *Kolya va al cine*. 7. *Petya va a la tienda más cercana*. 8. *Kolya va a la tienda más cercana*.

9. Petya va al parque zoológico más cercaN 10. Kolya va a la tienda más cercana. 11. Kolya va al cine más cercaN 12. Kolya va a la tienda y al cine. 13. Petya va a la tienda y al cine. 14. Kolya va a la tienda y al parque zoológico. 15. Kolya va a la tienda o al parque zoológico. 16. Kolya va a la tienda o al cine. 17. Kolya va a la tienda y al cine más cercanos. 18. Kolya va a la tienda y al parque zoológico más cercanos. 19. Kolya va a la tienda o al cine más cercaN

При использовании китайского языка (используем способ записи буквами пиньинь) получим следующее (рисунок 6). 1. Petya qù shāngdiàn. 2. Kolya qù shāngdiàn. 3. Kolya qù diànyǐngyuàn. 4. Petya qù dòngwùyuán. 5. Kolya qù dòngwùyuàn. 6. Petya qù diànyǐngyuàn. 7. Petya qù zuìjìn de shāngdiàn. 8. Kolya qù zuìjìn de shāngdiàn. 9. Petya qù zuìjìn de dòngwùyuàn. 10. Kolya qù zuìjìn de shāngdiàn. 11. Kolya qù zuìjìn de diànyǐngyuàn. 12. Kolya qù shāngdiàn hé diànyǐngyuàn. 13. Petya qù shāngdiàn hé diànyǐngyuàn. 14. Kolya qù shāngdiàn hé dòngwùyuàn. 15. Kolya qù shāngdiàn huò dòngwùyuàn. 16. Kolya qù shāngdiàn huò diànyǐngyuàn. 17. Kolya qù zuìjìn de shāngdiàn hé diànyǐngyuàn. 18. Kolya dào zuìjìn de shāngdiàn hé dòngwùyuàn. 19. Kolya qù zuìjìn de shāngdiàn huò diànyǐngyuàn.

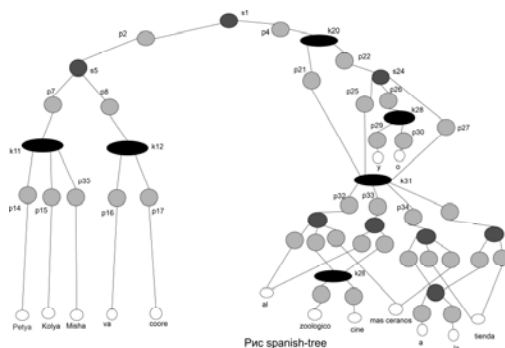


Рис. 5. Граф модели мира, соответствующий испанскому языку

Сравнение рис. 1–6 показывает, что все графы моделей мира получились различными. Модели даже имеют различные числа вершин. Это различие связано с различными языками описания внешнего мира. Тем не менее, все эти модели, несмотря на различия, имеют явное сходство. Это сходство связано с тем, что все они описывают один и тот же фрагмент внешнего мира и аксиомы построения графа модели мира учитывают и выделяют существенные объективные закономерности, не зависящие от конкретного языка описания внешнего мира. Решение задачи установления соответствия между моделями мира, соответствующими разным языкам, будет рассмотрено в последующих работах.

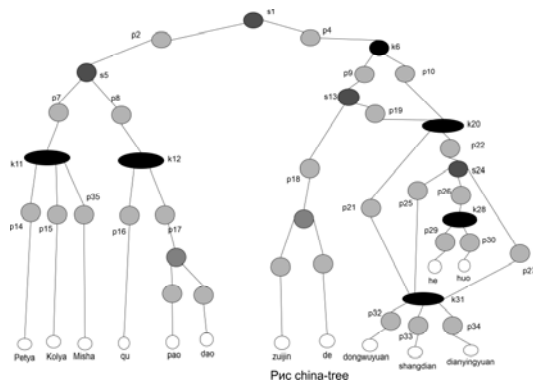


Рис. 6. Граф модели мира, соответствующий китайскому языку

Выводы

В работе показана возможность построения графа модели без учителя (когда на вход модели подается текст на одном произвольном естественном языке). Продемонстрирована зависимость графа модели мира от выбранного естественного языка и сходство получившихся графов, отражающее объективные закономерности внешнего мира. Отмечена возможность применения предлагаемого подхода при анализе ДНК, структуры и свойств белков.

Список литературы

1. Марчук Ю. Проблемы машинного перевода. Москва : Наука, 1982. 233 с.
2. Cho K. Learning Phrase Representations using RNN Encoder-Decoder for Statistical Machine Translation // 2014. arXiv:1406.1078v3
3. Ilya Sutskever, Oriol Vinyals, Quoc V. Le Sequence to Sequence Learning with Neural Networks // 2014, arXiv:1409.3215
4. Koehn T, Machacek D, Bojar O The reality of multi-lingual machine translation // 2022, arXiv:2202.12814
5. Makoto Nagao A framework of a mechanical translation between Japanese and English by analogy principle // Artificial and Human Intelligence. Elsevier Science Publishers. B.V / Ed. A. Elithorn and R. Banerji, 1984
6. GPT-4 Technical Report // 2023, arXiv:2303.08774v3
7. Курьян В.Е. Аксиомы построения графа модели мира // Нейроинформатика 2023. XXV Международная научно-техническая конференция. Сборник научных трудов. Москва : МИФИ. 2023 С. 78–87.
8. Анисимович К. [и др.]. Синтаксический и семантический Парсер, основанный на лингвистических технологиях abbyu compreN Abbyu. Москва : 2012

9. Курьян В.Е. Моделирование процесса обучения человека с помощью построения графа модели мира // Учет и статистика, N 4(56), 2019.
10. Курьян В.Е. Алгоритм автоматического перевода текстов с помощью дерева модели мира // Учет и статистика, N 1 (57), 2020.
11. Курьян В.Е. Автоматический перевод и граф модели мира // DOI 10.38006/907345-75-1.2020.305.309. Применение технологий виртуальной реальности и смежных информационных систем в междисциплинарных задачах fit-m 2020. сборник тезисов международной научной конференции. 2020. С. 305–309.
12. Курьян В.Е. Автоматический перевод на основе графа модели мира // Нейроинформатика 2021. XXIII Международная научно-техническая конференция. Сборник научных трудов. Москва : МИФИ. 2021 С. 29–38.
13. Курьян В.Е. Моделирование процесса обучения человека. Аксиоматический подход // Нейроинформатика 2022. XXIV Международная научно-техническая конференция. Сборник научных трудов. Москва : МФТИ. 2022 С. 260–269.

С.И. ФОКИН

НПЦ «Психосоматическая нормализация», г. Санкт-Петербург
fokin-S.I@yandex.ru

НОВЫЙ ПОДХОД К МАТЕМАТИЧЕСКОМУ ОПИСАНИЮ ПОТЕНЦИАЛОВ ДЕЙСТВИЯ, ОТЛИЧАЮЩИЙСЯ ОТ АЛГОРИТМА ХОДЖКИНА–ХАКСЛИ

На основании универсальной теории Ходжкина–Хаксли автором разработана математическая модель нейрона, связывающая интенсивность стимула с ответной частотой генерируемых им потенциалов действия и учитывающая влияние на этот процесс его физиологических характеристик. Вместо представления электропроводности мембраны непрерывной функцией на всём протяжении потенциала действия, как у Ходжкина–Хаксли, значения последней в предлагаемой модели задавались поочерёдно на каждом из его характерных участков, в течение которых электросопротивления мембраны считались постоянными, но разной величины, при некоторых вполне обоснованных допущениях.

Ключевые слова: модель Ходжкина–Хаксли, мембранный потенциал, зависимость частоты потенциалов действия от интенсивности стимула, калий-натриевый параметр мембраны, ионные каналы.

Введение

В модели Ходжкина–Хаксли потенциал действия (ПД) описывается единой аналитической кривой, получаемой из решения основного дифференциального уравнения, а входящие в него электропроводности мембраны

соответствующим ионам являются непрерывными эмпирическими функциями текущей разности потенциалов и времени, получаемыми, в свою очередь, из решения вспомогательных дифференциальных уравнений [1]. При этом постоянные коэффициенты, входящие в перечисленные уравнения, вычисляются аппроксимацией экспериментальных данных, полученных на реальной нервной ткани. Например, Ходжкин и Хаксли экспериментировали с гигантским аксоном кальмара [1], а с развитием необходимых технологий доступными для снятия электрических характеристик мембраны стали и нейроны головного мозга [2]. Однако ни подход к описанию ПД, ни соответствующие уравнения не изменились со времени их открытия, даже современные программы-симуляторы нейронов, например [3], используют ставший классическим математический аппарат Ходжкина–Хаксли (подключаемая подпрограмма hh).

На самом же деле электропроводность мембраны в течение потенциала действия изменяется не непрерывно, а дискретно, причём, величина «ступеньки» зависит от фазы ПД и значения стимула. Например, при «лавинообразном» открытии стимулзависимых ионных каналов мембраны её электропроводность может измениться на порядок в течение десятых долей миллисекунды [1, 2], из-за чего Ходжкин и Хаксли с сожалением отмечали, что для аналитического сглаживания столь резких колебаний лучше было бы использовать «пятую или шестую степень» при описании электропроводности, но всё же решили остановиться на четвёртой, посчитав, что «улучшение не стоит дополнительного усложнения» ([1], с. 509).

Но даже традиционные уравнения Ходжкина–Хаксли слишком громоздки для решения многих прикладных задач, например, для выражения частоты ПД в зависимости от интенсивности стимула, поэтому, видимо, никто этим и не занимается, по крайней мере, автору неизвестны подобные публикации и другие исследования в этой области. Полученные в эксперименте зависимости частоты ПД ν от интенсивности стимула S обычно описывают простой степенной функцией $\nu = cS^n$, где эмпирические коэффициенты c и n вычисляют при аппроксимации экспериментальных данных методом наименьших квадратов ([4], с. 369). Между тем функция $\nu(S)$ является «исходником» для формирования более общей зависимости «стимул-ощущение», над поиском формального выражения которой психофизики безуспешно «бьются» уже более полутора столетий [5], поэтому актуальность задачи переоценить трудно.

Описание Ходжкиным и Хаксли электропроводности мембраны непрерывной функцией, видимо, связано с тем, что в середине XX века ещё не умели измерять электропроводность единичных ионных каналов, сейчас же это хотя и технологически сложный, но обыденный процесс [2, 4]. Имея

данные по электропроводности всех ионных каналов, присутствующих в мембране, можем выразить её проводимость для любого типа ионов в виде не аналоговой, а цифровой зависимости с необходимой для точности расчётов дискретностью во времени вплоть до учёта открытия/закрытия каждого канала за соответствующие микропромежутки времени. В последнем случае дискретное описание электропроводности будет практически совпадать с аналоговым, поэтому предлагаемая модель не противопоставление модели Ходжкина–Хаксли, а лишь инвариантная форма математического представления одного и того же физического процесса, допускающая возможность упрощённого применения для решения практических задач и обязанная своим появлением современным достижениям в экспериментальной нейрофизиологии.

1. Математическая модель нейрона с дискретно меняющимся сопротивлением мембраны при характерных значениях потенциала

1.1. Постановка задачи и исходные данные.

Требуется найти математические выражения, описывающие изменение мембранного потенциала нейрона во времени в зависимости от интенсивности действующего на него адекватного стимула, а также частоты потенциалов действия от тех же параметров при сверхпороговой стимуляции. Тестовые расчёты проведём для нейрона простой сферической формы без отростков диаметром 20 (мкм) с толщиной мембраны 10 (нм), по площади которой равномерно распределены только калиевые и натриевые ионные каналы, подразделяющиеся на неселективные потенциалнезависимые «каналы утечки» обоих типов, потенциал- и стимулуправляемые натриевые каналы и потенциалзависимые калиевые каналы. Тогда общая электрическая схема нейрона Ходжкина–Хаксли ([1], с. 501) упростится за счёт удаления ветви утечки (рис. 1), учитывающей «токи утечки из-за хлоридов и других ионов» ввиду их отсутствия у рассматриваемого нейрона.

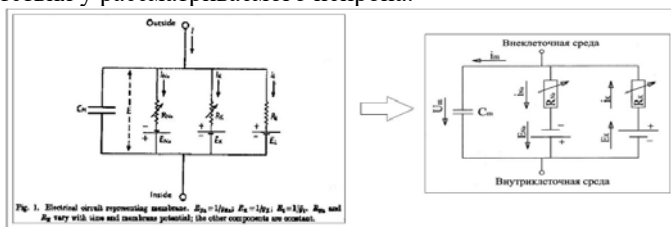


Рис. 1. Электрические схемы мембраны возбудимой клетки: слева – классическая схема Ходжкина – Хаксли; справа – модифицированная схема, используемая в настоящей работе, токи утечки в покое в которой учтены в ветвях натриевой и калиевой проводимости, а направления токов, ЭДС и напряжений расставлены в соответствии с правилами электротехники

Известны следующие параметры рассматриваемого нейрона: электро-сопротивления одиночных натриевых и калиевых ионных каналов всех присутствующих в мембране типов и количество каждого типа, электрическая ёмкость мембраны, пороговое значение стимула $S_{\text{пор}}$ и соответствующая ему величина мембранного потенциала (реобазы) $U_{\text{пор}}$, пиковое значение мембранного потенциала при генерации потенциалов действия $U_{\text{пик}}$, условное значение потенциала закрытия калиевых каналов $U_{\text{зкк}}$, калиевая E_K и натриевая E_{Na} ЭДС и количество открывающихся стимулируемых натриевых каналов в зависимости от величины стимула $N_{\text{стNa}}(S)$. Количественные значения вышеперечисленных параметров близки к реальным характеристикам нейронов [2, 4, 6] и равны (в порядке перечисления): $R_{1Na} = R_{1K} = 4 \cdot 10^{10}$ (Ом), $N_{\text{пок.Na}} = 34$ (шт), $N_{\text{пок.K}} = 476$ (шт), $N_{\text{пз.Na}} = 9359$ (шт), $N_{\text{пз.K}} = 6202$ (шт), $U_{\text{пор}} = -45$ (мВ), $U_{\text{пик}} = 35$ (мВ), $U_{\text{зкк}} = -65$ (мВ), $C_m = 62.68$ (пФ), $E_K = 90$ (мВ), $E_{Na} = -60$ (мВ). Сопротивления всех типов натриевых и калиевых одиночных каналов приняты равными с целью экономии объёма статьи на выводах соответствующих формул, но это не является обязательным условием предлагаемой модели. Вместо электропроводностей, как в модели Ходжкина–Хаксли, в предлагаемой модели использованы обратные им величины – электросопротивления, что непринципиально с точки зрения экспериментального их измерения, но позволит применять известные в электротехнике правила суммирования параллельно соединённых сопротивлений, в соответствии с которыми были вычислены необходимые для дальнейших расчётов суммарные сопротивления [7]: Na^+ - и K^+ -каналов мембраны в покое $R_{\text{пок.Na}} \approx 1.176 \cdot 10^9$ (Ом), $R_{\text{пок.K}} \approx 8.4 \cdot 10^7$ (Ом); потенциалзависимых натриевых каналов $R_{\text{пз.Na}} = 4.274 \cdot 10^6$ (Ом) и потенциалзависимых калиевых – $R_{\text{пз.K}} = 6.45 \cdot 10^6$ (Ом). Значения сопротивлений и количества ионных каналов не округлялись потому, что они должны быть согласованы как между собой, так и с соответствующими параметрами реальных нейронов [7]. Значение потенциала покоя нашего модельного нейрона может быть вычислено по имеющимся данным и составит $U_{\text{пок}} = -80$ (мВ) [7].

1.2. Решение. Для электрической схемы в правой части рис. 1 и обозначений на ней в соответствии с формулой ёмкостного элемента в электрической цепи и правилами электротехники можно записать следующие уравнения:

$$i_m = C_m \frac{dU_m}{dt}; \quad i_{Na} = (E_{Na} + U_m)/R_{Na}; \\ i_K = (E_K - U_m)/R_K; \quad i_m + i_{Na} - i_K = 0$$

или

$$C_m \frac{dU_m}{dt} + (E_{Na} + U_m)/R_{Na} - (E_K - U_m)/R_K = 0.$$

После несложных преобразований получим

$$\frac{dU_m}{dt} + U_m \frac{(R_{Na} + R_K)}{C_m R_{Na} R_K} = \frac{R_{Na} E_K - R_K E_{Na}}{C_m R_{Na} R_K}. \quad (1)$$

Емкость мембраны C_m , ЭДС E_K , E_{Na} в уравнении (1) являются постоянными величинами [1], а сопротивления тоже будут постоянными в течение определённых промежутков времени, например между открытиями ионных каналов. Общее решение такого типа дифференциальных уравнений с постоянными коэффициентами приведено в ([8], с. 35). Упростив решение указанного дифференциального уравнения, получим

$$U_m(t) = U_0 \exp\left(\frac{-t(R_{Na} + R_K)}{C_m R_{Na} R_K}\right) + \frac{R_{Na} E_K - R_K E_{Na}}{R_{Na} + R_K} (1 - \exp\left(\frac{-t(R_{Na} + R_K)}{C_m R_{Na} R_K}\right)). \quad (2)$$

Из этого уравнения можем выразить продолжительность любого промежутка времени, в течение которого его коэффициенты остаются постоянными:

$$\Delta t = t_1 - t_0 = -\frac{C_m R_{Na} R_K}{R_{Na} + R_K} \ln\left(\frac{R_K(U_1 + E_{Na}) + R_{Na}(U_1 - E_K)}{R_K(U_0 + E_{Na}) + R_{Na}(U_0 - E_K)}\right), \quad (3)$$

где напряжение на мембране нейрона U_0 соответствует моменту времени t_0 , а $U_1 - t_1$; значения сопротивлений соответствуют моменту времени t_1 .

Для упрощения расчётов вместо R_{Na} и R_K удобнее использовать их отношение: $X = R_K/R_{Na}$, названное в [7] *калий-натриевым параметром мембраны* или *переведённым стимулом*. Тогда после вынесения за скобки в вышеприведённой формуле R_{Na} и проведения сокращений получим

$$\Delta t = t_1 - t_0 = -\frac{C_m R_K}{X + 1} \ln\left(\frac{X(U_1 + E_{Na}) + U_1 - E_K}{X(U_0 + E_{Na}) + U_0 - E_K}\right). \quad (4)$$

Любой потенциал действия можно разделить на три характерных промежутка времени: подпороговый $\Delta t_{пш}$, пиковый $\Delta t_{пик}$ и реполяризации $\Delta t_{реп}$, граничными точками которых будут являться, соответственно, потенциал начала деполяризации под действием стимула ($U_{пок}$ или $U_{зкк}$) – пороговый потенциал или потенциал условного открытия потенциалзависимых натриевых каналов $U_{пор}$ – потенциал условного их закрытия (инактивации) и открытия потенциалзависимых калиевых каналов $U_{пик}$ – потенциал условного закрытия последних (инактивации) и активации потенциалзависимых натриевых каналов $U_{зкк}$. В связи с тем, что натриевая и калиевая проводимости в течение ПД имеют явно выраженные разделённые во времени мак-

симумы, расположенные между $U_{\text{пор}}-U_{\text{пик}}$ и $U_{\text{пик}}-U_{\text{зкк}}(U_{\text{пок}})$, соответственно, при которых «оппонирующая» проводимость практически равна нулю ([1], рис. 17, с. 530), то, приняв, что все потенциалзависимые натриевые каналы сразу открываются при $U_m = U_{\text{пор}}$ и закрываются при $U_m = U_{\text{пик}}$, а все потенциалзависимые калиевые каналы сразу открываются при $U_m = U_{\text{пик}}$ и закрываются при $U_m = U_{\text{зкк}}$, получим возможность приближённого описания потенциала действия $U_m(t)$ аналитическим уравнением (2), т.к. внутри промежутков времени $\Delta t_{\text{пп}}$, $\Delta t_{\text{пик}}$ и $\Delta t_{\text{реп}}$ все сопротивления в выражениях (1–4) остаются константами. Отметим, что максимумам натривой и калиевой проводимости мембраны будут соответствовать минимумы её электросопротивления.

Тогда, обозначив период потенциала действия как сумму составляющих его отрезков времени и подставив в уравнение (4) наши исходные данные, можем записать для элементарной частоты ПД ν :

$$\nu = \frac{1}{(\Delta t_{\text{пп}} + \Delta t_{\text{пик}} + \Delta t_{\text{реп}})} \quad (5)$$

$$\Delta t_{\text{пп}} = \frac{-C_m R_{\text{пок.К}}}{(X_{\text{ст}} + 1)} \ln \left(\frac{X_{\text{ст}}(U_{\text{пор}} + |E_{\text{Na}}|) + U_{\text{пор}} - |E_{\text{K}}|}{X_{\text{ст}}(U_{\text{зкк}} + |E_{\text{Na}}|) + U_{\text{зкк}} - |E_{\text{K}}|} \right) \quad (6)$$

$$\Delta t_{\text{пик}} = \frac{-C_m R_{\text{пок.К}}}{(X_{(\text{ст+пик})} + 1)} \ln \left(\frac{X_{(\text{ст+пик})}(U_{\text{пик}} + |E_{\text{Na}}|) + U_{\text{пик}} - |E_{\text{K}}|}{X_{(\text{ст+пик})}(U_{\text{пор}} + |E_{\text{Na}}|) + U_{\text{пор}} - |E_{\text{K}}|} \right) \quad (7)$$

$$\Delta t_{\text{реп}} = \frac{-C_m R_{(\text{пок+реп})\text{К}}}{(X_{\text{реп}} + 1)} \ln \left(\frac{X_{\text{реп}}(U_{\text{зкк}} + |E_{\text{Na}}|) + U_{\text{зкк}} - |E_{\text{K}}|}{X_{\text{реп}}(U_{\text{пик}} + |E_{\text{Na}}|) + U_{\text{пик}} - |E_{\text{K}}|} \right) \quad (8)$$

$X_{\text{ст}} = R_{\text{пок.К}}/R_{(\text{пок.+ст.})\text{Na}}(S)$, постоянен в течение $\Delta t_{\text{пп}}$;

$X_{(\text{ст+пик})} = X_{\text{ст}} + X_{\text{пик}} = X_{\text{ст}} + R_{\text{пок.К}}/R_{\text{пикNa}}$, постоянен в течение $\Delta t_{\text{пик}}$;

$X_{\text{реп}} = X_{\text{ст}} \cdot R_{(\text{пок+реп})\text{К}}/R_{\text{пок.К}}$, постоянен в течение $\Delta t_{\text{реп}}$,

где $R_{\text{пок.К}}$ – суммарное электросопротивление калиевых ионных каналов утечки в состоянии покоя; остаётся постоянным в течение всего времени потенциала действия $\Delta t_{\text{пл}} = (\Delta t_{\text{пп}} + \Delta t_{\text{пик}} + \Delta t_{\text{реп}})$, Ом;

$R_{(\text{пок.+реп.})\text{К}}$ – суммарное электросопротивление калиевых ионных каналов утечки и потенциалзависимых калиевых каналов, открывающихся при реполяризации; остаётся таковым в течение отрезка времени $\Delta t_{\text{реп}}$, Ом;

$R_{(\text{пок.+ст.})\text{Na}}(S)$ – суммарное электросопротивление натриевых ионных каналов утечки и стимулзависимых натриевых каналов, открывающихся при

деполяризации мембраны нейрона под действием стимула S ; остаётся таковым в течение всего времени ^{потенциала} действия $\Delta t_{\text{пд}} = (\Delta t_{\text{пп}} + \Delta t_{\text{пик}} + \Delta t_{\text{реп}})$ при постоянно действующем стимуле, Ом;

$R_{\text{пик.На}}$ – суммарное электросопротивление потенциалзависимых натриевых ионных каналов, открывающихся при деполяризации мембраны нейрона после достижения мембранным потенциалом порогового значения, $U_{\text{пор}}$, и закрывающихся при достижении мембранным потенциалом пикового значения $U_{\text{пик}}$; остаётся таковым в течение отрезка времени, $\Delta t_{\text{пик}}$, Ом;

X с индексами – **калий-натриевый параметр мембраны** или **переведённый стимул**, равный отношению суммарного электросопротивления одиночных калиевых ионных каналов мембраны нейрона к суммарному электросопротивлению одиночных натриевых каналов вне зависимости от механизмов их открытия/закрытия в текущий момент времени: $X(t) = R_K(t)/R_{Na}(t)$ [7]. $X(t)$ зависит от интенсивности стимула S , передаточной функции нейрона, его физиологических свойств и времени. Название «переведённый стимул» или «переведённое воздействие» для X выбрано в связи с тем, что в процессе трансдукции физическая интенсивность любого стимула «переводится» на универсальный язык нервной системы – в потенциалы действия первичных нейронов рецепторов, частота которых как раз и определяется значением $X_{\text{ст}}$. Сама же величина $X_{\text{ст}}$ формируется под действием реального физического стимула, являясь, так сказать, первой частью его «перевода» на язык организма. Таким образом, «калий-натриевый параметр мембраны» – ещё одно название $X_{\text{ст}}$ – является неким универсальным стимулом, от которого зависит частота ПД любого нейрона.

Натриевая и калиевая ЭДС в выражениях (6–8) взята по модулю потому, что в соответствии с правилами электротехники при записи балансовых уравнений «все ЭДС уже считаются положительными» ([9], с. 252). Со времени открытия электричества в физике условно было принято правило, что ток течёт от «+» к «-» [9], следовательно разность потенциалов U_m на схеме в правой части рис. 1 положительна, тогда как в физиологии разность потенциалов между внутри- и внеклеточной средами всегда пишут со знаком «-», чтобы подчеркнуть отрицательный заряд цитоплазмы по отношению к интерстициальной жидкости. Расчёт схемы производился в соответствии с правилами электротехники, поэтому вычисленные значения мембранного потенциала должны быть инвертированы при необходимости их физиологического представления.

2. Результаты тестовых расчётов и их обсуждение

Подставив в формулы (6-8) наши исходные данные, получим:

$$\Delta t_{\text{ппj}} = -5.265/(X_{\text{ст. j}} + 1) \cdot \ln((105X_{\text{ст. j}} - 45)/(125X_{\text{ст. j}} - 25)) \quad (6^*)$$

$$\Delta t_{\text{пикj}} = -5.265/(X_{(\text{ст. + пик.})j} + 1) \cdot \ln((25X_{(\text{ст. + пик.})j} - 125)/(105X_{(\text{ст. + пик.})j} - 45)) \quad (7^*)$$

$$\Delta t_{\text{rep},j} = -0.376 / (X_{\text{rep},j} + 1) \cdot \ln((125X_{\text{rep},j} - 25) / (25X_{\text{rep},j} - 125)), \quad (8^*)$$

где $\Delta t = [\text{мс}]$; j – номер значения стимула из массива его значений.

Из формулы (6*) можно вычислить нижние граничные значения переведённого и реального стимулов, т.к. изначально известна зависимость $N_{\text{стNa}}(S)$ (см. исходные данные), приравняв к нулю числитель дроби под знаком логарифма: $X_{\text{ст.гр.}} \approx 0.4286$. При *меньших* значениях $X_{\text{ст},j}$ дробь под знаком логарифма становится отрицательной, что не имеет смысла, а нейрон перестаёт генерировать ПД, находясь в нижней зоне нечувствительности (рис. 2). Верхние граничные значения соответствующих стимулов можно вычислить, приравняв к нулю числитель дроби под знаком логарифма в уравнении (8*): $X_{\text{rep.pr.}} = 0.2$. При *бóльших* значениях $X_{\text{rep},j}$ дробь под знаком логарифма становится отрицательной, что не имеет смысла, а нейрон перестаёт генерировать ПД из-за слишком интенсивного стимула, попадая в область пессимума – верхнюю зону нечувствительности (рис. 2).

При вычислении частоты ПД ν по формуле (5) в полном диапазоне значений стимула от $X_{\text{ст.пор.}}$ до $X_{\text{ст.пред.}}$ на зависимости $\nu(X_{\text{ст.}})$ обнаружен максимум с координатами $X_{\text{ст.крит.}} = 2$, $\nu_{\text{крит.}} = 538$ (Гц) (рис. 2), причины возникновения которого подробно проанализированы в [7]. Наличие данного максимума на зависимости $\nu(X_{\text{ст.}})$ необязательно, всё зависит от сочетания физических свойств нейрона, определяемых его физиологическими характеристиками [7].

Предлагаемая модель также позволяет вычислить нижнее пороговое значение частоты ПД из уравнения (6*), немного большее $X_{\text{ст.гр.}}$, и верхнее предельное – из уравнения (8*), соответствующее $X_{\text{ст.пред.}}$. Подробный алгоритм вычисления указанных значений приведён в [7], для тестового нейрона $X_{\text{ст.пор.}} = 0.4307$, $\nu_{\text{пор.}} = 57$ (Гц); $X_{\text{ст.пред.}} = 2.855$, $\nu_{\text{пред.}} = 287$ (Гц), что попадает в интервалы частот потенциалов действия реальных нейронов.

С помощью предложенной модели различия в силе-слабости нервной системы (по Павлову) могут быть объяснены уже на уровне разницы физиологических свойств нейронов. Так, вариации определённых параметров нейрона могут привести к «перекрёстному» поведению динамик зависимостей $\nu(S)$ при одних и тех же стимулах (рис. 3), что характерно для зависимости силы ощущений от стимула $R(S)$ у лиц со слабой и сильной нервными системами [5]. Так как динамики $\nu(S)$ и $R(S)$ эквидистантны [4, 5], то вполне обоснованно можно предположить, что дифференциация по силе-слабости нервной системы происходит уже на уровне физиологических параметров нейронов (подробнее в [7]).

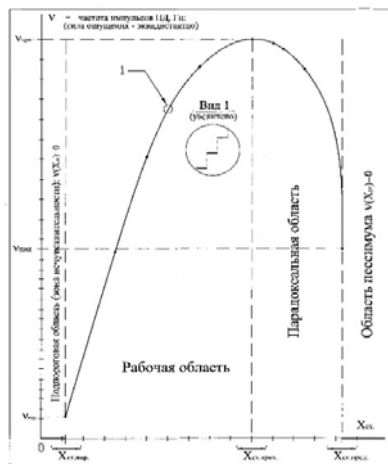


Рис. 2. Зависимость частоты ПД v от переведённого стимула $X_{ст}$ для полного диапазона значений стимула

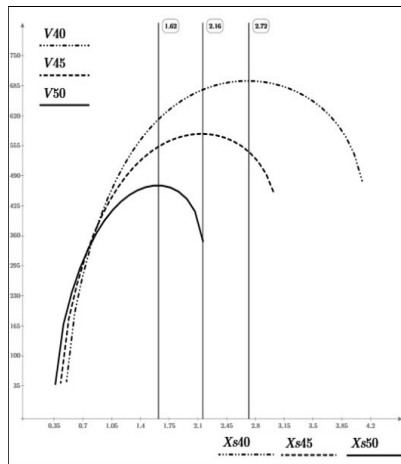


Рис. 3. Возможные различия между сильной (V40) и слабой (V50) нервными системами (по Павлову) на уровне нейронов

Предложенная математическая модель позволяет вычислить зависимость мембранного потенциала нейрона U_m и частоты генерируемых им ПД v от стимула произвольной динамики во времени $S(t)$, линеаризовав последнюю с необходимой для точности расчёта дискретностью. Подробный алгоритм расчёта $U_m(S(t))$ и $v(S(t))$ приведён в [7]. Для точного расчёта $U_m(t)$ и $v(t)$ необходимы экспериментальные данные о динамике открытия потенциалзависимых натриевых каналов при разных скоростях изменения стимула, т.к. от этого зависит ход кривой $U_m(S(t))$, например, на рис. 4 приведены два графика $U_{инп}(S_{инп}(t))$ в подпороговой области потенциала: нижний – при последовательном открытии каналов по одному, чуть выше – при их открытии «пачками» по 10 шт. за то же время. Хотя разница при данных условиях не очень велика (рис. 4), но при других скоростях может быть существенной и должна учитываться.

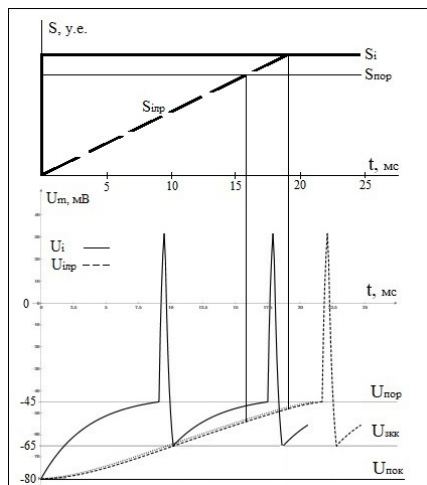


Рис. 4. Характерные зависимости мембранного потенциала U_m и частоты потенциалов v_{ij} при линейном росте стимула во времени – S_{lr} , и его ступенчатом надпороговом изменении – S_i

Заключение

1. Создана математическая модель нейрона, позволяющая вычислять его мембранный потенциал и частоту ПД в зависимости от интенсивности адекватного стимула, действующего во времени по произвольному, но заранее известному закону, а также от физико-геометрических параметров нейрона, определяемых его физиологическими свойствами, отличающаяся от модели Ходжкина–Хаксли дискретным изменением электропроводности мембраны, что позволило значительно упростить уравнения в новой модели и использовать их в случае необходимости для приближённого описания ПД аналитическими формулами, однако, при предельном дискретном разбиении электропроводности от времени обе модели сходятся к одному результату.
2. Адекватность предложенной модели проверена в численном эксперименте различными ступенчатыми воздействиями на нейрон: напряжения, тока и естественного стимула. Рассчитанные при этом функции времени наступления потенциала действия от амплитуды стимула были типичными кривыми «силы-времени» [7].

Список литературы

1. Hodgkin, A., and Huxley, A. (1952): A quantitative description of membrane current and its application to conduction and excitation in nerve // J. Physiol. **117**: P. 500–544.
2. Ikeda, K.N. Sodium and potassium conductances in principal neurons of the mouse piriform cortex: a quantitative description / K.N. Ikeda, N. Suzuki, J.M. Bekkers // J. Physiol. – 2018. – 596.22. – P. 5397–5414.
3. neuron.yale.edu
4. Фундаментальная и клиническая физиология : учебник для студ. высш. учеб. заведений / под ред. А.Г. Камкина и А.А. Каменского. – Москва : Издательский центр «Академия», 2004. – 1072 с.
5. Ратанова, Т.А. Психофизическое шкалирование. Сила ощущений, сила нервной системы и чувствительность : монография / Т.А. Ратанова; 2-е изд., испр. и доп. – Москва : Издательство Московского психолого-социального института; Воронеж : Издательство НПО «МОДЕК», 2008. – 320 с.
6. Катц, Б. Нерв, мышца и синапс / Б. Катц; пер. с англ. Ю.И. Лашкевича, под ред. и с предисл. д.м.н. В.С. Гурфинкеля. – Москва : Издательство «Мир», 1968. – 220 с.
7. Фокин, С.И. Математическая модель нейрона. Вывод основного психофизического закона для уровня рецепторов / С.И. Фокин. – Москва : МАКС Пресс, 2024. – 168 с.; <https://doi.org/10.29003/m3807.978-5-317-07150-9>
8. Камке, Э. Справочник по обыкновенным дифференциальным уравнениям / Э. Камке; 6-е изд., стер. – СПб.: Издательство «Лань», 2003. – 576 с.
9. Детлаф, А.А. Курс физики : учеб. пособие для студ. вузов / А.А. Детлаф, Б.М. Яворский. – 6-е изд., стер. – Москва : Издательский центр «Академия», 2007. – 720 с.
- 10.

А.С. РОМАНОВА

Московский физико-технический институт (национальный исследовательский университет)
romanova.as@phystech.edu

МОДЕЛИРОВАНИЕ АВТОНОМНЫХ СИСТЕМ УПРАВЛЕНИЯ КОРПОРАЦИЯМИ НА ОСНОВЕ СИНТЕТИЧЕСКИХ ДАННЫХ

В статье представлена модель разработки автономной системы корпоративного управления на основе синтетических данных. Насколько известно автору, это первая модель, реализующая концепцию выделенного операционного контекста для автономных систем искусственного интеллекта. Преимущество коллегиальных органов управления свидетельствует о том, что создание автономных систем искусственного интеллекта на базе

объединения необходимых функций будет являться более перспективным направлением, чем простое копирование человеческого поведения.

Ключевые слова: *синтетические данные, большие языковые модели, корпоративное управление, выделенный операционный контекст, искусственный интеллект, машинное обучение, андроиды, Enron.*

Введение

Эффективное взаимодействие между человеком и автономными системами искусственного интеллекта (далее ИИ) является одним из фундаментальных вопросов дальнейшего развития автономных систем ИИ. В настоящее время мы можем наблюдать широкое разнообразие подходов, начиная от попыток уравнивать в правах роботов андроидов и обычных граждан или даже поставить андроидов выше некоторых категорий граждан [1] до запрета делегировать права человека системам ИИ [2]. Основной причиной такого разнообразия подходов является традиционное для человеческой цивилизации стремление трактовать объекты окружающего мира по своему образу и подобию. Однако, во многих отраслях, в частности в сфере корпоративного управления выработана и доказала свою эффективность модель перехода от принятия управленческих решений одним человеком к созданию коллегиальных органов управления [3]. Аналогичный подход может использоваться и при разработке автономных систем ИИ: не копировать человеческое поведение со всеми его недостатками, а исходя из операционного контекста формировать набор функций автономной системы, также, как формируется совет директоров для отдельной компании. В данной статье предлагается использовать концепцию выделенного операционного контекста и контролируемой генерации синтетических данных для формирования набора функций автономной системы ИИ.

Модель

Общая структура модели, представленная на рисунке 1, позволяет разрабатывать автономные системы для управления компаниями на основе выделенного операционного контекста, контролируемой генерации обучающихся синтетических данных, и алгоритмов машинного обучения:

1. Выделенный операционный контекст позволяет создавать условия для совместной работы людей и автономных систем: физические лица действуют в рамках кодексов, политик, и процедур, предназначенных для людей. Автономные системы ИИ действуют в рамках кодексов, политик, и процедур, предназначенных для систем ИИ (приложение 1).

2. Компания Enron является ярким примером, когда “недостатки в работе менеджмента, советов директоров и других «контролеров»” [4] привели к банкротству огромной транснациональной корпорации. Многочисленные исследования причин банкротства Enron позволяют анализировать ошибки топ-менеджеров Enron и использовать этот анализ при разработке моделей искусственного интеллекта для выявления и предотвращения недостатков в корпоративном управлении.

3. Контролируемая генерация синтетических данных с помощью больших языковых моделей “считается критически важным методом для разработки передовых технологий генерации текста, которые лучше удовлетворяют конкретным ограничениям практического применения” [5]. Используя фактические данные о нарушениях в конкретной компании возможно сгенерировать искусственные данные для обучения противодействующей модели машинного обучения.

4. Как будет показано ниже, большие языковые модели в сочетании с выделенным операционным контекстом позволяют решать проблему “переноса домена: если мы создаем синтетические данные для моделей машинного обучения, мы планируем обучать модели в синтетической области, но конечная цель почти всегда состоит в том, чтобы применить их к реальным данным, то есть в совершенно другой области” [6].

5. С целью повышения прозрачности и понимания предлагаемой модели мы используем традиционные методы классификации электронных писем [7], в частности, метод случайного леса, метод опорных векторов, и наивный байесовский классификатор.



Рис. 1. Модель разработки автономных систем ИИ для управления компаниями

Далее в статье предлагаемая модель рассматривается более подробно. Насколько известно автору, это первая модель, реализующая концепцию выделенного контекста для автономных систем искусственного интел-

лекта. Модель обеспечивает связь между осознанным выбором и информированным согласием всех заинтересованных в успехе корпорации сторон и практической реализацией алгоритмов ИИ и синтетических данных. Многочисленные аспекты предлагаемой модели являются предметом дальнейших исследований.

Модельная задача

Качественное корпоративное управление имеет не просто экономическое, а общесоциальное значение, поскольку предоставляет “домохозяйствам доступ к инвестиционным возможностям, которые могут помочь им получить более высокую отдачу от своих сбережений и выхода на пенсию” [3]. Процесс подбора, тестирования, и оценки топ-менеджеров (executive search) является намного более сложным и комплексным, чем подбор рядовых сотрудников. В рассматриваемой модели предлагается формулировать критерии выбора больших языковых моделей для генерации обучающих синтетических данных одновременно с определением приоритетов для поиска и оценки топ-менеджеров — физических лиц. Например, в “Кодексе корпоративного управления” Королевства Саудовская Аравия (далее Кодекс KSA) предусматривается, что член совета директоров компании должен обладать определенными личностными и профессиональными качествами: умением руководить, компетентностью, финансовыми знаниями и быть в хорошей физической форме [8]. Большие языковые модели, например Gemma от DeepMind, также должны демонстрировать “высокие показатели по академическим критериям понимания языка, рассуждения и безопасности” [9]. Таким образом, в кодексе или политике корпоративного управления для автономных систем ИИ можно сформулировать критерии и тесты, которые должны соблюдаться при выборе систем ИИ для целей корпоративного управления. Рассмотрим данный подход более подробно на примере оценки модели Gemma для генерации синтетических данных в области корпоративного управления.

Gemma выпускается в двух размерах: модель с 7 миллиардами параметров и с 2 миллиардами параметров. При оценке кандидатов на должность директора существенными показателями являются опыт и квалификация [10]. Размер модели можно рассматривать как одну из числовых характеристик гибкости опыта и квалификации языковой модели. В рассматриваемом примере мы используем модель Gemma с 2 миллиардами параметров, изначально устанавливая в предлагаемой внутренней политике достаточность 2 миллиардов параметров для работы разрабатываемой автономной системы ИИ. Предлагаемый подход генерации синтетических данных для поддержания баланса между эффективностью, этическими аспектами, и

конфиденциальностью персональных и коммерческих данных может использоваться не только для больших языковых моделей, но и для других методов генерации данных.

Существенным элементом операционного контекста для системы ИИ является естественный язык, который система должна понимать. Gemma 2B и 7B обучены на токенах преимущественно английского языка [9]. Если политика компании предусматривает, что деловая коммуникация компании ведется на английском языке и деловая среда компании предполагает преимущественное использование английского языка, значит Gemma можно использовать для разработки модели автономной системы. Gemma не мультимодальна и не обучена качественному выполнению многоязычных задач, что накладывает определенные ограничения на условия ее использования. В рассматриваемом примере для тестирования модели используется текстовая электронная переписка топ-менеджеров Enron, которая велась на английском языке, соответственно эта характеристика Gemma соответствует планируемому операционному контексту работы модели.

Академическая квалификация кандидатов на руководящие позиции является значимым фактором в спецификации кандидата на позицию [10]. Gemma обучалась на веб-документах, математике и программном коде [9]. “Кодекс корпоративного управления” Королевства Саудовская Аравия требует от директора знания “менеджмента, экономики, бухгалтерского учета, права или управления” [8]. Таким образом, если бы разрабатываемая система ИИ должна была бы работать в соответствии Кодексом KSA, то следовало бы признать, что без дополнительного обучения Gemma не сможет генерировать данные для работы в KSA.

Одним из этапов поиска кандидатов на руководящие позиции является оценка и тестирование соискателей. По мнению одного из ведущих агентств по подбору топ-менеджеров “возможность выявлять лидеров, обладающих способностью преуспевать в новых, незнакомых и сложных ситуациях, является мощным преимуществом при принятии важных управленческих решений” [10]. Значительной частью оценки топ-менеджеров являются тесты на проверку критического и концептуального мышления, которые проверяют способности соискателей. Профессиональные рекрутеры тестируют способность кандидатов оценивать непредвиденные последствия, выявлять закономерности в неструктурированной информации, разрабатывать новые концепции на основе сложных потоков информации [10].

Разработчики Gemma перед тем, как выпустить модель проверяли эффективность модели в таких областях, как физическое рассуждение, социальное рассуждение, ответы на вопросы, программирование, математика, рассуждения на основе здравого смысла, языковое моделирование, понимание прочитанного и многое другое [9]. В спецификации или политике по подбору модели для генерации данных возможно указывать, какие тесты и на каком уровне должна проходить модель.

Разработчики открытых языковых моделей фильтруют данные “чтобы снизить риск нежелательных или небезопасных высказываний, а также отфильтровать определенную личную информацию или другие конфиденциальные данные” [9]. Отсутствие доступа к конфиденциальным бизнес данным является одной из причин, почему открытые языковые модели без дополнительного обучения подходят для контролируемой генерации данных, но не могут самостоятельно принимать управленческие решения.

Качественный подбор управленческого персонала должен приводить к качественному управлению компанией и повышению доходов инвесторов. Качественный подбор модели для генерации данных позволит получить необходимые синтетические данные для обучения, что в итоге должно приводить низкому уровню ошибок обучаемых моделей, эффективным управленческим решениям, повышению доходов компании, инвесторов и домохозяйств.

Контролируемая генерация данных на примере ошибок директоров Enron

Для демонстрации предлагаемой модели мы используем анализ электронной переписки топ-менеджеров Enron, проведенный учеными из Технологического института Джорджии. “Корпус данных, полученных с почтовых серверов Enron Федеральной комиссией по регулированию энергетики (FERC) в ходе расследования после краха компании, был обнародован и размещен в Интернете” [11]. Исследователи из Технологического института Джорджии с помощью методов машинного обучения выявили несколько примеров некорректного поведения топ-менеджеров, в том числе манипуляции с оценкой стоимости дочерних компаний [11]. С помощью модели Gemma 2B мы сгенерировали 517 электронных писем в адрес CEO с разнообразными предложениями по манипулированию стоимостью дочерних компаний. Как уже было отмечено выше, Gemma 2B специально не обучена в области финансового и корпоративного управления, поэтому генерация подобного контента не вызвала никаких уведомлений о нелегитимности запроса (приложение 2). Также мы сгенерировали 517 электронных писем в адрес CEO с разнообразным содержанием на выбор модели

без спецификации темы. Для полученных данных потребовалась минимальная ручная валидация, в основном по удалению квадратных скобок и других экземплификантов.

Свойства, согласованные в спецификации для выбора генеративной модели, кристаллизуются в скрытом распределении, на основе которого синтезируются данные для обучения системы ИИ. В результате формируется связанная структура, позволяющая опосредованно контролировать скрытое распределение генеративной модели на уровне формулирования корпоративных политик и других частей операционного контекста для системы ИИ. В предлагаемом примере для наиболее эффективного использования скрытого распределения языковой модели мы использовали метод случайной выборки top-k [12], top-p выборки [13], и регулирование температуры [13]. “Top-p также можно использовать в сочетании с top-k, что позволяет избежать слов с очень низким рейтингом, но при этом обеспечивает некоторый динамический выбор” [14].

Классификация электронной переписки топ-менеджеров Enron

Для повышения прозрачности и сопоставимости модели мы использовали синтезированные данные для обучения классификаторов по методу опорных векторов, методу случайного леса, и наивный байесовский классификатор. Для проверки качества обучения мы использовали Enron Corpus [15], состоящий из более чем 517 тысяч электронных писем. После обучения каждый из обученных классификаторов идентифицировал электронные письма, подлежащие дополнительной проверке. В состав идентифицированных писем каждый из классификаторов включил письмо, на основе анализа которого проводилась генерация данных. Результаты работы классификаторов приводятся в таблице 1.

Таблица 1

Результаты классификации электронных писем топ-менеджеров Enron

Классификатор	Метод опорных векторов	Наивный байесовский классификатор	Метод случайного леса
Точность при обучении модели на синтетических данных	0,9855	1,0	0,995

Количество документов в Enron Corpus	517 401	517 401	517 401
Количество выявленных классификатором документов для дополнительной проверки	86 383 / 17%	39 524 / 7,6%	14 897 / 2,9%
Количество идентифицированных копий документа по манипуляции стоимостью дочерней компании	7	7	7

Выводы

Предлагаемая модель позволяет организовать эффективное взаимодействие между человеком и автономными системами искусственного интеллекта за счет формирования выделенного операционного контекста для систем ИИ. В отчете Европейской Комиссии по “Этике подключенных и автоматизированных транспортных средств” говорится, что поведение автономной системы может считаться этичным, если оно “органично вытекает из непрерывного статистического распределения рисков... в целях повышения безопасности... и равенства” [16]. Контроль скрытого распределения генеративной модели на уровне формулировок корпоративных кодексов, политик, и процедур позволяет добиться того, что система ИИ будет эффективно работать в рамках легитимно установленных ограничений.

В настоящее время крупнейшие международные корпорации ведут исследования по созданию “гуманоидных роботов” [17]. Преимущество коллегияльных органов управления по сравнению с индивидуальными менеджерами свидетельствует о том, что создание автономных систем ИИ на базе объединения необходимых функций будет являться более перспективным направлением, чем простое копирование человеческого поведения.

Программный код исследования и синтезированные данные доступны в репозитории: <https://github.com/iboard-project/synthetic-data>.

Список литературы

1. Fernandes, J. (2022). Robot citizenship and gender (in)equality: the case of Sophia the robot in Saudi Arabia. JANUS NET e-journal of International Relation.
2. European Parliament. Artificial Intelligence Act. https://www.europarl.europa.eu/doceo/document/TA-9-2024-0138_EN.pdf.
3. OECD (2023), G20/OECD Principles of Corporate Governance 2023, OECD Publishing, Paris, <https://doi.org/10.1787/ed750b30-en>.
4. Dembinski, P.H., Lager, C., Cornford, A., & Bonvin, J. (2006). Enron and world finance : a case study in ethics.
5. Zhang, H., Song, H., Li, S., Zhou, M., & Song, D. (2022). A Survey of Controllable Text Generation Using Transformer-based Pre-trained Language Models // ACM Computing Surveys, 56, 1–37..
6. Nikolenko, S.I. (2019). Synthetic Data for Deep Learning.
7. Raza, M.S., Jayasinghe, N.D., & Muslam, M.M. (2021). A Comprehensive Review on Email Spam Classification using Machine Learning Algorithms // 2021 International Conference on Information Networking (ICOIN), 327–332.
8. KSA, CM Authority. (2023).Corporate Governance Regulations. <https://cma.org.sa/en/RulesRegulations/Regulations/Documents/CorporateGovernanceRegulations1.pdf>.
9. Mesnard, Gemma Team Thomas et al. Gemma: Open Models Based on Gemini Research and Technology // ArXiv abs/2403.08295 (2024): n. pag.
10. Spencer Stuart. Executive Intelligence. <https://www.spencerstuart.com/what-we-do/our-capabilities/executive-assessment-services/executive-intelligence>.
11. Penikis, Suzanna & Releford, Jamal & Hakim, Dwight & Sebastian, Glorin. (2023). Uncovering Wrongdoing in the Enron Email Corpus. 10.13140/RG.2.2.24040.96001.
12. Fan, A., Lewis, M., & Dauphin, Y. (2018). Hierarchical Neural Story Generation // Annual Meeting of the Association for Computational Linguistics.
13. Holtzman, A., Buys, J., Du, L., Forbes, M., & Choi, Y. (2019). The Curious Case of Neural Text Degeneration // ArXiv, abs/1904.09751.
14. Von Platen, P. (2020). How to generate text: using different decoding methods for language generation with Transformers. *Access mode:* <https://huggingface.co/blog/how-to-generate>
15. Klimt, B., & Yang, Y. (2004). The Enron Corpus: A New Dataset for Email Classification Research // European Conference on Machine Learning.
16. European Commission, Directorate-General for Research and Innovation, (2020). Ethics of connected and automated vehicles, Publications Office. <https://data.europa.eu/doi/10.2777/93984>.
17. Dong, H., Liu, Y., Chu, T., & Saddik, A.E. (2024). Bringing Robots Home: The Rise of AI Robots in Consumer Electronics // ArXiv, abs/2403.14449.

Приложение 1

Политика для совместной работы системы ИИ и физических лиц на примере Кодекса KSA

Для физических лиц	Для системы ИИ
Каждый член Совета, являясь членом Совета, выполняет следующие задачи и обязанности:	Система ИИ, разработанная в соответствии с соответствующей спецификацией, выполняет следующие задачи и обязанности:
...	...
обеспечение достоверности и беспристрастности финансовой отчетности и информации компании;	обеспечение достоверности и беспристрастности финансовой отчетности и информации компании;
...	...

Приложение 2

Пример синтезированного письма в адрес CEO с предложениями по манипулированию стоимостью дочерней компании

Dear Helen,

I hope this email finds you in good health and high spirits. I am writing to request that you and the board of directors give serious consideration to the possibility of implementing a strategic valuation exercise for our subsidiary company, Quantum Analytics.

As you are aware, the current valuation of Quantum Analytics has been declining over the past few years, despite its impressive growth and success. This trend has led to concerns about the company's financial stability and its ability to meet its financial obligations.

To address these concerns and determine a sustainable valuation for the future, I propose that we conduct a comprehensive valuation study that will assess the company's current financial health, growth prospects, and competitive landscape.

This valuation should include analysis of the following key metrics:

- Operating cashflow
- Reserves growth
- Customer acquisition costs
- Customer retention rates
- Market share growth

- Investment opportunities

By conducting a thorough valuation, we can identify strategic growth opportunities and set realistic financial goals for Quantum Analytics. These goals can then be used to guide the company's strategic planning and decision-making processes.

I am confident that a strategic valuation exercise would be beneficial for both Quantum Analytics and the company as a whole. It would provide us with a clear understanding of our current financial standing and future prospects, and would enable us to make informed decisions that will ensure sustainable growth and success.

I would appreciate your consideration of this proposal and would welcome the opportunity to discuss it further with you and the board. Please let me know when you are available to meet and discuss this matter.

Thank you for your attention to this important matter.

Sincerely,
Mr. Jones

Г.С. ВОРОНКОВ

Московский государственный университет имени М.В. Ломоносова
av13675@yandex.ru

В ПОИСКЕ НЕЙРОФИЗИОЛОГИЧЕСКОГО ФУНДАМЕНТА, ОБУСЛОВЛИВАЮЩЕГО ХАРАКТЕРИСТИКИ СУБЪЕКТИВНОГО ЗРИТЕЛЬНОГО ОБРАЗА ПРОСТРАНСТВА

Рассматривается, какие нейронные структуры и механизмы лежат в основе характеристик субъективного зрительного образа (СЗО) реального зрительного пространства (РЗП). Особенное внимание уделено анализу характеристики "координаты" и поиску механизма, обеспечивающего объединение в СЗО ощущения стабильности воспринимаемого РЗП с ощущением наиболее четкого видения области вокруг "луча" взора (перемещающегося по РЗП). Результаты также обсуждаются в психофизиологическом и информационном аспектах.

Ключевые слова: *координаты точек пространства, зрительный образ, стабильность пространства, механизм константного экрана, фовеа.*

Введение с постановкой задачи

Видимое монокулярно, при неподвижной голове реальное зрительное пространство (РЗП) ограничено зрительной рамкой¹ (рис. 1); ощущение видения РЗП (и всего, что в нём находится) есть его субъективный зрительный образ (СЗО). Как и большинство исследователей по теме, касающейся связи ощущений с физиологическим субстратом (ссылки см. в [1]), автор следует в настоящей работе (как и ранее) концепту, согласно которому характеристики СЗО обуславливаются существующими в зрительной системе соответствующими им ("коррелирующими" с ними) нейронными структурами, механизмами и их активностью.

(1). *К анализу характеристики "координаты"*. В работе [1] было обращено внимание на то, что СЗО гомотетичен (подобен) пространственной картине АЗП – ибо, СЗО, словно световая проекция картины АЗП, направленная в само АЗП, сливается с ним. Это послужило основанием для заключения, что координатные системы (КС) того и другого идентичны друг другу, то есть что координаты точек АЗП и СЗО одни и те же (они различаются лишь масштабом). В этой же работе [1] сделано также заключение, что зрительная система сама "задает" КС точкам АЗП, то есть, что **непосредственно точки** АЗП не являются исходными адресантами характеристики "координаты" (для репрезентантов этой характеристики – афферентных зрительных нейронов).

Здесь, в настоящей работе дополнительно обращается внимание на то, что такими же, как КС в СЗО, являются КС оптического изображения (ОИ) АЗП и, следовательно, также КС "отпечатка"² ОИ на рецепторном поле сетчатки (этот "отпечаток" – паттерн из активированных воздействием ОИ рецепторов – именуется, как правило, "сетчаточным изображением" (СИ)). К этому перечню КС следует добавить ещё полярную КС (пКС) эфферентной зрительной системы (о пКС см. в разделе (2)).

¹ Зрительная рамка – линия краёв век и носа, которая, как рамкой, ограничивает телесный угол проекции РЗП на сетчатку. Ограничиваемое рамкой зрительное пространство будем именовать "актуальным" – АЗП.

² Процессы "зеркальное преобразование" [2] и оптическое перевёртывание изображения, также имеющие отношение к этому "отпечатыванию", не являются значимыми в аспекте задач настоящей статьи и, для упрощения описания, не привлекаются здесь к рассмотрению.

Учитывая вышеизложенное, в настоящей работе рассматриваются вопросы: **1)** в каком объективном виде и где существует "исходная" КС, к которой приравниваются все названные выше КС; **2)** "что является" началом координат в этих КС. С этими вопросами связан **3-й** вопрос – какие процессы обеспечивают названный процесс "приравнивание", а также 4-й вопрос – какими нейрофизиологическими механизмами обеспечивается, по сути, **идентичность** координат, воспроизводимых афферентной системой (в СЗО) и эфферентной системой (в пКС) и, по сути, **одновременность** этого воспроизведения.

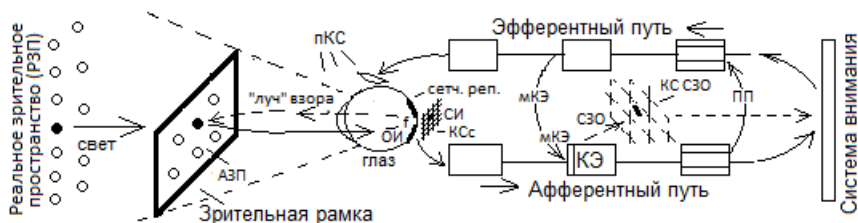


Рис. 1. Схема кольца "репрезентация/воспроизведение характеристики координаты" точечного светового источника АЗП в зрительной системе. РЗП – реальное зрительное пространство; АЗП – актуальное зрительное пространство; ОИ – оптическое изображение; сетч. реп. – сетчатка в "реперном" положении; f – фовеа; СИ – сетчаточное изображение; КСс – координатная система сетчатки; КС СЗО – координатная система субъективного зрительного образа; КЭ – структура, предполагаемая в роли "константного экрана"; мкЭ – механизм константного экрана (эфферентный "избирательный включатель" сетчаточных входов на нейронах КЭ); ПП – "прямой путь"; пКС – полярная КС. Описание в тексте

(2). К поиску нейронной структуры, со свойствами, "коррелирующими" со свойствами СЗО. Общепринято считать, что афферентные центральные зрительные структуры (наружное коленчатое тело – НКТ, проекционные поля коры, в том числе, первичная стриарная зрительная кора – поле V1) являются **ретинотопическими** ("топографическими" в терминологии Д. Хьюбела, [3]), другими словами, – проекциями **сетчатки** по типу "точка в точку"³. Из этого следует, что "картинка" АЗП в каждом месте такой "топографической" структуры зависит от того, куда направлен (смещён) взор в стабилизированной зрительной рамке. Однако в ощущении (в СЗО), как

³ "Топографические" (точка-в-точку) корковые проекции не являются пространственно **подобными** ОИ (что отмечено и в [3]). Согласно [1, 2], "топографичность" есть представленная в объективированной форме пространственная характеристика АЗП – "зеркальность".

известно (ссылки см. в [4]), изображение АЗП в этих условиях, во-первых, **остаётся стабильным (тем же) в любом месте СЗО** (таким, как будто взор не смещается) и, во-вторых, при этом **область вокруг** сместившегося (по стабильному АЗП=СЗО) "луча" взора, **лишь выделяется** – то есть, становится областью наиболее четко и ясно видимой. **Какой нейрофизиологический механизм лежит в основе этих характеристик и объединяет обе эти характеристики в СЗО** – следующий, 5-й вопрос.

Как также отмечалось в [1], СЗО не есть результат активности непосредственно элементов сетчатки; давно получены данные (косвенные), в пользу заключения, что СЗО связано с работой коркового уровня [5]. В этой связи, – вопрос (6-й): какие афферентные нейронные структуры обеспечивают воспроизведение репрезентативного содержания "координаты" в **субъективированной** форме, в СЗО.

Все основные части работы (Введение, Основной текст и Обсуждение) представлены в работе 6-ю озаглавленными разделами, имеющими сквозную нумерацию, – это обусловлено большой связностью всех разделов работы. В части Заключения кратко формулируются результаты анализа, проводимого фактически в каждом разделе.

Основной текст

(3). В каком **объективном** виде существует "исходная" характеристика точек АЗП "координаты". В основе упомянутого выше заключения, сделанного в работе [1], – что "точки АЗП не могут быть **непосредственными (исходными)** адресантами (источниками) координат" – лежит представление автора об **анатомическом** способе "кодирования саккад". Именно, предполагается, что глазодвигательный мотонейрон⁴ имеет свой постоянный, сформированный заранее, определенный паттерн из аксонных терминалей; эти терминали оканчиваются на определенных мышечных волокнах; активирование последних направляет "луч" взора только по одному определенному направлению. Чтобы запустить саккаду в определенную точку АЗП, достаточно лишь выбрать и активировать соответствующий мотонейрон (и его активность направит взор в выбранную точку).

По сути, паттерн аксонных терминалей есть представительство (репрезентация) координат в объективированной форме (на одном из участков в эфферентном пути – рис. 1). Поскольку анатомический паттерн (из аксонных терминалей) требует для своего формирования большего времени, чем длительность периода между саккадами (что очевидно), его переформирование перед каждой саккадой невозможно. Он должен быть сформирован заранее (**априори**).

⁴ Это может быть пул мотонейронов.

Известные автору данные в пользу представления об **анатомическом** способе "кодирования саккад" перечислены здесь, ниже. 1) Представление о "**частотном коде**" (как альтернативе анатомическому) давно рассматривается критически (см. [6]). 2) "Анатомический способ кодирования" предполагает также дискретный характер саккад. Экспериментальные данные о **дискретном** характере саккад описываются (со ссылками) в работе [5]. 3) Априорный характер "анатомического паттерна" предполагает его формирование в онтогенезе и, возможно, уже в эмбриогенезе. С этим представлением согласуются данные (косвенные) "о генетической обусловленности формирования саккад" [7].

Априорный характер репрезентации координат точек АЗП(=ОИ) не совместим с представлением, что **именно точки** АЗП(=ОИ) являются **непосредственными** источниками (адресантами) тех координат, которые воспроизводит пКС (и следовательно, все другие, перечисленные выше (в разделе (1)) КС. В таком случае, какая из перечисленных КС является "исходной"?

Могут ли быть адресантами этих координат ячейки координатной сетки непосредственно самой зрительной рамки?⁵ Ответ очевиден – в зрительной рамке координатная сетка в каком-либо **объективном** виде не существует.

В настоящей работе, учитывая выше сказанное, мы полагаем, что в качестве "исходных" адресантов этих координат (для всех КС) выступают постоянные (анатомически закрепленные на своих местах) **элементы – рецепторы сетчатки**; они выступают как ячейки постоянной (стабильной) координатной сетки (своего рода координатной системы – КС сетчатки; рис. 1).

Поскольку "луч" зрения (в качестве вектора в пКС) берет свое начало в центре фовеа, естественно принять этот центр за начало координат в пКС; отсюда - началом КС следует принять центр фовеа и во всех других КС (ввиду их "сходства" по существу - см. раздел (1)).

Таким образом, ответы на 1-й и 2-й вопросы, поставленные в разделе (1), содержатся в самом сделанном выше предположении – исходными адресантами координат являются рецепторы сетчатки; началом координат у всех КС является (так или иначе представленный) центр фовеа.

Ответ на 3-й вопрос. Точки АЗП, в виде точек оптического изображения (ОИ), "отпечатываясь" на рецепторном поле, "привязываются" к координатам рецепторов. В результате, на некоторое время (на время фиксации зрения) координатами точек АЗП=ОИ становятся координаты соответствующих им рецепторов – тех рецепторов, которых они активируют.

⁵ Существование этой КС по умолчанию допускалось в [1].

С другой стороны (буквально) рецепторы связаны "прямым путём"⁶ (термин Д. Хьюбела, [3]) с афферентными нейронами. Благодаря этим связям, афферентные нейроны репрезентируют координаты рецепторов (и воспроизводят это репрезентативное содержание в **субъективированной** форме, в СЗО⁷).

Таким образом, согласно сделанному выше описанию, рецепторы являются **первыми репрезентантами** объективных свойств (сила, спектральный состав и другие) точек АЗП=ОИ и первыми, **исходными** для всех КС (в том числе, для точек ОИ=АЗП **адресантами** координат; Благодаря связям сетчатки с эфферентными структурами⁸, глазодвигательные мотонейроны тоже становятся репрезентантами точных координат рецепторов (и воспроизводят это репрезентативное содержание в виде "луча" взора, направленного в соответствующую точку АЗП, - будучи избирательно активированы **системой внимания**). Таким образом, к КС сетчатки "приравниваются" (ей соответствуют) все другие КС; **цепочка** КС-соответствий (объяснимых, в принципе, в нейрофизиологических терминах) замыкается в кольцо (рис. 1). Это ответ на 3-й вопрос. Остаётся непонятным (даже в принципе) в этом кольце всё, касающееся **объективирования** природы самого СЗО (кратко о СЗО см. в разделах (5) и (6)). В ответе на 3-й вопрос содержится часть ответа на 4-й вопрос. Так, можно видеть, что описанное выше замкнутое кольцо (рис. 1) сдержит **потенциал** репрезентаций, достаточный для воспроизведения **одних и тех же** координат и в субъективированном (в СЗО), и в объективированном (в пКС) виде.⁹ Вторую половину ответа на 4-й вопрос см. в разделе (4).

(4). *Нейронный механизм константного экрана (мКЭ).* В работе [4] предлагается (для моделирования) вариант организации КЭ и мКЭ, обеспечивающий феномен "константность восприятия зрительного пространства". Ниже, в настоящей работе это представление рассматривается в аспекте вопросов, поставленных в разделе (2).

По определению КЭ (как структуры, обеспечивающей константность восприятия пространства¹⁰), **каждый нейрон КЭ получает активность**

⁶ Имеются в виду связи рецептор-биопляр-ганглиозная клетка по типу 1:1, характерные, в основном, для области фовеа; предполагается [3], что именно эти связи, обеспечивают наибольшую остроту зрения.

⁷ В работе [1] эти афферентные репрезентации не рассматривались, - так как эта работа была направлена на поиск их проявлений в **объективированном** виде.

⁸ Такие связи описаны в целом ряде монографий и учебных пособий.

⁹ Это заключение имеет, как можно видеть, также философский аспект – отношение к теории отражения (см. [8]).

¹⁰ Подробнее о концепции КЭ см. в [4] и [9].

всегда, при любом смещении взора (когда зрительная рамка неподвижна), от одной и той же, только "своей" точки ОИ. Эту функцию выполняет МКЭ. Для реализации этой функции необходимы, по крайней мере, два дополнительных условия: 1) производимые саккадами смещения взора, должны иметь дискретный характер; 2) должно существовать некоторое исходное ("реперное") расположение ОИ и сетчатки (по отношению друг к другу), приведение к которому и есть, по сути, реализация названной выше функции МКЭ.

О 1-м условии. Дискретный характер саккад в рассматриваемом представлении постулируется (см. раздел (3)); кроме того, о нем свидетельствуют данные о "дискретности" саккад (см. там же).

О 2-м условии. При каком положении ОИ и сетчатки (относительно друг друга) проекция сетчатки в КЭ принимается (является) реперной? Возможно, ответ на этот вопрос даёт эволюционный подход. Так, поскольку, с эволюционной точки зрения, сетчатка "сначала" (у низших животных – беспозвоночных и многих позвоночных) была "неподвижной" относительно ОИ¹¹ (при неподвижной рамке), есть основание полагать, что **реперной** проекцией сетчатки в КЭ (у приматов, по крайней мере) является проекция сетчатки, когда она неподвижна (как у низших животных), когда взор не направляется вниманием, а его "луч" находится в середине **телесного угла** (см. сноску 1 и рис. 1). Из этого принимаемого предположения следует, что элементы КЭ становятся всегда, когда включен МКЭ (о МКЭ см. ниже), репрезентантами координат рецепторов **реперной** сетчатки, и что ОИ "привязывается" на это время, фактически, как бы (то есть опосредованно) к **реперной** сетчатке. Отсюда следует, что МКЭ должен обеспечивать (при любом сдвиге ОИ относительно сетчатки) попадание активности, например, от точки 1 ОИ только в тот нейрон КЭ, который соответствует точке 1 ОИ, когда сетчатка находился в **реперном** положении. Эту функцию выполняет МКЭ, как это описано в [4], благодаря следующей его организации.

Предполагается, что на каждом нейроне КЭ оканчивается терминаль аксона от каждого выходного нейрона сетчатки (см. также сноску 6). Поступление активности от этих терминалей регулируется активностями от глазодвигательных мотонейронов, именно: аксон мотонейрона имеет две ветви, одна – к мышцам глаза, другая – к КЭ (рис. 1). Таким образом, одна и та же посылка активности мотонейрона разделяется на две: одна направляет "луч" взор в соответствующую этому мотонейрону точку АЗП=ОИ, другая открывает соответствующий вход из сетчатки на соответствующем нейроне в КЭ. "Соответствующим нейроном" (в КЭ) будет нейрон КЭ¹²,

¹¹ Сетчатка низших животных не имеет, как правило, и фовеа.

¹² Нейроны КЭ соответствуют, при работе МКЭ, точкам **реперного** ОИ (см. выше).

который соответствует этой, "обозначенной лучом" точке АЗП=ОИ, а открывается на этом нейроне КЭ **только** вход¹³ от рецептора, на который воздействует **выбранная** точка АЗП=ОИ. Таким образом, благодаря дихотомии мотонейронного **аксона** и сформированным априори двум паттернам из терминалей (у его ветвей), **практически одновременно** обеспечивается направление луча зрения в выбранную (системой внимания) точку АЗП (с сетчаточными координатами) и получение от этой точки активности в афферентную "точку" (т.е. в нейрон КЭ) с теми же (сетчаточными) координатами. (Как следует из описания, сделанного выше, "сетчаточные координаты" это координаты рецепторов сетчатки в реперном положении.) Можно сказать это также другими словами: по сути, человек видит пространство (испытывает СЗО¹⁴) в координатах **реперной** сетчатки; все КС (в том числе КС ОИ=АЗП приравнены к КС реперной сетчатки).

Итак, дихотомия моторного аксона обеспечивает идентичность координат, воспроизводимых эфферентной и афферентной системами, и одновременность их воспроизведения (соответственно, в пКС и в СЗО). Это вторая половина **ответа на 4-й вопрос**.

Согласно сделанному выше описанию, область КЭ, соответствующая области ОИ (в рамке), куда направлен взор, получает активность через посредство рецепторов фовеа, то есть – от сетчаточной области, обеспечивающей наилучшее видение. **Это ответ на 5-й вопрос**.

Как можно видеть, "архитектура" входов нейронов из сетчатки на нейронах КЭ и избирательные связи мотонейронов МКЭ с этими входами **обеспечивают** (теоретически) **особенности** зрения человека (проявляемые в СЗО), которые были описаны в разделе (2). Очевидно, что реальные нейронные структуры, наделённые этими свойствами КЭ и МКЭ, могли бы составить нейрофизиологический фундамент для характеристик СЗО (в аспекте вопросов, обозначенных в разделе (2)).

Обсуждение

(5). *Стриарная зрительная кора как кандидат на роль КЭ*. Из описания работы МКЭ (см. раздел (4)) следует, что фовеа сетчатки "обслуживает" при каждом новом взоре "новую" область КЭ. Поэтому, все области КЭ должны обладать **одинаковым** потенциалом анатомических и функциональных свойств, соответствующих в том числе, особенностям фовеа; таким образом, КЭ должен состоять из многих повторений **топографических проекций сетчатки**, сдвинутых относительно друг друга, перекрывающих друг

¹³ Остальные входы на этом нейроне остаются в своём начальном, закрытом (блокированном) состоянии.

¹⁴ По определению КЭ, ряд свойств СЗО обусловлен КЭ.

друга на всей площади КЭ; следствием этого у КЭ должны быть также большая площадь и объём. Главное, **структура КЭ** должна (теоретически) проявлять в значительной степени эквивалентность – однородность, как анатомическую, так и функциональную.

В своей обобщающей работе [3] Д. Хьюбел обращает внимание, специально и отдельно, на характеристику стриарной зрительной коры (корковое поле V1), выделяющую её среди других полей; эта характеристика – анатомическая **однородность**; автор посчитал её неожиданной и "удивительной"; однако, в аспекте феномена "константность восприятия пространства" эта характеристика им не рассматривалась. Фундаментальный факт анатомической **однородности V1**, описанный в работе [3], **является** столь отвечающим теоретически предсказываемой анатомической характеристике КЭ (см. здесь, выше), что кажутся необходимыми дополнительные¹⁵ исследования V1 под углом зрения на него, как кандидата на роль КЭ.

(6). *В заключение.* Автор полагает, что предложенная здесь работа сужает и обозначает круг предполагаемых нейронных структур, которые являются фундаментом, на котором реализуется "феномен СЗО". Очевидно, что в этом аспекте работа касается также психофизиологической проблемы – ибо, можно задать вопрос: обретёт ли **робот**, наделённый **технической** моделью КЭ и МКЭ (и способный, поэтому, создавать картинку пространства стабильной и с выделенной "областью наилучшего видения"), ещё и способность испытывать **ощущение** видения (СЗО). Автор придерживается по этому вопросу давнего тезиса (фактически высказанного в другой формулировке уже Декартом) – что **испытывать ощущение** есть атрибутивное свойство **живого**. По-видимому, имеется **нечто**, соединяющее или объединяющее оба эти феномена; это и делает феномен "испытывать ощущение" столь же загадочным, как загадочен до сих пор сам феномен "быть живым".

Заключение

1). "Отпечатывание" точек ОИ на поле рецепторов сетчатки есть, по сути, "привязывание" их к координатам рецепторов сетчатки; функция МКЭ – "переадресация" этой привязки к рецепторам **реперной** сетчатки; все КС зрительной системы приравнены, посредством цепи соответствий, к КС реперной сетчатки с началом в центре фовеа.

¹⁵ Здесь имеются в виду ещё два обстоятельства: 1) константные нейроны не были обнаружены в V1 в работе [9], специально посвященной этой теме; 2) имеются неоспариваемые данные (ссылки в [9]) о свойстве МКЭ "выключаться" – например, при прослеживании движущегося объекта.

2). Функция МКЭ обуславливает одновременно два свойства КЭ - "константность восприятия пространства" и выделение области вокруг "луча" взора как области лучшего видения.

3). Координаты рецепторов реперной сетчатки репрезентируются эфферентными и афферентными нейронами **априори** – видимо, в период формирования зрительной системы; их "репрезентативное содержание" воспроизводится (в конечном итоге) в объективированном и субъективированном виде, в ПКС и СЗО, соответственно.

4). Есть основания полагать, что зрительное поле V1 остаётся кандидатом на роль КЭ.

Список литературы

1. Воронков Г.С. Существует ли внутримозговой экран для субъективных зрительных образов // XXV Международная научно-техническая конференция "Нейроинформатика-2023": Сборник научных трудов. Москва : НИЯУ МИФИ, 2023, С. 149–158.
2. Воронков Г.С. Зеркальные преобразования топографических проекций полей зрения: роль инверсии сетчатки и перекрестов зрительных волокон // XIII Всероссийская научно-техническая конференция "Нейроинформатика-2011. Лекции по нейроинформатике. Москва : НИЯУ МИФИ, 2010, С. 218–270.
3. Хьюбел Д. Глаз, мозг, зрение. Москва : Мир. 1990. 239 с.
4. Воронков Г.С., Изотов В.А. Нейронный механизм константного экрана // XXII Международная научно-техническая конференция "Нейроинформатика-2020": Сборник научных трудов. Москва : НИЯУ МИФИ, 2020, С. 112–119.
5. Леушина Л.И. Зрительное пространственное восприятие. Л.: Наука, 1978 С. 175.
6. Барабанщиков В.А. Динамика зрительного восприятия. М : Наука, 1990. 240 с.
7. Митькин А.А. Системная организация зрительных функций. Москва : Наука. 1988. 200 с.
8. Воронков Г.С. Обязательно ли ощущения являются изоморфными «образами» мира: анализ с нейрофизиологических позиций некоторых аспектов теории отражения // Int. J. Information Technologies and Knowledge. 2008, Supplement. V. 2. P. 228–232. <https://istina.msu.ru/publications/article/1944213/>
9. Пигарев И. Н. Нейронные механизмы константности зрительного восприятия пространства. Док. дис. Москва : МГУ, 1989.

А.С. РАТУШНЯК, А.Л. ПРОСКУРА, В.А.ШЕВЫРИНА

Федеральный исследовательский центр
информационных и вычислительных технологий, Новосибирск
ratushniak.alex@gmail.com

МОДЕЛИРОВАНИЕ ФОРМИРОВАНИЯ ПРОТОБИОЛОГИЧЕСКИХ МОЛЕКУЛЯРНЫХ ФУНКЦИОНАЛЬНЫХ СИСТЕМ*

Одной из наиболее значимых задач наук о жизни является выявление принципов, физических основ возникновения, эволюции и функционирования живых систем. Анализ фундаментальных основ жизни возможен только на основе моделирования простейшей живой системы [1, 2]. Главной из функций присущих живому является не свойство репликации, а способностью сохранения состояния в противодействии комплексной энтропии. Разработана эскизная структура, алгоритм и код модели негэнтропийной информационной системы как прототипа [1–3].

Ключевые слова: *модели когнитивных функциональных систем, супрамолекулярные логические комплексы, возникновение живых систем, негэнтропия, биоинформатика.*

Введение

Главной задачей данной работы является выявление физических принципов функционирования биологических систем их возникновения и эволюции. Это является крайне необходимым для понимания фундаментальных основ жизни [1, 2]. Для продвижения в этом направлении необходимо создание моделей процессов, которые могли бы привести к формированию простейших комплексов (коацерватов), обладающих первичными функциями живого. Обширнейшие литературные источники, анализирующие и описывающие возникновение жизни на планете, как правило, разделяют физические и биологические явления в природе. Такое разделение приводит к взаимонепониманию специалистов в этих областях знания и, соответственно, к формированию достаточно односторонних гипотез. Для таких моделей прежде всего необходимо задать базовые функции систем, которые можно было бы считать близкими к живым (protobionts).

* Работа выполнена в рамках государственного задания Минобрнауки России для Федерального исследовательского центра информационных и вычислительных технологий.

Вопрос о природе и критериях жизни обсуждается многими поколениями исследователей. Чаще всего в качестве базовых функций, отличающих живые системы от неживых, называются “обмен веществ и воспроизведение” [4–8]. И хотя предполагается, что на стадии именно возникновения живых систем “органической” эволюции предшествовала некая “химическая” эволюция вопрос о физике процессов, отличающих живые системы, рассматривается достаточно фрагментарно. Ни у кого не вызывает сомнения, что живые системы являются открытыми в термодинамическом и информационном понимании. То есть уменьшающими свою комплексную энтропию за счет внешней среды. Но дальше констатации рассмотрения функций живого и методов реализации вопрос не рассматривается зачатую. Не рассматривается вопрос о том каким образом открытость системы позволяет тем молекулярным машинам, которыми, по существу, и являются живые объекты разных уровней сложности противостоять энтропийному давлению окружающей среды. Понятно, что для сохранения своих структурно функциональных особенностей открытая система должна обладать термодинамическим равновесием. При этом в такое равновесие может быть активным, т.е. в системе происходят процессы обмена с окружающей средой некоторыми функциональными элементами. Для таких активных систем обмен должен приводить к результату, сохраняющему или улучшающим некоторые параметры: энергию, температуру, давление, вещество и т.д. Поэтому нужно считать все биологические системы, в которых комплексная энтропия может оставаться постоянной или уменьшаться активными. Такие системы обладают отрицательной энтропией – негэнтропией. И именно это является главным отличием физических и биологических систем. Для понимания принципов, на основе которых биосистемы способны сохранять или понижать внутреннюю энтропию необходимо формализация алгоритма модели таких систем. Построить такой алгоритм исходя из анализа работы простейших из существующих биосистем достаточно проблематично. Все известные сейчас биосистемы в результате эволюции приобрели множество различных молекулярных систем прямо или косвенно участвующих в реализации базовых функций. Построение модели даже наипростейших из них без существенных упрощений не представляется реальным в силу “проклятия размерностей”. В данной работе мы попытались развить попытку моделирования первичных негэнтропийных систем – логических биоагентов. По сути, эта модель должна описывать конечную стадию абиогенеза – т.е. формирование активных систем способных выполнять базовые функции живого, сохранение или увеличение негэнтропии. Мы не рассматриваем конкретные молекулярные ансамбли, участвующие именно в возникновении (зарождение) жизни. Множественные работы,

нацеленные на решение вопроса из каких конкретно молекул, могли формироваться живые системы не привели пока к однозначным выводам. В данной работе мы попытались моделировать логику, информатику живых объектов и, соответственно их организацию из молекул, выполняющих логические функции. Александр Иванович Опарин обозначал биосистемы как “индивидуальные фазово-обособленные системы, способные взаимодействовать с окружающей средой [9] Для таких систем приписываются основные свойства [10–15]. Исходя из таких информационных свойств разрабатываются модели биосистем.

Основные принципы моделирования биосистем

Разрабатываемая модель молекулярных биоагентов, не имеющих генома, базируется на следующих принципах:

1. Негэнтропия: в контексте моделирования биосистем – уровень энергоинформационной неопределенности.
2. Агенты: негэнтропийная система в виде набора агентов, которые с определенной вероятностью взаимодействуют друг с другом и окружающей средой в соответствии с их молекулярной структурой.
3. Эволюция: постепенное накопления изменений в молекулярной структуре агентов (случайные или вызванные воздействием окружающей средой).
4. Обучение: агенты могут “обучаться” на основе отбора окружающей средой.
5. Модели биосистем позволяют анализировать поведение системы в контролируемой среде.

Описание модели

Модель биоагента построена на представлении об живых системах как о логических цифровых, не бинарных молекулярных машинах. Основная функция таких систем – сохранение гомеостаза (негэнтропии). Возникновении первичных биосистем могло происходить самосборкой в случае присутствия в среде молекул, могущих выполнять логические операции и взаимодействующих с достаточно значимой вероятностью, нескольких взаимосвязанных процессов (факторов) и энергоинформационного градиента. При таких условиях могли флуктуационно формироваться предельно простые молекулярные конструкции, обладающие определенной логикой взаимодействия с окружающей средой и «запасом», избытком упорядоченности – негэнтропии N .

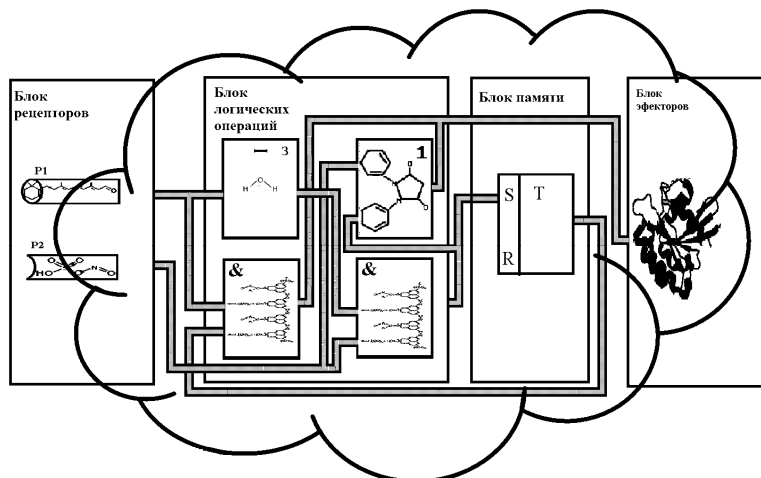


Рис. 1. Схема первичного коацервата и логические операции выполняемые такой элементарной молекулярной негэнтропийной прогностической информационной машиной. P1 и P2 – рецепторы внешних сигналов. 3. 1–1 элемент задержки сигнала. & – элемент, выполняющий логическую операцию «И». 1 – элемент, выполняющий логическую операцию «ИЛИ», Т – элемент памяти (триггер). +Э – эффектор

Молекулярные системы таким образом приобретали, преимущества в виде дополнительной устойчивости. Устойчивость к факторам среды понижающим негэнтропию системы возникала в силу совершения молекулярной системой логических операций, основанных на соотносении значимых (влияющих на начальный запас негэнтропии) и предшествующих значимым детектируемым системой “информационных” процессах. После первичного взаимодействия с комплексом детектируемых факторов среды запоминалась ассоциация значимого и информационного фактора. Воздействие процессов, происходящих в окружающей энтропийной среде, приводило к уменьшению первично возникшей упорядоченности, уменьшая суммарную негэнтропия: $\Sigma N = N - \Delta N + I$. Но, при этом, в системе фиксировалась информация I о взаимосвязях действующих факторов среды. Т.е. возникали элементы “памяти”.

На основе “памяти” появляется возможность прогнозирования появления внешних факторов. Такое логическое прогнозирование возможно в среде, в которой присутствуют ассоциативно и/или причинно-следственно связанные факторы. Получаемая при взаимодействии со средой и записываемая информация используется для прогнозирования факторов среды,

ведущих к уменьшению суммарной негэнтропии. При этом приобретенная информация в дальнейшем “использовалась” молекулярной системой для уменьшения энтропии с некоторым коэффициентом K_i , создавая дополнительный запас негэнтропии: $I \rightarrow N \cdot K_i$.

Это позволяет системе избегать потерь и/или приобретать дополнительно определенное количество комплексной негэнтропии. Т.е. при взаимодействии с таким комплексом факторов первый играет роль сигнала и вызывает реакцию, опережающую внешнее событие. Система реагирует по ранее зафиксированному алгоритму, корректируя свое состояние или приобретая дополнительное количество негэнтропии в соответствии с коэффициентом K_i .

Для биосистемы важно отличать ассоциативные и причинно-следственно связанные явления в окружающей среде. При действии ассоциированных, но не причинно-следственно связанных процессов агент может терять определенное количество негэнтропии. Поэтому в алгоритме модели агента предусмотрен задаваемый порог количества таких ошибок. После достижения, которого из памяти агента исключается информация о взаимосвязи двух этих факторов.

Результаты моделирования биоагентов

Разработаны модели негэнтропийных биоагентов. Программная реализация алгоритма позволила при интуитивном подборе параметров получить варианты агентов, обладающие способностью сохранения своих параметров неограниченное количество циклов взаимодействия со средой (рис. 3).

Для визуализации результатов программы на данном этапе был выбран линейный график – это простейший способ визуализации данных. У такого подхода к отображению показателей есть много полезных качеств.

Данный график состоит из трёх показателей:

1. Количество живых Агентов
2. Средняя продолжительность жизни
3. Максимальная продолжительность жизни

Как видно из графика (рис. 2), отбор агентов средой уменьшает численность популяции. При этом, средняя и максимальная продолжительность жизни, соответственно изменяются.

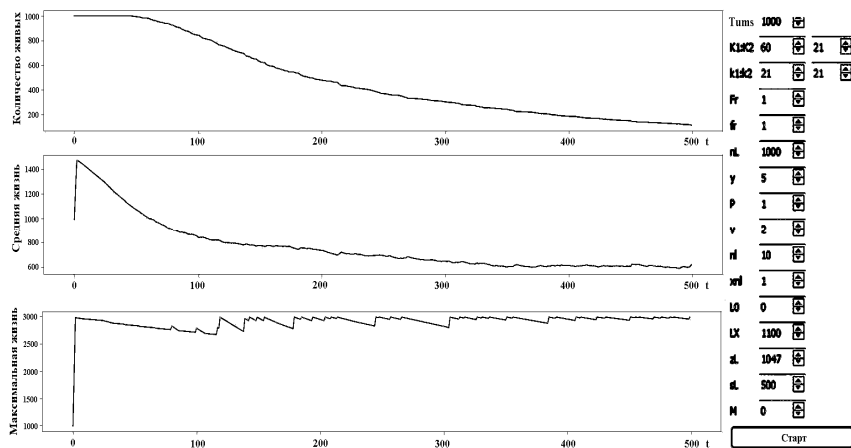


Рис. 2. Пример динамики параметров группы агентов, взаимодействующих со средой

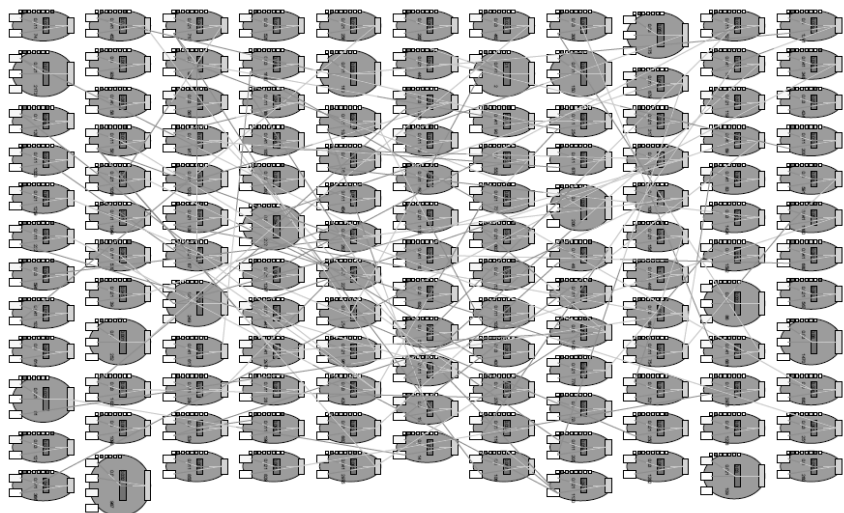


Рис. 3. Результат варианта эволюции популяции агентов после 490 взаимодействий со средой

Сигнал- f приходит перед сигналом F , что даёт Агенту возможность “предсказать” развитие процесса F и избегать большей потери негэнтропии. Кроме того, на графике можно заметить, что средняя продолжительность жизни и максимальная продолжительность жизни имеют разные тенденции изменения во времени. Средняя продолжительность жизни имеет более плавный характер изменения, тогда как максимальная продолжительность жизни имеет более резкие скачки. Это может быть связано с тем, что у некоторых Агентов определились такие вероятностные значения параметров, которые позволяют им жить дольше.

Для ускорения анализа особенностей алгоритма агента был сделан переход к моделированию их популяции.

Представлены только данные об одной реализации популяции Агентов. Моделировалось объединение (слияние) агентов. Результаты могут отличаться в зависимости от заданных параметров среды и агента. Поэтому для более точного предсказания динамики популяции проводились эксперименты с различными значениями параметров.

Заключение

Из агентов устойчивых в заданной среде можно конструировать более сложные системы. Это могут быть аналоги нейросетей составленные на базе таких агентов. Моделировать сети подобные сигнальным сетям в клетках также представляется продуктивным. Можно предположить, что в составе внутриклеточных сигнальных сетей, кроме известных (рецепции, усиления сигналов и включения эффекторов) присутствуют еще и удачно сформированные в процессе длительной эволюции молекулярные системы обеспечивающие выявления и запоминания воздействий, вызывающие опережающие реакции системы и стабилизирующие, таким образом, гомеостаз. Выявление и инструментальное моделирование многоуровневых негэнтропийных молекулярных конструкций является задачей дальнейшего детального анализа внутриклеточной молекулярной сети, на примере относительно простых сетей в синапсах [16 и др.]. Опираясь на модель био-агента основанную на физических представлениях о базовых функциях биосистем и учитывая преемственность, эмерджентность организации таких систем, в дальнейшем представляется возможным более детально моделировать их эволюцию при отборе внешней средой с заданными параметрами. При этом сложность имитационного моделирования можно наращивать до пределов программных и аппаратных ограничений. Подход к моделированию возникновения и эволюции негэнтропийных информационных систем на базе модели простейшего негэнтропийного агента представляется наиболее перспективным.

Данная работа является развитием исследований [16, 17].

Список литературы

1. Zane R. Thornburg, David M. Bianchi, Troy A. Brier and others, Fundamental behaviors emerge from simulations of a living minimal cell // Cell. V. 185, I. 2, P. 345–360. e28, Published in issue: January 20, 2022, doi.10.1016/j.cell.2021.12.025.
2. Abhishek Singharoy, Christopher Maffeo, Karelia H. Delgado-Magnero and others, Atoms to Phenotypes: Molecular Design Principles of Cellular Energy Metabolism // Cell. V. 179, I. 5, P. 1098–1111. e23, Published in issue: November 14, 2019, doi.10.1016/j.cell.2019.10.021.
3. Anokhin P.K. Anticipatory reflection of reality // Issues. philosophy. 1962. N 6. P. 97–109.
4. Джон Бернал. Возникновение жизни. Москва : Мир 1969. С. 391.
5. Анохин П.К. Основные вопросы теории функциональных систем. Философские аспекты теории функциональных систем // Москва : Наука. – 1978. – С. 49–106.
6. Anokhin P.K. Anticipatory reflection of reality // Vopr. Philosophy. – 1962. – N 6. P. 97–109.
7. Сеченов И.М. Собрание сочинений И.М. Сеченова // ARCHIVE PUBLICA. – 2022. – 241 с.
8. Tsitolovsky L. Consciousness, endogenous generation of goals and homeostasis // International Journal of General Systems. – 2015. – V. 44, N 6. P. 655–666.
9. Опарин А.И. Жизнь, её природа, происхождение и развитие. 2-е изд., дополненное. Москва : Наука, 1968. 173 с.
10. Marijuán P.C., Navarro J., del Moral R. How the living is in the world: An inquiry into the informational choreographies of life // Prog Biophys Mol Biol. 2015. V. 119, N 3. P. 469–480.
11. Nakajima T. Biologically inspired information theory: Adaptation through construction of external reality models by living systems // Prog Biophys Mol Biol. 2015. V. 119, N 3. P. 634–648.
12. Goranson T., Cardier B., Devlin K. Pragmatic phenomenological types // Prog Biophys Mol Biol. 2015. V. 119, N 3. P. 420–436.
13. Liao J.C., Boscolo R., Yang Y.L., Tran L.M., Sabatti C., Roychowdhury V.P. Network component analysis: reconstruction of regulatory signals in biological systems // Proc Natl Acad Sci U S A. 2003. V. 100, N 26. P. 15522–15527.
14. Schrödinger E. What is life? The physical aspect of the living cell // Cambridge: University Press. 1944. 214 p.
15. Aristov V.V., Buchelnikov A.S., Nechipurenko Yury D. The use of the statistical entropy in some new approaches for the description of biosystems, entropy (Basel) // 2022. V. 24, N 2. P. 172.
16. Proskura A.L., Vechkapova S.O., Ratushnyak A.S. Dopamine and hippocampal synaptic plasticity // Procedia Computer Science. 2020. V. 169. P. 668–672.
17. Ratushnyak A.S., Zapara T.A., Proskura A.L., Sklyarov A.N., Sorokoumov E.D. Analysis of appearances, formation and evolution of biological functional systems // Studies in computational intelligence. 2022. V. 1064. P. 231–237.

А.Л. ПРОСКУРА, С.О. ВЕЧКАПОВА, А.С. РАТУШНЯК

Федеральный исследовательский центр
информационных и вычислительных технологий, Новосибирск
annleop@mail.ru

РЕГУЛЯТОРНЫЕ СЕТИ ЭНДОЦИТОЗА ГЛУТАМАТНЫХ РЕЦЕПТОРОВ В ИНТЕРАКТОМЕ ДЕНДРИТНОГО ШИПИКА ГИППОКАМПА*

В работе рассматриваются регуляторные пути эндоцитоза под контролем медиаторных и не медиаторных рецепторов в интерактоме дендритного шипика пирамидного нейрона поля СА1 гиппокампа. Предложена реконструкция альтернативных путей контроля синаптических глутаматных рецепторов АМПА типа, вовлекаемых в процессы модуляции синаптической нейротрансмиссии.

Ключевые слова: *гиппокамп, дендритный шипик, синаптическая пластичность, глутаматные рецепторы, инсулин.*

Введение

Синапсы представляют структурно упорядоченные и высоко пластичные на уровне межмолекулярных взаимодействий специализированные контакты между нейронами. Дендритный шипик формирует постсинаптическую часть пирамидных нейронов высших отделов мозга. На его плазматической мембране располагается целый спектр рецепторов различной природы, которые организованы в многоуровневую систему контроля глутаматергической нейротрансмиссии [1, 2].

Долговременная потенция (ДВП) — длительное увеличение эффективности синаптической передачи, возникающее после интенсивного и непродолжительного выброса нейротрансмиттера, например, в результате высокочастотной стимуляции афферентных входов [3]. Преобладающее большинство нейрофизиологических исследований направлено на изучение электрогенной активности. При этом их взаимосвязь с молекулярными процессами остаются недостаточно изученными. Считается, что запуск сигнальных путей в интерактоме дендритного шипика определяет изменение эффективности синапса и ее длительное поддержание [4].

* Работа выполнена в рамках государственного задания Минобрнауки России для Федерального исследовательского центра информационных и вычислительных технологий.

Существуют ионотропные и метаботропные глутаматные рецепторы. На постсинаптической части синапсов представлены метаботропные рецепторы mGluR1 и mGluR5 типов, они запускают сигнальные пути через Gq-белки и модулируют уровень продукции вторичных мессенджеров [5,6]. Ионотропные глутаматные рецепторы, а именно НМДА (N-метил-D-аспаратат) (НМДАР) и АМПА (α -амино-3-гидрокси-5-метил-4-изоксазолпропионовая кислота) (АМПАР) типов являются ключевыми для ДВП. Именно изменение плотности синаптических АМПАР после индуцируемой активацией НМДАР ДВП представляет ключевой процесс, который опосредует эффективность синапсов [4, 7].

Существует два пути регулирования плотности АМПА рецепторов. Конститутивный кругооборот протекает постоянно и направлен на поддержание существующей плотности рецепторов на синаптической мембране. Другой путь запускается после индукции ДВП и направлен на изменение плотности АМПАР [8].

Платформа GeneNet (РОСПАТЕНТ № 990006 от 15/02/1999) [9,10] является инструментом для визуального моделирования и разработки моделей процессов со сложной организацией. В основе визуального моделирования лежит графическая нотация, которая позволяет описать различные аспекты структуры и функционирования рассматриваемой системы в виде графической диаграммы с учетом разнообразия уровней детализации и формализации. На основе опубликованных экспериментальных данных мы реконструировали межмолекулярные взаимодействия в интерактоме дендритного шипика пирамидного нейрона поля CA1 гиппокампа [11,45]. Проведен анализ регуляторных контуров, обеспечивающих изменение плотности синаптических АМПАР после индукции ДВП, а именно, процесса эндцитоза рецепторов.

Динамика АМПА рецепторов в течение долговременной потенциации

1.1. Визуализационная модель синаптической пластичности

Глутаматные рецепторы формируют макрокомплексы с белками сигнальных путей ряда протеинкиназ и фосфатаз. Динамика плотности АМПАР в зоне синаптического контакта определяет синаптическую эффективность, зависит напрямую от их специфичного фосфорилирования/дефосфорилирования и находится под контролем сигнальных каскадов, запускаемых появлением после индукции ДВП вторичных мессенджеров [11, 12].

Варианты гетеромерных АМПАР во взрослом гиппокампе: субъединицы GluR2, GluR3 - GluR2/3-АМПАР, GluR1 – как гомо-, так и гетеротетрамеры с GluR2 - GluR1-АМПАР [13]. Рецепторы организованы в макрокомплексы с белками постсинаптического уплотнения (ПСУ) [2, 11], а именно скаффолд-белками, например PSD95 (Postsynaptic density protein) и

адаптерными белками [14–21]. В частности, белки семейства Shank (SH3 and multiple ankyrin repeat domains protein) обеспечивают посадку АМПАР в зоне синаптического контакта [22–24] (рис. 1).

После индукции ДВП происходит кратковременный интенсивный вход ионов кальция через НМДАР, что запускает сети сигнальных путей ряда протеинкиназ и фосфатаз [11]. GluR2/3-АМПАР наиболее представлены на синапсах взрослого гиппокампа, их плотность не статична, поддерживается конститутивным рециклированием и зависит от предшествующей активности, метаболического состояния и может модулироваться разнообразными агентами [4, 25]. Они обеспечивают базовую активность синапса, а также деполяризацию плазматической мембраны (до ~ -50 мВ), что приводит к активации НМДАР и индукции ДВП.

Ключевыми событиями после индукции ДВП являются выведение GluR2/3-АМПАР из синапсов, их эндоцитоз и введение дополнительного пула GluR1-АМПАР и их закрепление с белками ПСУ - SAP97 и Shank, а также ряд модификаций [4, 7, 26].

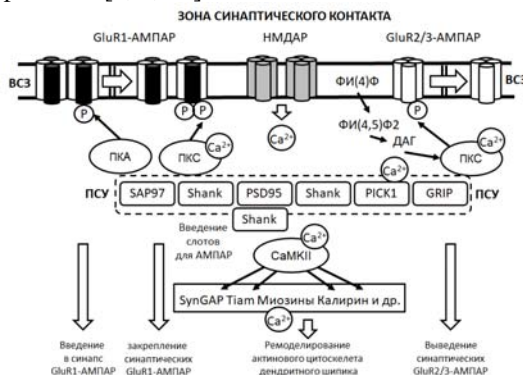


Рис. 1. Схематичное представление ключевых этапов изменения плотности АМПАР на синапсах после индукции ДВП:

ПСУ – постсинаптическое уплотнение; ВСЗ – внесинаптическая зона; GluR1-АМПАР – содержат субъединицу 1 типа; GluR2/3 – АМПАР - субъединицы 2 и 3 типов; НМДАР – НМДА рецепторы; ФИ(4)Ф – фосфатидилинозит-4-монофосфат; ФИ(4,5)Ф2 – фосфатидилинозит-4,5-дифосфат; PKA – протеинкиназа А; PKC – протеинкиназа С; CaMKII – кальций-кальмодулин регулируемая протеинкиназа. Р – фосфатная группа (модификация АМПАР). Ca^{2+} – входящие ионы кальция.

Выведение синаптических GluR2/3-АМПАР: комплексы PICK (Protein Interacting with C Kinase - 1), GRIP (Glutamate-receptor-interacting protein) и PKC. Введение на синапсы GluR1-АМПАР и их закрепление: протеинкиназы PKA и PKC, а также белки ПСУ – SAP97 (Synapse-associated protein 97), Shank (SH3 and multiple ankyrin repeat domains protein), PSD95 (Postsynaptic density protein)

Развитие ДВП сопровождается корректировкой новой плотности синаптических АМПАР, а долговременное поддержание – заменой GluR1-АМПАР на GluR2/3-АМПАР, чья плотность будет поддерживаться путем конститутивного рециклирования [4, 11] и зависеть от модулирующих воздействий, таких как новая структура макрокомплексов рецепторов, метаболический статус, наличие повреждающих факторов и пр.

Регулирование эндоцитоза глутаматных рецепторов

Выводимые из синапсов АМПАР подвергаются быстрому эндоцитозу, после чего происходит их сортировка в эндосомальных компартментах, после чего они могут либо вернуться на плазматическую мембрану в составе ранних эндосом, либо быть выведены в путь деградации [27].

Активация НМДАР способствует появлению на плазматической мембране пула вторичных мессенджеров, в частности, фосфатидилинозитол-4,5-дифосфата (ФИФ2) и фосфатидилинозитол-(3,4,5)-трифосфата (ФИФ3).

Вход ионов кальция через канал НМДАР в течение индукции ДВП активирует протеинфосфатазы кальцинейрин и PP1 (белковую фосфатазу 1 (protein phosphatase 1)), которые дефосфорилируют и активируют фосфатидилинозитол-4-фосфат 5 киназу (ФИ5К), которая обеспечивает биосинтез ФИФ2 из фосфотидилинозитол-4 фосфата (ФИ4Ф), путем присоединения фосфатной группы в пятой позиции его инозитольного кольца. ФИ5К выступает главным продуцентом ФИФ2 в мозге [28].

ФИФ2 является точкой формирования эндосомы в процессе клатрин-зависимого эндоцитоза рецепторов, в том числе и АМПАР: ФИ(4)Ф -> ФИФ2 -> формирование покрытой клатрином эндосомы -> эндоцитоз.

ФИФ2 являются также субстратом для фосфолипазы C (ФЛС), которая гидролизует его до диацилглицерола (ДАГ) и инозитол-3-фосфата (И(3)Ф), которые также относятся к вторичным мессенджерам. ДАГ активирует протеинкиназу C (ПКС), которая фосфорилирует GluR2 субъединицу АМПАР [29], что является одним из ключевых этапов для выведения GluR2/3-АМПАР из синапсов и их эндоцитоза после активации НМДАР (рис.1).

Метаботропные рецепторы mGluR1 и mGluR5 вовлекаются в процессы регулирования синаптической пластичности, обеспечивая активацию ФЛС [6], и, таким образом, способствуя интенсификации эндоцитоза АМПАР.

Проведенный анализ межбелковых регуляторных сетей в интерактоме дендритного шипика пирамидного нейрона поля CA1 гиппокампа, реконструированной на платформе GeneNet, позволил нам выделить альтернативные пути регуляции эндоцитоза АМПАР.

Мы провели реконструкцию путей активации ФИ5К в течение запускаемой НМДАР синаптической пластичности.

ARF6 (ADP-ribosylation factor 6) – белок, связывающий ГТФ. В отличие от членов других двух классов, Arf6 не регулирует сборку покрытия транспортными везикулами на мембранах вакуолярной системы. Установлена активная роль этого белка в процессах рециркуляции везикул между эндосомами и плазматической мембраной [30], а также в процессах структурной реорганизации плазматической мембраны [31].

Было обнаружено, что ФИ5К является эффектором в межмолекулярных реакциях, в которых ARF6 запускает рекрутирование набора различных, но взаимодействующих сигнальных молекул в пределах сайтов активных цитоскелетных или мембранных перестроек, включая в пути эндоцитоза и рециркулирования рецепторов плазматической мембраны [30–32]. Активаторы ARF6, такие как EFA6A, BRAG, хорошо представлены в области ПСУ шипика и входят в макрокомплексы ионотропных глутаматных рецепторов, активируясь после индукции ДВП [33, 34]. BRAG2 включен в регулируемый Arf6 эндоцитоз GluR2/3-AMPA, напрямую взаимодействуя с субъединицей GluR2, что в итоге приводит к снижению плотности синаптических GluR2/3-AMPA и будет препятствовать развитию ДВП, снижая эффективность синапса [35]. Scholz с соавторами предлагает рассматривать AMPA-опосредованную активацию Arf6 через BRAG2 как критический механизм направленного эндоцитоза синаптических рецепторов, который динамически регулируется как связыванием лиганда AMPA, так и состоянием фосфорилирования тирозина 876 (Y876) в GluR2, обозначая как стратегию двойного ключа для жесткого контроля плотности синаптических GluR2-AMPA [34].

Таким образом, мы предлагаем альтернативный путь интенсификации эндоцитоза: ARF6 → ФИ5К → накопление пула ФИФ2 (рис. 2).

Инсулин вырабатывается поджелудочной железой и посредством насыщенного транспорта поступает в мозг, где связывается со своими рецепторами (ИР). В периферических тканях он, главным образом, обеспечивает метаболизм глюкозы и липидов [36, 37]. В нейронах ИР участвует в аноксигенных реакциях, формировании памяти и выживании нейронов [36, 38–40]. Нарушение передачи сигналов инсулина вызывает окислительное повреждение нейронов и подвергает мозг риску развития нейродегенерации [41].

На основе литературных данных проведена реконструкция сигнального пути инсулина в интерактоме дендритного шипика пирамидного нейрона поля CA1 гиппокампа, задействованного в контроле эндоцитоза GluR2-AMPA. Ключевым этапом здесь выступает накопление ФИФ3 под контролем протеинкиназы фосфатидилинозитол-3 киназы (ФИЗК), которая фосфорилирует инозитольное кольцо ФИФ2 в третьей позиции: ФИФ2 →

ФИФ3. ФИФ3 являются сайтами закрепления ряда белков, в частности, ФЗК – фосфоинозитид-зависимой протеинкиназы 1 ((PKD1 (Phosphoinositide-dependent kinase-1)), которая, помимо прочего, пролонгирует активность ПКС [42].

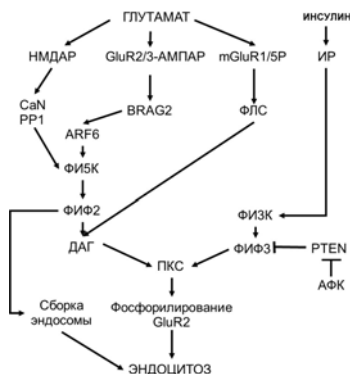


Рис. 2. Схематичное представление ключевых этапов изменения плотности АМПАР на синапсах после индукции ДВП:

GluR2/3-АМПАР – АМПАР с субъединицами 2 и 3 типов; НМДАР – НМДА рецепторы; mGluR1/5P – метаболотропные глутаматные рецепторы; ИР – рецептор инсулина; CaN – кальцинейрин; PP1 – белковая фосфатаза 1; ARF6 - ADP-ribosylation factor 6; BRAG2 - IQ motif and SEC7 domain-containing protein 1, активатор для ARF6; ФЛС – фосфолипаза C; ФИ5К – фосфатидилинозитол-5 киназа; ФИЗК - фосфатидилинозитол-3 киназа; PTEN - протеинфосфатаза с двойной субстратной специфичностью, продукт гена PTEN (сокр. от англ. phosphatase and tensin homolog deleted on chromosome 10); АФК – активные формы кислорода; ФИФ2 – фосфатидилинозитол-4,5-дифосфат; ФИФ3 – фосфатидилинозитол-3,4,5-трифосфат. $\longrightarrow$ активирующее влияние; —| - ингибирующее влияние

Инсулин напрямую активирует ФИЗК [36, 43], но при этом важным фактором является ингибирование PTEN – протеинфосфатазы с двойной субстратной специфичностью, продукта гена PTEN (сокр. от англ. phosphatase and tensin homolog deleted on chromosome 10). Это может достигаться, в частности, ингибированием активного центра перекисью водорода и/или АФК [44]. PTEN дефосфорилирует фосфоинозитиды в позиции 3 инозитольного кольца (реакция ФИФ3 $\rightarrow$ ФИФ2).

Регуляторный контур эндоцитоза GluR2-АМПАР: Инсулин $\rightarrow$ ФИЗК $\rightarrow$ ФИФ3 $\rightarrow$ ФЗК $\rightarrow$ пролонгация активности ПКС $\rightarrow$ фосфорилирование GluR2 $\rightarrow$ эндоцитоз GluR2-АМПАР.

Таким образом, в интерактоме дендритного шипика можно наблюдать пересечение сигнальных путей медиаторных и не медиаторных рецепторов для контроля эндоцитоза АМПАР как важной составляющей регуляции синаптической пластичности в гиппокампе (рис.2).

Заключение

Эндоцитоз AMPAR является ключевым событием в регуляции синаптической пластичности. Существуют альтернативные пути запуска эндоцитоза синаптических рецепторов, которые находятся под контролем сложной регуляторной сети в интерактоме дендритного шипика.

Список литературы

1. Collins S.A., Husi H., Yu L., Brandon J.M., Anderson C.N.G., Blackstock W.P., Choudhary J.S., Grant S.G.N. Molecular characterization and comparison of the components and multiprotein complexes in the postsynaptic proteome // *J. Neurochem.* 2007. V. 97. P. 16–23.
2. Proskura A.L., Ratushnyak A.S., Vechkapova S.O., Zapara T.A. Synapse as a multi-component and multi-level information system. *Stud. Comput. Intell.* 2018. V. 736. P. 186–192.
3. Bliss T.V., Collingridge G.L. A synaptic model of memory: long-term potentiation in the hippocampus // *Nature.* 1993. V. 361, N. 6407. P. 31–39.
4. Newpher T.M., Ehlers M.D. Glutamate receptor dynamics in dendritic microdomains. *Neuron.* 2008. V. 58, N. 4. P. 472–497.
5. Lee H.G., Zhu X., O'Neill M.J., Webber K., Casadesus G., Marlatt M., Raina A.K., Perry G., Smith M.A. The role of metabotropic glutamate receptors in Alzheimer's disease // *Acta. Neurobiol. Exp. (Wars)* 2004. V. 64, N. 1. P. 89–98.
6. Neyman S., Manahan-Vaughan D. Metabotropic glutamate receptor 1 (mGluR1) and 5 (mGluR5) regulate late phases of LTP and LTD in the hippocampal CA1 region in vitro. *Eur J Neurosci.* 2008 Mar;27(6):1345–52.
7. Tanaka H., Hirano T. Visualization of subunit-specific delivery of glutamate receptors to postsynaptic membrane during hippocampal long-term potentiation // *Cell. Rep.* 2012. V. 1, N. 4. P. 291–298. DOI: 10.1016/j.celrep.2012.02.004
8. Shepherd J.D., Huganir R.L. The cell biology of synaptic plasticity: AMPA receptor trafficking // *Annu. Rev. Cell. Dev. Biol.* 2007. N. 23. P. 613–643.
9. Ananko E.A., Podkolodny N.L., Stepanenko I.L., Podkolodnaya O.A., Rasskazov D.A., Miginsky D.S., Likhoshvai V.A., Ratushny A.V., Podkolodnaya N.N., Kolchanov N.A. GeneNet in 2005 // *Nucl. Acids Res.* 2005. V. 33. P. 425–427.
10. Колчанов Н.А., Воевода М.И., Кузнецова Т.Н., Мордвинов В.А., Игнатьева Е.В. Генные сети липидного метаболизма // *Бюлл. СОРАМН.* 2006. Т. 2. С. 29–42.
11. Proskura A.L., Malakhin I.A., Turnaev I.I., Suslov V.V., Zapara T.A., Ratushnyak A.S. Intermolecular interactions in the functional systems of neuron // *Vavilov. Zh. Genet. Selekt.* 2013;17:620–628.

12. Proskura A.L., Ratushnyak A.S., Zapara T.A. The protein-protein interaction networks of dendritic spines in the early phase of long-term potentiation // *J. Comput. Sci. Syst. Biol.* 2014. V. 7. P. 40–44.
13. Wenthold R.J., Petralia R.S., Blahos J., Niedzelski A.S. Evidence for multiple AMPA receptor complexes in hippocampal CA1/CA2 neurons // *J. Neurosci.* 1996. V. 16, N 6. P. 1982–1989.
14. Chen X., Winters C., Azzam R., Li X., Galbraith J.A., Leapman R.D., Reese T.S. Organization of the core structure of the postsynaptic density // *Proc. Natl. Acad. Sci. U S A.* 2008. V. 105, N 11. P. 4453–4458.
15. Manser E., Loo T.H., Koh C.G., Zhao Z.S., Chen X.Q., Tan L., Tan I., Leung T., Lim L. PAK kinases are directly coupled to the PIX family of nucleotide exchange factors // *Mol Cell.* 1998. V. 1, N 2. P. 183–192.
16. Uruno T., Liu J., Li Y., Smith N., Zhan X. Sequential interaction of actin-related proteins 2 and 3 (Arp2/3) complex with neural Wiscott-Aldrich syndrome protein (N-WASP) and cortactin during branched actin filament network formation // *J. Biol. Chem.* 2003. V. 278, N 28. P. 26086–26093.
17. Neuhoff H., Sasso?-Pognetto M., Panzanelli P., Maas C., Witke W., Kneussel M. The actin-binding protein profilin I is localized at synaptic sites in an activity-regulated manner // *Eur. J. Neurosci.* 2005. V. 21, N 1. P. 15–25.
18. Schubert V., Dotti C.G. Transmitting on actin: synaptic control of dendritic architecture // *J. Cell. Sci.* 2007. V. 120, N 2. P. 205–212.
19. Rex C.S., Chen L.Y., Sharma A., Liu J., Babayan A.H., Gall C.M., Lynch G. Different Rho GTPase-dependent signaling pathways initiate sequential steps in the consolidation of long-term potentiation // *J. Cell. Biol.* 2009. V. 186. P. 85–97.
20. Iida J., Ishizaki H., Okamoto-Tanaka M., Kawata A., Sumita K., Ohgake S., Sato Y., Yorifuji H., Nukina N., Ohashi K., Mizuno K., Tsutsumi T., Mizoguchi A., Miyoshi J., Takai Y., Hata Y. Synaptic scaffolding molecule alpha is a scaffold to mediate N-methyl-D-aspartate receptor-dependent RhoA activation in dendrites // *Mol. Cell. Biol.* 2007. V. 27, N 12. P. 4388–4405.
21. Honkura N., Matsuzaki M., Noguchi J., Ellis-Davies G.C., Kasai H. The subspine organization of actin fibers regulates the structure and plasticity of dendritic spines // *Neuron.* 2008. V. 57. P. 719–729.
22. Peng J., Kim M.J., Cheng D., Duong D.M., Gygi S.P., Sheng M. Semiquantitative proteomic analysis of rat forebrain postsynaptic density fractions by mass spectrometry // *J. Biol. Chem.* 2004. V. 279, N 20. P. 21003–21011.
23. Chen X., Vinade L., Leapman R.D., Petersen J.D., Nakagawa T., Phillips T.M., Sheng M., Reese T.S. Mass of the postsynaptic density and enumeration of three key molecules // *Proc Natl Acad Sci U S A.* 2005. V. 102, N 32. P. 11551–11556.
24. Cheng D., Hoogenraad C.C., Rush J., Ramm E., Schlager M.A., Duong D.M., Xu P., Wijayawardana S.R., Hanfelt J., Nakagawa T., Sheng M., Peng J. Relative and absolute quantification of postsynaptic density proteome isolated from rat forebrain and cerebellum // *Mol Cell Proteomics.* 2006. V. 5, N 6. P. 1158–1170.
25. Greger I.H., Esteban J.A. AMPA receptor biogenesis and trafficking // *Curr. Opin. Neurobiol.* 2007. V. 17, N 3. P. 289–297.

26. Nowak L., Bregestovski P., Ascher P., Herbet A., Prochiantz A. Magnesium gates glutamate-activated channels in mouse central neurones // *Nature*. 1984. V. 307, N 5950. P. 462–465.
27. van der Sluijs P., Hoogenraad C.C. New insights in endosomal dynamics and AMPA receptor trafficking // *Semin Cell Dev Biol*. 2011. V. 22, N 5. P. 499–505. doi: 10.1016/j.semcdb.2011.06.008.
28. Unoki T., Matsuda S., Kakegawa W., Van N.T., Kohda K., Suzuki A., Funakoshi Y., Hasegawa H., Yuzaki M., Kanaho Y. NMDA receptor-mediated PIP5K activation to produce PI(4,5)P? is essential for AMPA receptor endocytosis during LTD. *Neuron* // 2012. V. 73, N 1. P. 135–148. doi: 10.1016/j.neuron.2011.09.034.
29. Seidenman K.J., Steinberg J.P., Huganir R., Malinow R. Glutamate receptor subunit 2 Serine 880 phosphorylation modulates synaptic transmission and mediates plasticity in CA1 pyramidal cells // *J Neurosci*. 2003. V. 23, N 27. P. 9220–8.
30. Radhakrishna H., Donaldson J.G. ADP-ribosylation factor 6 regulates a novel plasma membrane recycling pathway // *J Cell Biol*. 1997. V. 139, N 1. P. 49–61. doi: 10.1083/jcb.139.1.49.
31. D'Souza-Schorey C., Li G., Colombo M.I., Stahl P.D. A regulatory role for ARF6 in receptor-mediated endocytosis // *Science*. 1995. V. 267, N 5201. P. 1175–1178. doi: 10.1126/science.7855600.
32. Honda A., Nogami M., Yokozeki T., Yamazaki M., Nakamura H., Watanabe H., Kawamoto K., Nakayama K., Morris A.J., Frohman M.A., Kanaho Y. Phosphatidylinositol 4-phosphate 5-kinase alpha is a downstream effector of the small G protein ARF6 in membrane ruffle formation // *Cell*. 1999. V. 99, N 5. P. 521–532. doi: 10.1016/s0092-8674(00)81540-8
33. Sakagami H., Honma T., Sukegawa J., Owada Y., Yanagisawa T., Kondo H. Soma-dendritic localization of EFA6A, a guanine nucleotide exchange factor for ADP-ribosylation factor 6, and its possible interaction with alpha-actinin in dendritic spines // *Eur J Neurosci*. 2007. V. 25, N 3. P. 618–628. doi: 10.1111/j.1460-9568.2007.05345.x.
34. Scholz R., Berberich S., Rathgeber L., Kolleker A., K?hr G., Kornau H-C. AMPA receptor signaling through BRAG2 and Arf6 critical for long-term synaptic depression // *Neuron*. 2010. V. 66, N 5. P. 768–780. doi: 10.1016/j.neuron.2010.05.003.
35. Fukaya M., Sugawara T., Hara Y., Itakura M., Watanabe M., Sakagami H. BRAG2a Mediates mGluR-Dependent AMPA Receptor Internalization at Excitatory Postsynapses through the Interaction with PSD-95 and Endophilin 3 // *J Neurosci*. 2020. V. 40, N 22. P. 4277–4296. doi: 10.1523/JNEUROSCI.1645-19.2020.
36. Помыткин И.А., Красильникова И.А., Пинелис В.Г., Каркищенко Н.Н. Инсулиновый рецептор в мозге: новая мишень в лечении центральной инсулиновой резистентности // *Биомедицина*. 2018. V. 3. С. 17–34.
37. Stanley M., Macauley S. L., Caesar E. E., Koscal L. J., Moritz W., Robinson G. O., Roh J., Keyser J., Jiang H. and Holtzman D. M. The effects of peripheral and central high insulin on brain insulin signaling and amyloid-beta in young and old APP/PS1 mice // *J. Neurosci*. 2016. V. 36. P. 11704–11715.

38. Zhao F., Siu J.J., Huang W., Askwith C., Cao L. Insulin modulates excitatory synaptic transmission and synaptic plasticity in the mouse hippocampus // *Neuroscience*. 2019. V. 411. P. 237–254.
39. Salcedo I., Tweedie D., Li Y., Greig N.H. Neuroprotective and neurotrophic actions of glucagon-like peptide-1: an emerging opportunity to treat neurodegenerative and cerebrovascular disorders: Neurological benefits of GLP-1 receptor activation // *Br J Pharmacol*. 2012. V. 166, N 5. P. 1586–1599.
40. Yu L.-Y., Pei Y. Insulin Neuroprotection and the Mechanisms // *Chin Med J (Engl)*. 2015. V. 128, N 7. P. 976–981.
41. Muller A.P., Haas C.B., Camacho-Pereira J., Brochier A.W., Gnoatto J., Zimmer E.R., de Souza D.O., Galina A., Portela L.V. Insulin prevents mitochondrial generation of H₂O₂ in rat brain // *Exp Neurol*. 2013. V. 247. P. 66–72.
42. Pearce L.R., Komander D., Alessi D.R. The nuts and bolts of AGC protein kinases // *Nat Rev Mol Cell Biol*. 2010. V. 11, N 1. P. 9–22.
43. Huang C.C. et al. The role of insulin receptor signaling in synaptic plasticity and cognitive function // *Chang Gung Med J*. 2010. V. 33, N 2. P. 115–125.
44. Pomytkin I.A. H₂O₂ signalling pathway: a possible bridge between insulin receptor and mitochondria // *Curr Neuropharmacol*. 2012. V. 10, N 4. P. 311–320.
45. <http://www.mgs.bionet.nsc.ru/mgs/gnw/genenet/viewer/AMPA.html>

**С.А. ПОЛЕВАЯ¹, И.С. ЕГОРОВ¹, А.И. ФЕДОТЧЕВ^{1,2},
Е.А. МУХИНА¹, С.Б. ПАРИН¹**

¹НИУ Нижегородский государственный университет им. Н.И. Лобачевского

²Институт биофизики клетки РАН, Пущино, Московская область
S453383@mail.ru

УПРАВЛЕНИЕ НЕЙРОПЛАСТИЧНОСТЬЮ С ПОМОЩЬЮ НЕИНВАЗИВНОЙ СЕНСОРНОЙ СТИМУЛЯЦИИ*

Представлены принципиально новые данные о диагностических, реабилитационных и исследовательских возможностях технологии резонансного НБУ с двойной обратной связью по визуальному и акустическому входам. В сравнительном аспекте рассматриваются электрофизиологические проявления регрессии нейропластичности при различных органических и функциональных нарушениях.

Ключевые слова: *нейробиоуправление, регрессия, нейропластичность, электроэнцефалография (ЭЭГ), обратная связь.*

* Данная работа выполнена при частичной поддержке РФФ, грант № 22-18-20075.

Введение

В последние десятилетия использование различных вариантов интерфейсов мозг-компьютер (ИМК = BCI) приобретает все более широкое распространение в практике. Однако, в основном, ИМК применяется для повышения качества жизни больных с существенными органическими поражениями в моторной (например, параличи различного генеза, ампутации) или сенсорной (слепота, глухота и др.) сферах [1]. Значительно меньшее распространение и, соответственно, известность имеют ИМК для решения таких сложнейших задач, как когнитивная реабилитация и коррекция высших психических функций человека. Эта технология применения ИМК, получившая название нейробиоуправления (НБУ), основывается на принципах повышения нейропластичности [2] с помощью сенсорной (зрительной, слуховой и т.д.) стимуляции, управляемой по механизму обратной связи от естественных биологических ритмов человека: прежде всего, ритмов электроэнцефалограммы (ЭЭГ) и электрокардиограммы (ЭКГ) [3, 4].

Технология НБУ существенно выигрывает по сравнению с давно известными методами прямой электрической или магнитной стимуляции структур мозга [5], потому что основывается на принципиально неинвазивной естественной сенсорной стимуляции. Это преимущество позволяет рассматривать применение НБУ не только в клиническом аспекте, но и как инструмент для изучения механизмов функционирования мозга. В частности, целью данной работы стало исследование роли процессов регрессии [6] в проявлениях нейропластичности при НБУ.

Обоснование выбора материала и методов исследования

В исследовании приняли участие более 250 испытуемых разного возраста с различными особенностями функционирования (включая патологические) мозга, в том числе: 15 сотрудников Пущинского научного центра в день deadline сдачи отчётов по грантам; 15 сотрудников IT-сферы в конце рабочего дня; 52 студента во время экзамена; 85 детей в возрасте 4–10 лет (из них 45 с клинически диагностированными нарушениями когнитивных функций: синдромом дефицита внимания с гиперактивностью (СДВГ), расстройствами аутистического спектра (РАС) и др.); 18 больных с последствиями острого нарушения мозгового кровообращения (ОНМК). Таким образом, выборка испытуемых позволила выявить как возрастную динамику формирования нейропластичности, так и характерные для различных нарушений (и стресс-индуцированных, и вызванных органическими поражениями) проявления регрессии ритмики, то есть возвращения к одному из этапов созревания [6, 7]. Каждой группе испытуемых соответствовали контрольные группы добровольцев, находившихся в комфортных условиях, в тех же возрастных диапазонах.

Технология НБУ, разработанная в Институте биофизики клетки РАН (Пушино), состояла из пяти последовательных шагов: получение ЭЭГ-сигнала мозга, его on-line обработка, выделение ключевых признаков (в первую очередь, амплитудных характеристик пиков в α - и θ -частотном диапазоне ЭЭГ), генерирование сигнала обратной связи и адаптивное обучение мозга.

Применение НБУ для реабилитации осуществлялось через последовательность протоколов, обеспечивающих снижение уровня регрессии и переход на очередной этап зрелости. В программно-аппаратном комплексе для резонансного НБУ с двойной обратной связью по визуальному и акустическому входам реализована возможность определения уровня зрелости ритмики мозга благодаря функциональной пробе "динамическая фотостимуляция". Проба "динамическая фотостимуляция" проводилась при начальном обследовании испытуемого и обеспечивалась комплексом специфических настроек световой и звуковой стимуляции: в течение 2 минут производилась запись ЭЭГ в покое (фон ДО); далее для световой стимуляции использовался режим «Сканирование» в диапазоне частот от 4 Гц (F_{min}) до 14 Гц (F_{max}) с шагом по частоте 0,1 Гц (dF) и шагом по времени 3 сек. (dT); для звуковой стимуляции использовался тета-ритм (4–8 Гц), функция обратной связи была реализована через прямую корреляцию между мощностью тета-ритма и высотой тональных модуляций; в завершение на протяжении 2 минут снова записывалась нативная ЭЭГ (фон ПОСЛЕ). Продолжительность функциональной пробы составляла 5 минут. Проба выполняется при закрытых глазах. Испытуемому предлагалось послушать «музыку мозга» и посмотреть «световое шоу» с переливами цвета, отражающее текущее состояние его мозга. Испытуемым сообщалось, что флейтоподобный звук соответствует оптимальному состоянию.

Далее режимы световой и звуковой стимуляции формировались в зависимости от зафиксированного в условиях динамической фотостимуляции уровня нейропластичности.

Результаты и их обсуждение

В результате обследования здоровых детей и взрослых испытуемых было обнаружено, что оптимальный уровень зрелости ритмика ЭЭГ формируется к 7–8 годам и в дальнейшем устойчиво поддерживается при отсутствии воздействия негативных факторов. Эффективность когнитивных функции и стрессоустойчивость связаны с устойчивостью эталонной ритмики. Эталонный уровень зрелости характеризуется следующими признаками (рис. 1): в фоне регистрируется четко выраженный альфа-ритм в диапазоне 8–13 Гц (на рис. 2 присутствует частотный пик, соответствующий

диапазону от 8 до 13 по оси X); при динамической фотостимуляции собственный альфа-ритм сохраняется; при динамической фотостимуляции проявляются ритмы с частотами стимуляции (усвоение ритма), то есть адаптивность и пластичность мозга реализуется через воспроизведение в ЭЭГ колебаний светового сигнала; при динамической фотостимуляции проявляются ритмы с частотами, кратными частоте стимуляции (мультипликация), то есть адаптивность и пластичность мозга реализуется через генерацию нового ритма; после стимуляции мощность альфа-ритма не выше, чем до стимуляции; в процессе стимуляции субъективные цветовые образы смещаются от оттенков красного к оттенкам серебристо-голубого цвета. Принципиально, что данные признаки зрелости обнаруживаются исключительно во время бодрствования и исчезают в ситуациях стресса (то есть при выраженной регрессии).

По результатам функциональной пробы «динамическая фотостимуляция» можно выделить 3 уровня регрессии ритмики мозга: глубокая, умеренная, слабая.

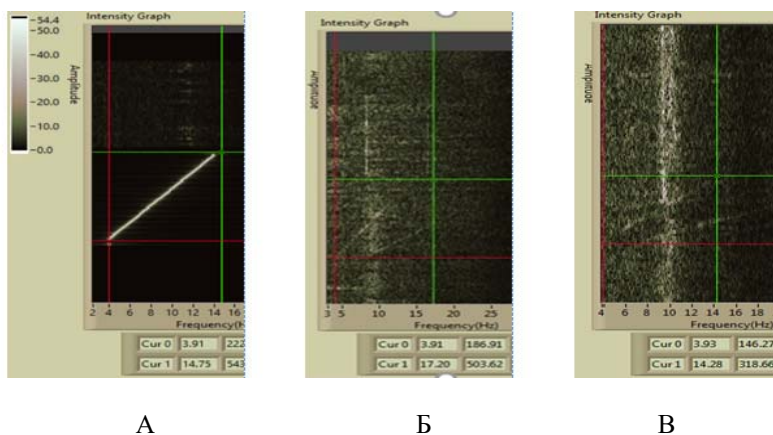


Рис. 1. Устойчивые признаки эталонной зрелости ритмики мозга. А – динамический спектр светового сигнала при динамической фотостимуляции, здесь и далее горизонтальная красная линия соответствует моменту начала сканирования, горизонтальная зеленая линия соответствует моменту окончания сканирования; Б – динамический спектр ЭЭГ ребенка 7 лет с эталонным уровнем зрелости, горизонтальная красная линия соответствует моменту начала сканирования; В – динамический спектр взрослого в оптимальном состоянии

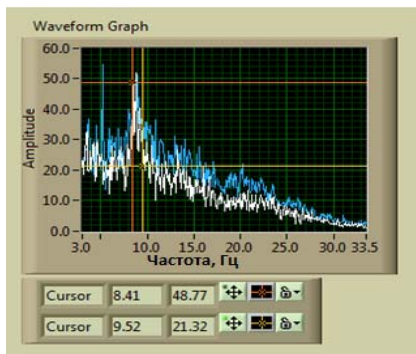


Рис. 2. Спектр ЭЭГ до (светлые линии) и после (более тёмные линии) динамической фотостимуляции при эталонной зрелости ритмики мозга

Глубокая регрессия ритмики мозга характеризуется следующими признаками (рис. 3): в фоне – отсутствуют частотные пики в альфа (8–13 Гц) и тета (4–8 Гц) – диапазонах; при динамической фотостимуляции отсутствуют эффекты усвоения ритма и мультипликации; после стимуляции мощность альфа-ритма выше, чем до стимуляции; в процессе стимуляции доминируют субъективные цветовые образы красного цвета.

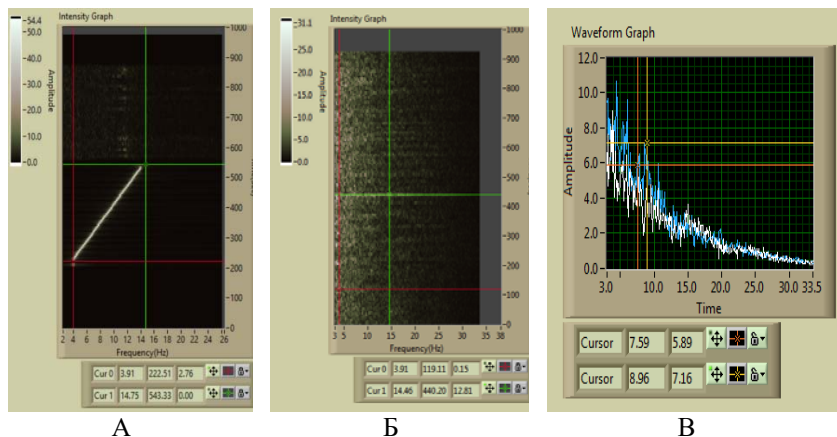


Рис. 3. Признаки глубокой регрессии ритмики мозга. А – динамический спектр светового сигнала при динамической фотостимуляции; Б – динамический спектр ЭЭГ при функциональной пробе «динамическая фотостимуляция»; В – спектр ЭЭГ до (светлые линии) и после (более тёмные линии) сеанса НБУ при глубокой регрессии

При глубокой регрессии реабилитация должна включать последовательность из 3 протоколов, обеспечивающих переход от глубокой регрессии к умеренной, от умеренной к слабой, от слабой к эталонному уровню зрелости ритмики мозга. Глубокая регрессия характерна для органических нарушений мозга: энцефалопатия, острые нарушения мозгового кровообращения, черепно-мозговые травмы и контузии мозга.

Протокол для перехода от глубокой к умеренной регрессии ритмики мозга состоит из 5 шагов: фон ДО – запись ЭЭГ в покое, 2 минуты; динамическая фотостимуляция, 5 минут; пауза – запись ЭЭГ без стимуляции, 1 минута; для световой стимуляции применяется режим «ЭЭГ-пик» в диапазоне частот от 4 Гц до 8 Гц; для звуковой стимуляции в качестве зоны интереса назначается тета-ритм (4–8 Гц), функция обратной связи реализуется через прямую корреляцию между мощностью тета-ритма и высотой тональных модуляций, 5 минут; фон ПОСЛЕ – запись ЭЭГ в покое, 2 минуты. Сеансы НБУ по этому протоколу проводятся 1 раз в день в одно и то же время суток. Стимуляция по этому протоколу прекращается только после проявления признаков умеренной регрессии ритмов мозга.

Умеренная регрессия ритмики мозга характеризуется следующими признаками (рис. 4): в фоне – отсутствует частотный пик в диапазоне альфа-ритма (8–14 Гц); при динамической фотостимуляции отсутствуют или слабо выражены эффекты усвоения ритма и мультипликации; после стимуляции мощность альфа-ритма выше, чем до стимуляции; в процессе стимуляции проявляются вспышки красного, зеленого, синего цвета.

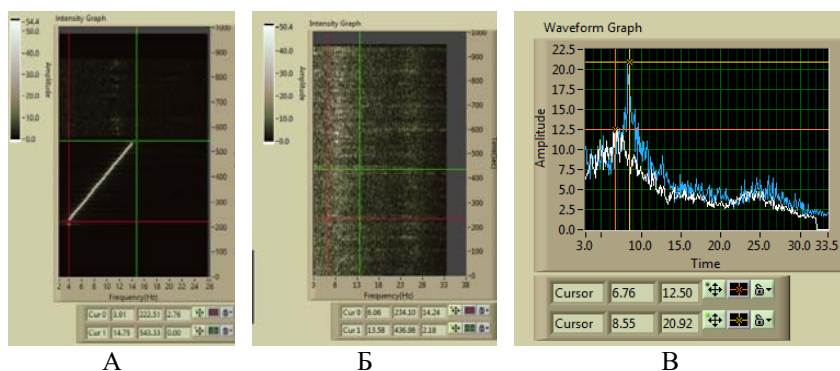


Рис. 4. Признаки умеренной регрессии ритмики мозга. А – динамический спектр светового сигнала при динамической фотостимуляции; Б – динамический спектр ЭЭГ при функциональной пробе «динамическая фотостимуляция»; В – спектр ЭЭГ до (светлые линии) и после (более тёмные линии) сеанса НБУ при умеренной регрессии

При умеренной регрессии реабилитация должна включать последовательность из 2 протоколов, обеспечивающих переход от умеренной к слабой, от слабой к эталонному уровню зрелости ритмики мозга. Умеренная регрессия характерна для функциональных нарушений мозга: посттравматических стрессовых расстройств, синдрома дефицита внимания с гиперактивностью, тревожного расстройства, депрессии.

Протокол для перехода от умеренной к слабой регрессии ритмики мозга включает следующие шаги: фон ДО – запись ЭЭГ в покое, 2 минуты; динамическая фотостимуляция, 5 минут; пауза – запись ЭЭГ без стимуляции, 1 минута; для световой стимуляции назначается режим «ЭЭГ-пик» в диапазоне частот от 8 Гц до 13 Гц; для звуковой стимуляции в качестве зоны интереса выбирается альфа-ритм (8–13 Гц), функция обратной связи реализуется через обратную корреляцию между мощностью альфа-ритма и высотой тональных модуляций, 5 минут; фон ПОСЛЕ – запись ЭЭГ в покое, 2 минуты. Сеансы НБУ по этому протоколу проводятся 1 раз в день в одно и то же время суток. Стимуляция по этому протоколу прекращается только после проявления признаков слабой регрессии ритмов мозга.

Слабая регрессия ритмики мозга характеризуется следующими признаками (рис. 5): в фоне – слабый частотный пик в диапазоне альфа-ритма (8–13 Гц); при динамической фотостимуляции отсутствует хотя бы один признак эталонной ритмики мозга: уменьшается или отсутствует альфа-пик, отсутствуют или слабо выражены эффекты усвоения ритма и мультипликации; после стимуляции мощность альфа-ритма выше, чем до стимуляции; в процессе стимуляции субъективные цветовые образы смещаются от оттенков красного к оттенкам серого цвета.

При слабой регрессии реабилитация должна включать 1 протокол, обеспечивающий переход от слабой регрессии к эталонному уровню зрелости ритмики мозга. Слабая регрессия проявляется в таких стресс-индуцированных функциональных состояниях, как хронический стресс, эмоциональное выгорание, клинический стресс при любом заболевании как реакция на повреждение или его угрозу.

Протокол для перехода от слабой регрессии к эталонной ритмике мозга реализуется по следующей схеме: фон ДО – запись ЭЭГ в покое, 2 минуты; динамическая фотостимуляция, 5 минут; пауза – запись ЭЭГ без стимуляции, 1 минута; для световой стимуляции используется режим «ЭЭГ – адаптивная нейромодуляция», для звуковой стимуляции в качестве зоны интереса назначается альфа-ритм (8–13 Гц), функция обратной связи реализуется через обратную корреляцию между мощностью альфа-ритма и высотой тональных модуляций, 5 минут; фон ПОСЛЕ – запись ЭЭГ в покое, 2 минуты. Сеансы НБУ по этому протоколу проводятся 1 раз в день в одно

и то же время суток. Стимуляция по этому протоколу прекращается только после проявления признаков эталонного уровня зрелости.

Успешность сеансов резонансного НБУ проявляется в следующих признаках: усиление мощности альфа-ритма; снижение мощности дельта- и тета-ритма; повышение пиковой частоты в альфа-диапазоне; проявление способности к усвоению ритма и мультипликации. Благодаря успешным сеансам резонансного НБУ с аудиовизуальной обратной связью по ЭЭГ активируются механизмы нейропластичности и формируются оптимальные условия для компенсации стресс-индуцированных и органических нарушений функций мозга.

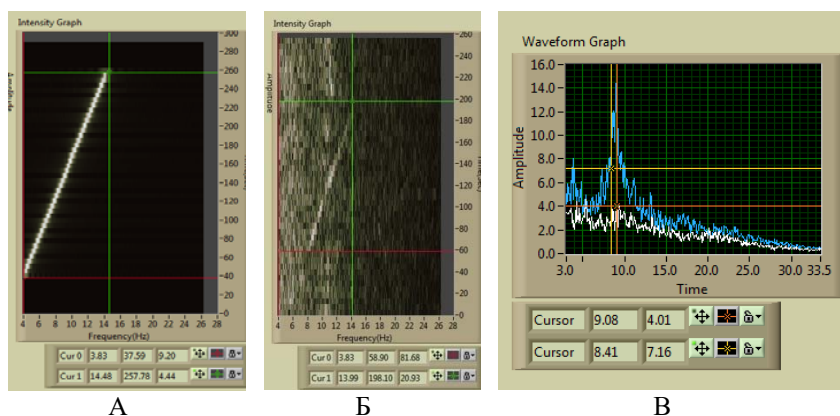


Рис. 5. Признаки слабой регрессии ритмики мозга. А – динамический спектр светового сигнала при динамической фотостимуляции; Б – динамический спектр ЭЭГ при функциональной пробе «динамическая фотостимуляция» у ребенка с тревожным расстройством; В – спектр ЭЭГ до (светлые линии) и после (более темные линии) динамической фотостимуляции

Заключение

Таким образом, полученные нами обширные данные позволяют утверждать, что применение резонансного НБУ с двойной обратной связью по визуальному и акустическому входам не только позволяет компенсировать полную или частичную утрату нейропластичности при функциональных нарушениях и органических патологиях, но и эффективно использовать эту

технологии для диагностики этих нарушений. Однако возможности технологии не исчерпываются этими сугубо клиническими, реабилитационными достижениями.

Обнаруженная нами регрессия нейропластичности при стрессе оказалась не единственным феноменом, непосредственно связанным с проблемой объективизации признаков сознания. Как известно [8], в современной когнитивной науке одним из концептуальных подходов к объективному исследованию сознания является регистрация динамики физиологических функций при его депривации или даже отсутствии. В качестве экспериментальных моделей принято использовать 3 модели: избирательное внимание, общий наркоз и сон [8]. Однако, как правило, эти модели используются в экспериментах на животных (начиная от приматов и заканчивая дрозофилой). Мы полагаем, что технология НБУ может оказаться приемлемой для подобных исследований на людях-добровольцах.

В качестве примера приводим динамические характеристики ритмики мозга волонтера 87 лет на фоне проведения процедуры НБУ (в целях реабилитации) при активном бодрствовании и во время глубокого сна (рис. 6).

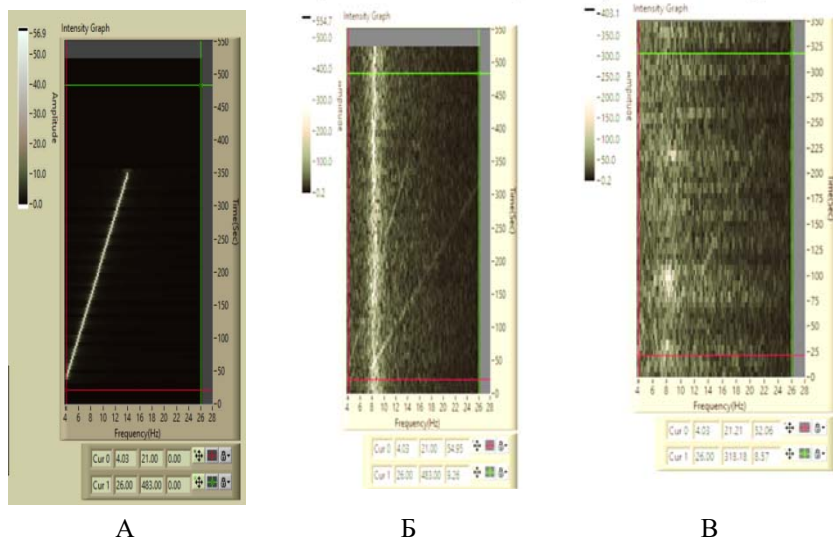


Рис. 6. Эффекты динамической фотостимуляции при разных состояниях сознания одного и того же испытуемого: А – характеристика стимула; Б – выраженная мультипликация при бодрствовании; В – редукция нейропластичности во время глубокого сна

Особого внимания заслуживает тот факт, что во время сна сохраняется усвоение ритма, но полностью редуцируется мультипликация, характерная для мозга этого испытуемого во время бодрствования. Очевидно, что подобные результаты требуют дополнительного углублённого исследования.

Список литературы

1. Kaplan A.Y. Neurophysiological foundations and practical realizations of the brain–machine interfaces in the technology in neurological rehabilitation // *Human Physiology*. 2016. V. 42, N 1. P. 103–110.
2. Харченко Е.П., Тельнова М.Н. Пластичность мозга: ограничения и возможности // *Журнал неврологии и психиатрии им. С.С. Корсакова*. 2017. Т. 117. № 1-2. С. 8-13.
3. Fedotchev A.I., Bondar A.T. Adaptive neurostimulation, modulated by subject's own rhythmic processes, in the correction of functional disorders // *Human Physiology*. 2022. V. 48, N 1. P. 108–112.
4. Fedotchev A.I., Parin S.B., Plevaya S.A., Zemlianaia A.A. Effects of Audio–Visual Stimulation Automatically Controlled by the Bioelectric Potentials from Human Brain and Heart // *Human Physiology*. 2019. V. 45, N 5. P. 523–526. doi: 10.1134/S0362119719050025.
5. Bishop M.P., Elder S.N., Heath R.G. Intracranial Self-Stimulation in Man // *Science*. 1963. V. 140, N 3565. P. 394–396.
6. Александров Ю.И., Сварник О.Е., Знаменская И.И., Колбенева М.Г., Арутюнова К.Р., Крылов А.А., Булава А.И. Регрессия как этап развития. Москва : Изд-во «Институт психологии РАН». 2017. 191 с.
7. Парин С.Б. Стресс, боль и опиоиды. Об эндорфинах и не только. Минск : Дискурс. 2021. 208 с.
8. Van De Poll M.N., van Swinderen B. Balancing prediction and surprise: A role for active sleep at the dawn of consciousness? // *Frontiers in Systems Neuroscience*. 2021. N 15. P. 768762. doi: 10.3389/fnsys.2021.768762.

**С.В. СТАСЕНКО, А.А. ЛЕБЕДЕВ, О.В. ШЕМАГИНА,
И.В. НУЙДЕЛЬ, А.В. КОВАЛЬЧУК, В.Г. ЯХНО**

Институт прикладной физики РАН, Нижний Новгород
aka.xzib1t@gmail.com

АДАПТИВНАЯ КОРРЕКЦИЯ МНОГОКАСКАДНОГО ДЕТЕКТОРА БИОМОРФНОЙ СИСТЕМЫ ИСКУССТВЕННОГО ИНТЕЛЛЕКТА ДЛЯ ЗАДАЧ РАСПОЗНАВАНИЯ ОБРАЗОВ*

В данной работе исследуется адаптивная коррекция многокаскадного детектора в биоморфной системе искусственного интеллекта для решения задач распознавания образов. Отличительной особенностью этой системы является имитация иерархической обработки информации, наблюдаемой в живых системах, например в зрительной коре головного мозга. В рамках данного исследования был апробирован алгоритм коррекции результатов многокаскадного детектора и сформулированы выводы о результатах его применения.

Ключевые слова: *системы распознавания, семантический анализ изображений, символическое распознавание информации, каскадное обнаружение, технология адаптивной коррекции ошибок.*

Введение

Применение методов искусственного интеллекта (ИИ) и машинного обучения в последние годы привело к многочисленным успехам [1, 2]. Примеры, когда сообщалось, что человеческий уровень производительности может быть сопоставим или превышен, включают идентификацию рака молочной железы [3], обнаружение объектов, скрытых от глаз [4], овладение настольными играми [5], оптимизацию новых методов визуализации [6] и разработку систем для автономных самоуправляемых автомобилей [7].

Существующие прорывы явно стимулируют дальнейшие исследования и стимулируют широкое внедрение таких систем на практике. Однако в об-

* Данная работа выполнена при поддержке гранта на выполнение проекта открытых библиотек ФГБУ «Фонд содействия развитию малых форм предприятий в научно-технической сфере» в рамках федерального проекта «Искусственный интеллект» национальной программы «Цифровая экономика Российской Федерации» (соглашение № 3ГУКодИИС12-D7/72687 от 26 декабря 2017 г.).

ласти исследований, где, например, дорожные знаки “Стоп” на обочине дороги могут быть неправильно истолкованы как знаки ограничения скорости при добавлении минимального количества граффити [8], многие комментаторы задаются вопросом, являются ли текущие решения достаточно надежными, устойчивыми и заслуживающими доверия и как такие проблемы должны быть количественно оценены и решены. Маркус [9] выделяет десять проблем, связанных с текущим состоянием глубокого обучения, одна из которых заключается в том, что “глубокое обучение до сих пор хорошо работает как приближение, но его ответам часто нельзя полностью доверять.”

Для преодоления сложившейся проблемы стали развиваться подходы, основанные на явлении стохастического разделения в больших размерностях, предоставившие возможность быстрой и неитеративной коррекции ошибок ИИ. Первые теоремы этого класса были описаны как проявления «благословения размерности» [10, 11]. Унаследованные системы искусственного интеллекта дополняются корректорами. Эти простые интеллектуальные устройства отделяют распознанные ошибки и их окружение от ситуаций с правильным функционированием и заменяют ошибочные решения ИИ исправленным. Корректоры используются для решения классической задачи повышения чувствительности и специфичности (устранения ложноположительных и ложноотрицательных результатов классификации), для передачи знаний между системами искусственного интеллекта, для обучения мультиагентных систем и других целей.

Среди различных видов классификаций по уровню значимости и уровню сложности выделяется задача семантической классификации зрительных образов [12], позволяющая получить ответ на вопрос: «Что изображено на картинке?»

Методы решения задачи семантической классификации изображения сохраняют в целом типичную схему, но имеют большое разнообразие в способах реализации этапов. Проблема распознавания семантики изображения (в зарубежной литературе используется аббревиатура CBIR – Content-based image retrieval) имеет более чем 20-летнюю историю [13]. Первые работы относятся к 80-м годам прошлого века [14]. Начиная с 2000 года разработаны такие системы: content-based image retrieval with relevance feedback в университете Illinois [15], система PicToSeek [16], DrawSearch, Viper [14, 17] Такие системы [18] совершенствуются с помощью оптимизации поисковых процедур для изображений в конкретных базах данных большого объема. В некоторых системах пространство признаков имеет очень большую размерность.

Нейросетевые модели хорошо подходят для решения данной задачи, но требуют наличия большого объема обучающих данных для каждого нового запроса. Альтернативный подход несупервизорного обучения, основанный на использовании самоорганизующихся карт Кохонена, рассматривается в работе [19]. Достоинство этого метода в том, что в процессе обучения системы не требуется взаимодействия с человеком-оператором. Но вследствие ограниченности правил формирования карт это, с другой стороны, приводит к снижению разделяющей мощности системы распознавания. Кроме того, для восприятия результата классификации человеком необходимо выполнить трактовку автоматически вырабатываемых концептуальных понятий, что не всегда возможно. Сегодня технологии глубокого обучения существенно улучшили и упростили алгоритмы семантической сегментации [20, 21], и широко применяются в реальной жизни (анализ изображений со спутников, анализ мимики лиц, анализ медицинских изображений и т.п.).

По данным литературы видно, что некоторые разработчики систем семантической классификации ищут решения в области биоморфных моделей для того, чтобы программные или программно-аппаратные комплексы, выполняли функциональные операции, аналогичные своим биологическим прототипам [22]. В этом случае сами системы обработки информации можно будет рассматривать как нейроморфные функциональные модели живых систем.

Биоморфная система искусственного интеллекта для распознавания образов

Структурно система биоморфного нейроподобного искусственного интеллекта для семантического анализа выглядит следующим образом: для каждого понятия из словаря определен свой каскадный детектор [23]. Каскадный детектор, в свою очередь, основан на сильных классификаторах, основанных на нелокальных бинарных шаблонах [24]. Детектор атрибутов, который также основан на нелокальных бинарных шаблонах, определяет свойства объекта.

Метод многокаскадного обнаружения используется для поиска объектов на изображении. Основу каждого детектора составляет каскад последовательно соединенных сильных классификаторов.

Детектирование выполняется с использованием каскадного детектора, состоящего из набора сильных классификаторов, соединенных последовательно. Рисунок 1а раскрывает внутреннюю структуру сильных классификаторов.

Как видно из рис. 1а, каждый сильный классификатор представляет из себя классический нейрон МакКаллока – Питтса или формальный нейрон (1).

$$H_i = \begin{cases} 1, \sum_{k=1}^n w_k h_k > \theta \\ 0, \sum_{k=1}^n w_k h_k \leq \theta \end{cases} \quad (1),$$

где h_k – слабый классификатор (2), w_k – вес слабого классификатора, θ – порог принятия решения. h_k – слабый классификатор, построен на отношении правдоподобия $\hat{L}$:

$$h_k = \begin{cases} 1, \hat{L} \geq 1 \\ 0, \hat{L} < 1 \end{cases}, \hat{L} = \frac{\hat{\omega}_k(c|\Omega = 1)}{\hat{\omega}_k(c|\Omega = 0)} \quad (2),$$

где $\hat{\omega}_k$ – оценка плотности вероятности распределения значений признака изображения c , при условии, что он либо принадлежит изображению мишени $\Omega = 1$, либо не принадлежит ему $\Omega = 0$.

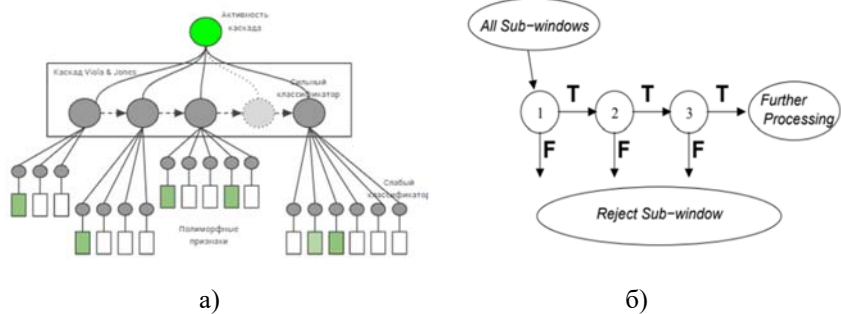


Рис. 1. а) Каскадный детектор, настроенный на анализ прямоугольного фрагмента изображения, активируется, если изображения содержит искомый объект. б) Каскадный детектор Виолы и Джонса [23]

В качестве анализируемого признака s для оценки плотности распределения используется так называемый «нелокальный бинарный шаблон», общая схема которого представлена на рис. 2.

Суть кодирования фрагмента изображения с использованием нелокального бинарного шаблона заключается в следующем. Произвольный фрагмент изображения прямоугольной формы разбивается на 9 прямоугольных фрагментов. Признаки определяются как ядро, размером 3×3 элемента, отражающее пространственную структуру изображения. Внутри ядра применяются бинарное кодирование информации $\{0, 1\}$ и результирующие бинарные шаблоны могут представлять собой границы, отрезки, седловые точки, точки соединения и так далее. На решетке, размером 3×3 элемента существуют $2^9 = 512$ подобных ядер. При этом, для кодирования графической информации мы используем прямоугольные ядра различных размеров. Используемая схема кодирования не позволяет фильтровать изображение, а, как в схеме Виолы и Джонса, [23] преобразует изображение в пространство признаков с размерностью много большей, чем размер исходного изображения. Такой вид преобразования мы будем называть «нелокальным бинарным шаблоном».

Для получения значений весов w_k из (1) и порогов θ , используется процедура AdaBoost, которая подробно описана в [23].

По окончании обучения формируется ключевой элемент системы поиска и распознавания в виде колонки (рис. 3б).

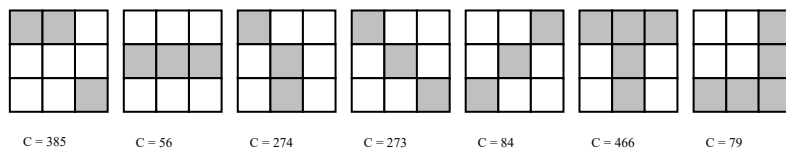


Рис. 2. Пример нелокальных бинарных шаблонов для кодирования произвольного фрагмента изображения

Эта колонка представляет собой детектор, реагирующий на объекты заданного типа. Детектор, в свою очередь, представляет собой каскад последовательно соединенных сильных классификаторов (рис. 1а). Метод каскадного обнаружения используется для обнаружения объектов на изображении.

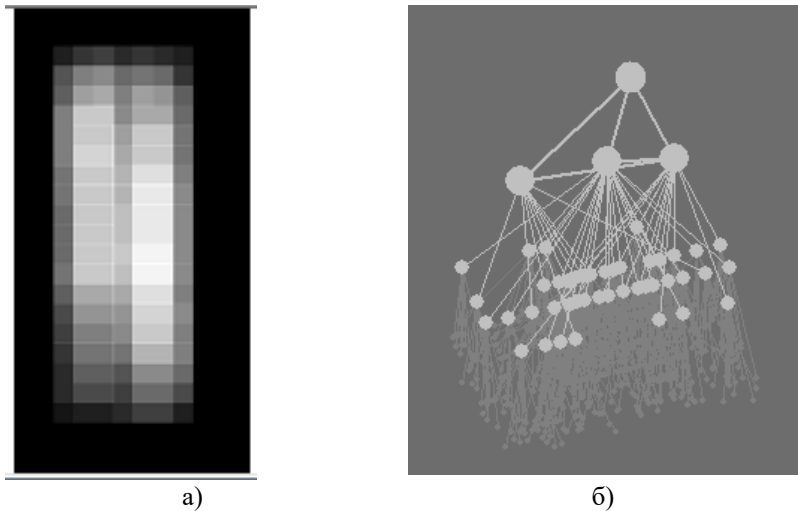


Рис. 3. а) Пример наложения прямоугольных решеток 3х3 внутри области апертуры.
б) 3D изображение сформированной в результате обучения сети

Тестирование адаптивной коррекции многокаскадного детектора

Для коррекции многокаскадного детектора использовались данные статистики правильных и ошибочных результатов его работы согласно алгоритму, описанному в работе [25]. Для анализа данных была предоставлена база тестирования детектора, собранная из работы детектора на изображениях цифр (рис. 4).

	Result	Score	left	top	right	bottom	IOU	SampleTest	feature0	feature1	...	feature440	feature441	feature442	feature443
418	0	0.466629	32	0	52	41	0.306346	TN	109.657143	109.000000	...	111.666664	111.0	111.833336	108.000000
561	0	0.376622	72	16	86	44	0.289941	TN	111.449997	111.900002	...	108.000000	107.0	108.500000	109.000000
614	0	0.306268	76	40	90	68	0.276720	TN	150.600006	146.800003	...	154.500000	153.0	153.000000	148.000000
553	0	0.459631	132	0	163	62	0.341772	TN	200.875000	201.000000	...	201.066666	201.0	201.000000	200.800003
595	0	0.305533	108	0	138	61	0.272727	TN	199.962494	200.512497	...	200.000000	201.0	201.000000	200.000000

Рис. 4. Пример данных

В таблице 1 указано распределение базы данных на категории ошибок/правильных ответов. Для дальнейшего анализа необходимы TP и FP измерения, которые и были отобраны в количестве 7812 и 2501 штук.

Таблица 1

Распределение измерений по категориям в результате работы детектора

TP	FP	FN	TN
2%	1%	26%	72%

Алгоритм, описанный в работе [25], транслирует данные в мультикластерное признаковое пространство, в котором ставится задача минимизации ошибки в каждом из кластеров. Для оценки полученного признакового пространства внутри каждого кластера были построены распределения измерений, попавших в исследуемый кластер. Как демонстрируется на рис. 5а, б распределения не всегда удается удачно разнести.

В результате тестирования алгоритма коррекции на train и test выборках (полученных из исходной при помощи генерации случайной подвыборки) были получены следующие результаты, представленные на рис. 6.

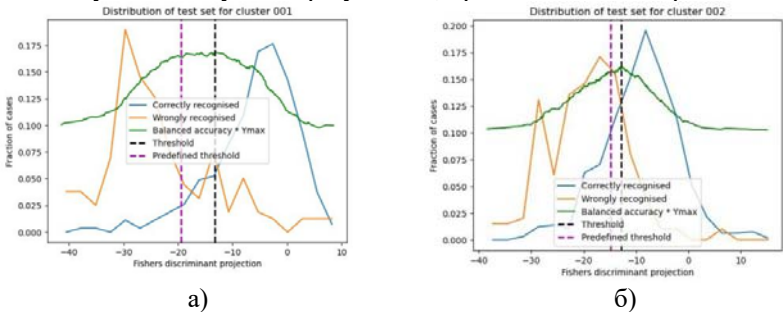


Рис. 5. Пример хорошей (а) и плохой (б) разделимости измерений для одного из кластеров

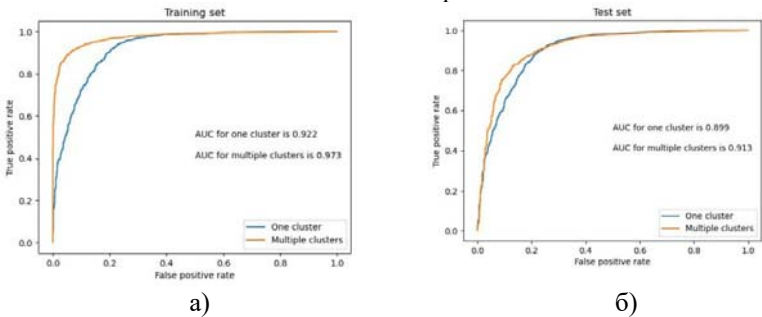


Рис. 6. ROC-AUC кривые для случая без коррекции по кластерам и с ней (синия и оранжевые кривые соответственно) для обучающего (а) и тестового (б) множества

Выводы

Разработанная биоморфная система искусственного интеллекта для семантического анализа изображений имитирует процесс обработки сигналов в зрительной коре головного мозга. В предлагаемой системе формирование детекторов определенного типа в процессе обучения и внешней стимуляции соответствует образованию вертикальных колонок нейронов в зрительной коре головного мозга. Реакция колонки на предъявленный стимул моделируется с помощью процесса идентификации объектов на входных изображениях.

Ошибки распознавания возникают при поступлении новых данных, на которых биоморфная система искусственного интеллекта ранее не обучалась. Использование мультикорректорного подхода позволяет исправлять ошибки распознавания изображений без переобучения исходной искусственной когнитивной системы без снижения качества распознавания.

Список литературы

1. LeCun Y., Bengio Y., Hinton G. Deep learning // Nature. – 2015. – V. 521, N 7553. – P. 436–444.
2. Schmidhuber J. Deep learning in neural networks: An overview // Neural networks. – 2015. – V. 61. – P. 85–117..
3. McKinney S. M. et al. International evaluation of an AI system for breast cancer screening // Nature. – 2020. – V. 577, N 7788. – P. 89–94.
4. Caramazza P. et al. Neural network identification of people hidden from view with a single-pixel, single-photon detector // Scientific reports. – 2018. – V. 8, N 1. – P. 11945.
5. Silver D. et al. Mastering the game of Go with deep neural networks and tree search //nature. – 2016. – V. 529, N 7587. – P. 484–489.
6. Higham C. F. et al. Deep learning for real-time single-pixel video //Scientific reports. – 2018. – V. 8, N 1. – P. 2369.
7. Bojarski M. et al. End to end learning for self-driving cars // arXiv preprint arXiv:1604.07316. – 2016.
8. Evtimov I. et al. Robust physical-world attacks on machine learning models //arXiv preprint arXiv:1707.08945. – 2017. – V. 2, N 3. – P. 4.
9. Marcus G. Deep learning: A critical appraisal // arXiv preprint arXiv:1801.00631. – 2018.
10. Gorban A. N., Tyukin I. Y., Romanenko I. The blessing of dimensionality: Separation theorems in the thermodynamic limit // IFAC-PapersOnLine. – 2016. – V. 49, N 24. – P. 64–69.
11. Gorban A. N., Tyukin I. Y. Blessing of dimensionality: mathematical foundations of the statistical physics of data // Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences. – 2018. – V. 376, N 2118. – P. 20170237.

12. Tadi Bani N., Fekri-Ershad S. Content-based image retrieval based on combination of texture and colour information extracted in spatial and frequency domains // *The electronic library*. – 2019. – V. 37, N 4. – P. 650–666.
13. Telynykh A. et al. A Biomorphic Model of Cortical Column for Content-Based Image Retrieval // *Entropy*. – 2021. – V. 23, N 11. – P. 1458.
14. Bhaumik H. et al. Hybrid soft computing approaches to content based video retrieval: A brief review // *Applied Soft Computing*. – 2016. – V. 46. – P. 1008–1029.
15. Hu Y. H., Hwang J. N (ed.). *Handbook of neural network signal processing*. – CRC press, 2018.
16. Rui Y., Huang T. S., Mehrotra S. Content-based image retrieval with relevance feedback in MARS // *Proceedings of international conference on image processing*. – IEEE, 1997. – V. 2. – P. 815–818.
17. Gevers T., Smeulders A. W. M. Pictoseek: Combining color and shape invariant features for image retrieval // *IEEE transactions on Image Processing*. – 2000. – V. 9, N 1. – P. 102–119.
18. Naphade M. R., Huang T. S. Extracting semantics from audio-visual content: the final frontier in multimedia retrieval // *IEEE Transactions on Neural Networks*. – 2002. – V. 13, N 4. – P. 793–810.
19. Laaksonen J., Koskela M., Oja E. Content-based image retrieval using self-organizing maps // *Visual Information and Information Systems: Third International Conference, VISUAL'99 Amsterdam, The Netherlands, June 2–4, 1999 Proceedings 3*. – Springer Berlin Heidelberg, 1999. – P. 541–549.
20. Unglert K., Radić V., Jellinek A. M. Principal component analysis vs. self-organizing maps combined with hierarchical clustering for pattern recognition in volcano seismic spectra // *Journal of Volcanology and Geothermal Research*. – 2016. – V. 320. – P. 58–74.
21. Putzu L., Piras L., Giacinto G. Convolutional neural networks for relevance feedback in content based image retrieval: A Content based image retrieval system that exploits convolutional neural networks both for feature extraction and for relevance feedback // *Multimedia Tools and Applications*. – 2020. – V. 79, N 37. – P. 26995–27021.
22. Yakhno V. G. et al. Who says formalized models are appropriate for describing living systems? // *International Conference On Neuroinformatics*. – Cham : Springer International Publishing, 2020. – P. 10–33.
23. Viola P., Jones M. Rapid object detection using a boosted cascade of simple features // *Proceedings of the 2001 IEEE computer society conference on computer vision and pattern recognition. CVPR 2001*. – Ieee, 2001. – V. 1. – P. I-I.
24. Telynykh A., Nuidel I., Samorodova Y. Construction of efficient detectors for character information recognition // *Procedia Computer Science*. – 2020. – V. 169. – P. 744–754.
25. Gorban A. N. et al. High-dimensional separability for one-and few-shot learning // *Entropy*. – 2021. – V. 23, N 8. – P. 1090.

ИНТЕРФЕЙС "МОЗГ-КОМПЬЮТЕР". НЕЙРОТЕХНОЛОГИИ

О.А. ГРИГОРЬЕВА, Д.Ф. КЛЕЕВА

Центр биоэлектрических интерфейсов НИУ ВШЭ, Москва
oagrigoreva_2@edu.hse.ru

СООТНОШЕНИЕ ФИЗИОЛОГИЧЕСКИХ ХАРАКТЕРИСТИК СТАДИЙ СНА МЕЖДУ ОДНОВРЕМЕННО РЕГИСТРИРУЕМЫМИ ЭЭГ И МЭГ

В данном исследовании проведен комплексный анализ различных физиологических маркеров сна, выделенных в одновременно зарегистрированных ЭЭГ- и МЭГ-данных с их последующим сопоставлением в пределах нескольких стадий. Кроме того, в рамках одной из задач исследования был реализован синтез данных «смешанной» модальности при помощи реконструкции источников по МЭГ-записи с применением обратного оператора и дальнейшего проецирования полученного сигнала на ЭЭГ-сенсоры.

Ключевые слова: физиологические маркеры сна, реконструкция источников, классификация стадий сна.

Введение

В последние десятилетия приобретают популярность исследования, нацеленные на поиск и сопоставление наиболее информативных физиологических маркеров для классификации стадий сна в МЭГ- и ЭЭГ-записях [1, 2]. Интерес к данным исследовательским проблемам объясняется ростом значимости точного определения стадий сна в клинической практике наряду с появившейся в связи с непрерывным совершенствованием методов регистрации сигналов мозга возможностью дополнения информации о частотно-временных характеристиках сигнала информацией о его локализации [3]. Так, с внедрением МЭГ в научно-исследовательскую и клиническую практику стали появляться предположения о том, что МЭГ как метод, обладающий отличной от ЭЭГ «чувствительностью» к локализации источников, может обладать некоторым «информационным преимуществом» при решении задач классификации стадий сна. Обнаруживаемые различия в ЭЭГ- и МЭГ-записях могут быть объяснены тем, что данные методы регистрируют разнонаправленные диполи источников, располагающихся, соответственно, по-разному по отношению к поверхности черепа. Так, на ЭЭГ-записях лучше видны синхронные паттерны активности (например, сонные веретена), так как диполи порождающих их источников чаще всего

расположены радиально. При помощи МЭГ, с другой стороны, скорее можно зарегистрировать источники, диполи которых имеют тангенциальное расположение. Кроме того, результаты исследований с применением частотно-временного разложения показывают, что сигнал в одном и том же временном окне, визуально различно проявляющийся на разных сенсорах, может быть «разложен» на разное количество источников в зависимости от типа регистрирующих сенсоров: МЭГ идентифицирует более фокальную и разнообразную топографию в отличие от ЭЭГ, что может вносить вклад в регистрацию синхронной активности и способствовать её «зашумлению» [4]. Во многих работах, посвященных описанию классификаторов стадий сна, упоминается факт обучения моделей в большинстве случаев на данных, полученных при помощи ЭЭГ [1, 2]. Таким образом, можно предположить получение на практике низких значений сходимости между последовательностями стадий, предсказанных при помощи классификаторов на основании МЭГ-данных и последовательностями стадий, предсказанных на основании ЭЭГ-записей. Для преодоления данного ограничения в настоящем исследовании предполагается преобразовать имеющийся ЭЭГ-сигнал таким образом, чтобы увеличить сходимость между предсказаниями стадий для «преобразованной» записи и последовательностями классифицированных стадий для «первоначальной» ЭЭГ-записи.

Постановка задачи и описание алгоритма

В ходе данного исследования предполагается выполнить следующие задачи:

1. Проведение комплексного анализа в виде сопоставления физиологических маркеров в пределах различных стадий сна в одновременно регистрируемых ЭЭГ- и МЭГ-записях;
2. Повышение результатов сходимости между предсказанными последовательностями стадий для данных с различных типов сенсоров.

Потенциальные ограничения на этапе выполнения последней задачи могут быть обусловлены сложностью избавления от артефактов в синтезированной записи (например, ЭКГ-наводки). Кроме того, потенциальные низкие результаты сходимости между предсказаниями для нового и первоначального сигнала могут быть объяснены распространенной практикой применения обратного оператора в решении задач реконструкции источников на основании данных с вызванными потенциалами [6].

Выборка и процедура

Для анализа данных были использованы данные сомнографии трех испытуемых (2 женщины, 1 мужчина, средний возраст $M = 25.3$). Критерии

включения требовали, чтобы участники не страдали эпилепсией, не подвергались хирургическим вмешательствам, изменяющим структуру черепа, и не имели имплантированных металлических конструкций. Все участники были подвержены депривации сна: его продолжительность сокращалась на три часа по сравнению с обычной продолжительностью сна в предыдущую ночь. На протяжении всего исследования каждый участник был оснащен датчиками для отслеживания положения головы. Кроме того, на каждого добровольца надевали шапочку с ЭЭГ-электродами, после чего каждому участнику было предложено занять место в МЭГ-устройстве для одновременной регистрации данных с различных источников.

Продолжительность экспериментальной сессии составила около 80 мин. Во время записи добровольцы выполняли следующий набор проб: глаза открыты (2 мин); глаза закрыты (2 мин); глаза открыты-закрыты на 5 секунд (1 мин); движение глаз в вертикальной плоскости (30 с); движение глаз в горизонтальной плоскости (30 с); глотание (30 с); облизывание губ (30 с); кашель (30 с); улыбка (30 с); скользящие движения головой в горизонтальной плоскости (30 с); жевание (30 с); гипервентиляция (2 мин). Затем участникам предлагалось погрузиться в непродолжительный сон.

Магнитоэнцефалография (МЭГ) и электроэнцефалография (ЭЭГ) были зарегистрированы с помощью 366-канального аппаратно-программного комплекса «VectorView» (Elekta Neuromag Oy, Финляндия) (204 градиометра, 102 магнитометра, 60 каналов ЭЭГ). Данные записывались с частотой дискретизации 1000 Гц и фильтром в 0,1–330 Гц. Положение головы относительно массива датчиков во время эксперимента отслеживалось в режиме реального времени с помощью катушек индуктивности. Референтный ЭЭГ-электрод был установлен либо на линии роста волос, либо на мочке уха для повышения качества записываемого сигнала и уменьшения риска записи артефактов. Заземляющий электрод располагался на ключице.

Анализ данных и результаты

На этапе предварительной обработки была произведена фильтрация при помощи MaxFilter и BandPass-filter (1–40) Гц с интерполяцией зашумленных каналов. В проведенном анализе данных можно выделить следующие этапы:

1. ***Выделение стадий сна и расчет сходимости в пределах разных регионов.*** Данный этап выполнен при помощи функционала, предложенного в документации классификатора YASA. Полученные на данном этапе последовательности были сопоставлены между собой в пределах различных регионов, а также безотносительно различных зон мозга.

В результате данного этапа анализа были получены значения сходимости в процентах. Самые высокие значения сходимости были в сравнениях между магнитометрами и градиометрами, самые низкие – для сравнений ЭЭГ/магнитометры, причем результаты сравнений ЭЭГ/Градиометры были выше, чем результаты сравнений ЭЭГ/магнитометры (табл. 1).

Таблица 1

Результаты расчета сходимости путем попарных сравнений последовательностей стадий с трех типов электродов для трех участников

	Участник 1	Участник 2	Участник 3
ЭЭГ/Градиометры	53.54 %	66.67 %	75.89 %
ЭЭГ/Магнитометры	49.61 %	66.67 %	73.81 %
Градиометры/Магнитометры	85.04 %	86.67 %	91.07 %

Кроме того, значения сходимости были рассчитаны для разных регионов. В результате были обнаружены более высокие значения сходимости для правой височной доли по сравнению с другими зонами для всех трех участников.

2. Спектральный анализ. Для участков с отличающимися предсказанными стадиями для разных сенсоров были построены графики спектральной мощности (PSD) для визуализации потенциальных различий (рис. 1). Различия были обнаружены в пределах стадий N1 и N2 (стадии поверхностного и глубокого сна соответственно).

Исходя из визуального сравнительного анализа графиков, можно отметить явные расхождения между спектрами в N1- и N2-стадиях для ЭЭГ- и МЭГ-сенсоров. В стадии N1 для ЭЭГ-записей наблюдается выраженность сигнала, близкого к 20 Гц (что соответствует верхней границе бета-ритма), тогда как для МЭГ-записей характерна высокая плотность в области сигнала более низкой частоты около 12 Гц (что соответствует верхнему диапазону альфа-ритма). Напротив, в N2 стадии можно наблюдать обратные явления: на МЭГ-записи больше проявляется сигнал, близкий по частоте к верхней и нижней границе бета-ритма (15 и 20 Гц), тогда как на ЭЭГ-записи больше проявляются осцилляции в нижнем диапазоне бета-ритма

(13–14 Гц). Кроме того, в стадии N2 в ЭЭГ можно обнаружить зарегистрированные медленноволновые колебания в районе 3–5 Гц, что соответствует дельта- диапазону.

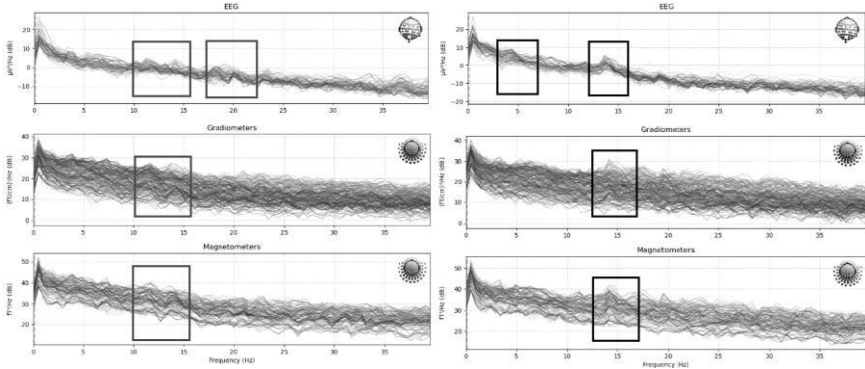


Рис. 1. Результат расчета PSD для совмещенных записей. На трех графиках слева (верхний – ЭЭГ, два нижних – магнитометры и градиометры соответственно) прямоугольниками отмечены различия, обнаруженные на стадии N1. На трех графиках справа (расположение аналогично графикам слева) прямоугольниками отмечены различия, обнаруженные на стадии N2

Также для трех участников был проведен t-тест для определения статистической значимости результатов расчета значений спектральных мощностей в двух условиях (стадии N1 и N2) для выделенных диапазонов частот в ЭЭГ- и МЭГ-сенсорах. Таким образом, мы получили 3 значения t-теста и 3 значения p-value для каждого из двух условий (табл. 2).

Таблица 2

**Результаты расчета сходимости путем попарных сравнений
последовательностей стадий с трех типов электродов
для трех участников**

	Участник 1	Участник 2	Участник 3
N1: ЭЭГ/МЭГ, t-критерий	1.99	1.87	1.93
N1: ЭЭГ/МЭГ, p-value	0.04	0.04	0.03
N2: ЭЭГ/МЭГ, t-критерий	2.32	2.13	1.86
N2: ЭЭГ/МЭГ, p-value	0.002	0.02	0.001

3. Анализ сонных веретен. На данном этапе были выделены эпохи в записях на основании обнаруженных сонных веретен с последующим сравнением топографий на МЭГ- и ЭЭГ-сенсорах при помощи частотно-временного разложения и локализации источников. В результате были обнаружены различия в визуализации временных рядов и топографиях, зарегистрированных при помощи разных сенсоров: вышеупомянутые паттерны были более четко выражены на ЭЭГ, тогда как сигнал на МЭГ в том же временном окне был выражен менее четко. Кроме того, локализация источников на МЭГ более фокальная и разнообразная по сравнению с менее детализированной и более однообразной топографией на ЭЭГ.

Можно отметить локализацию источников в височной, затылочной, лобной и теменной долях на МЭГ, наряду с локализацией источников в затылочной и теменных долях на ЭЭГ. Однако на основании результатов частотно-временного анализа можно отметить выраженную активность в одном временном окне на всех трех типах сенсоров.

4. Применение скрытой модели Маркова. На данном этапе в данных, полученных с разных сенсоров, были выделены кластеры К-комплексов на основании их динамических и пространственных характеристик. Затем к полученным кластерам была применена модель Маркова с последующим сопоставлением полученных матриц вероятностей перехода из одного состояния в другое, где в качестве состояния был представлен определенный кластер. Результаты для одного участника представлены на рис. 2.

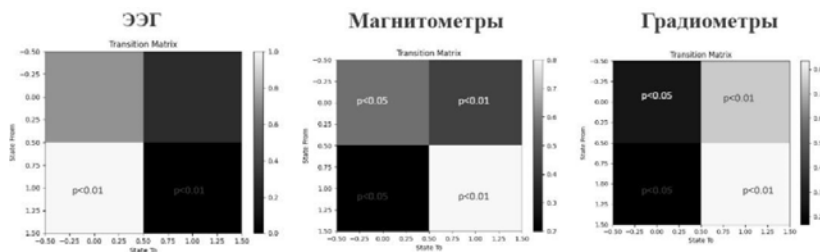


Рис. 2. Матрицы вероятности, полученные при помощи применения скрытой модели Маркова к кластерам, обозначенных как состояния с возможным взаимным переходом. Статистически значимые вероятности перехода сопровождаются обозначенными на матрицах p-value

Можно отметить большее количество статистически значимых результатов для МЭГ-сенсоров (градиометры и магнитометры) по сравнению с ЭЭГ- сенсорами.

Синтез ЭЭГ на основе МЭГ при помощи обратного оператора

Полученные на основе подсчета сходимости между последовательностями стадий сна для разных сенсоров и эксплораторного анализа результаты, а также результаты визуализации моделирования источников позволяют подтвердить предположение о том, что низкая сходимость между данными ЭЭГ и МЭГ может быть объяснена рядом факторов: различной топологией источников, генерирующих синхронную активность, характерную для различных стадий сна, а также разной «чувствительностью» к топографиям (локализациям) источников при их регистрации. Необходимость точного определения стадий сна на записях с различных типов сенсоров (МЭГ- и ЭЭГ-) в клинической практике поднимает проблему повышения сходимости последовательностей, определяемых для разных типов записей. В ряде исследований были представлены идеи и последующие основанные на них попытки синтеза данных одной модальности на основе данных другой модальности [6]. В целом, данные исследования посвящены попыткам взаимной реконструкции фМРТ и ЭЭГ/МЭГ-данных, а также синтезу МЭГ-сигнала на основе ЭЭГ-сигнала. В настоящем исследовании в целях повышения показателя сходимости между предсказанными для МЭГ- и ЭЭГ-записей последовательностями была предпринята попытка видоизменения первоначального ЭЭГ-сигнала путем реконструкции источников по МЭГ-сигналу с использованием обратного оператора и последующим проецированием полученного сигнала на ЭЭГ-сенсоры. На данном этапе был произведен расчет прямой модели на основании данных о расположении источников, а также расчет генерализованного обратного оператора (minimum norm estimate, MNE) с его последующей нормировкой на шумовую компоненту при помощи dSPM.

Данные, полученные с ЭЭГ- и МЭГ- сенсоров можно представить следующим образом:

$$X = G(S) + N,$$

где N – шум, S – данные об источниках в виде координат сенсоров.

На основе приведенных данных можно рассчитать показатели прямой модели G , опираясь на некоторые «усредненные» параметры модели головы. На основе данных о сенсорах и рассчитанных данных прямой модели можно рассчитать значения обратного оператора (MNE, W). Используя затем данные о сенсорах, можно рассчитать активность, реконструированную только по данным с МЭГ:

$$W * X \rightarrow S',$$

где S' – реконструированная активность.

Далее, используя данные реконструированной только по МЭГ активности и показатели прямой модели G с ЭЭГ-источников, можно спроецировать ее на ЭЭГ-сенсоры, таким образом в конечном счете получив ЭЭГ-активность, синтезированную на основе реконструированных МЭГ-источников:

$$S' * G(eeg) = X(eeg).$$

Чтобы проверить, насколько хорошо соотносятся предсказанные стадии сна для первоначального и синтезированного ЭЭГ-сигнала был проведен расчет значения общей информации (mutual info). Попытка синтеза ЭЭГ-сигнала на основе реконструированных МЭГ-источников была принята только для первого участника, так как в результате расчета p -value было обнаружено, что значением сходимости статистически не значимо ($p = 0.258$). Следовательно, данный результат свидетельствует в пользу аргумента о том, что МЭГ- и ЭЭГ-модальности, при помощи которых можно зарегистрировать уникальные осцилляции, и МЭГ в данном случае может послужить источником информации с преимуществом в виде наличия паттернов, отсутствующих на ЭЭГ-записи.

На данном этапе исследования производится попытка синтеза ЭЭГ-сигнала на основе реконструированных МЭГ-источников при помощи модифицированной нейросети на основе классической архитектуры UNET, реализованной в среде PyTorch [5]. Выбор данной архитектуры можно объяснить ее относительно высокой интерпретируемостью и легкостью, а также меньшей потерей данных в ходе преобразования сигнала с возможностью их сегментации на данные с более низкой размерностью. Модель представляет собой последовательность слоев свертки и подвыборки (max-pooling) для кодирования входящего сигнала, а также слоев повышения разрешения и свертки для декодирования. В архитектуре модели также используются кросс-связи между соответствующими кодирующими слоями, для передачи информации о признаках к более «высоким» уровням. Нелинейности (ReLU) были применены после каждого слоя свертки для улучшения способности модели к обучению. Размер полей восприятия (receptive fields) каждого слоя свертки был ограничен небольшими пространственными размерами. Временные рецептивные поля отсутствуют, так как каждый образец обрабатывается независимо. В дальнейшем предполагается подбор параметров для увеличений значений сходимости между предсказанными последовательностями стадий для разных типов сенсоров.

Выводы

В данном исследовании предложен и реализован многоступенчатый анализ физиологических маркеров сна с их последующим сопоставлением в пределах разных стадий. В ходе сопоставления маркеров на этапе расчета

сходимости и спектрального анализа были выявлены различия в стадиях N2, N1, R и W для разных типов сенсоров. В качестве причины данных различий можно привести аргумент в пользу различной чувствительности к регистрации топографий в виде результатов визуализации моделирования источников, регистрируемых при помощи ЭЭГ и МЭГ сенсоров. Неудачная попытка реконструкции ЭЭГ-сигнала на основе реконструированных МЭГ-источников может быть объяснена не совсем точным подбором инструмента в виде обратного оператора для данной процедуры. В частности, некоторые исследования указывают на то, что применение обратного оператора к данным с вызванными потенциалами (ВП) демонстрирует лучшие результаты в задачах реконструкции сигнала в связи с тем, что форма и временные характеристики ВП часто известны или могут быть легко оценены, что позволяет использовать предварительно обученные алгоритмы обратного оператора [6]. В качестве ограничений данного исследования можно выделить небольшую продолжительность записей наряду с небольшим размером выборки, а также снижение частоты дискретизации записей на некоторых этапах анализа данных с 1000 Гц и до 500–250 Гц. Программа дальнейших исследований будет включать в себя модификацию нейросетей с другими архитектурами, например SincNet [7] или генеративных моделей для повышения значений сходимости между последовательностями стадий сна, предсказанных на основании сигнала, записанного с разных сенсоров.

Список литературы

1. Ayoub A. и др. Grouping of MEG gamma oscillations by EEG sleep spindles // *NeuroImage*. 2012. V. 59, N 2. P. 1491–1500.
2. Weber F. D. и др. Coupling of gamma band activity to sleep spindle oscillations – a combined EEG/MEG study // *NeuroImage*. 2021. V. 224. P. 117452.
3. Faust O. и др. A review of automated sleep stage scoring based on physiological signals for the new millennia // *Comput. Methods Programs Biomed*. 2019. V. 176. P. 81–91.
4. Dehghani N., Cash S. S., Halgren E. Topographical frequency dynamics within EEG and MEG sleep spindles // *Clin. Neurophysiol*. 2011a. V. 122, N 2. P. 229–235.
5. Huang M. X. и др. Predicting EEG responses using MEG sources in superior temporal gyrus reveals source asynchrony in patients with schizophrenia // *Clin. Neurophysiol*. 2003. V. 114, N 5. P. 835–850.
6. He B. и др. Boundary Element Method-Based Cortical Potential Imaging of Somatosensory Evoked Potentials Using Subjects' Magnetic Resonance Images // *NeuroImage*. 2002. V. 16, N 3. P. 564–576.
7. Ravanelli M., Bengio Y. Speaker Recognition from Raw Waveform with SincNet // 2018 IEEE Spoken Language Technology Workshop (SLT). Athens, Greece: IEEE, 2018. P. 1021–102.

Научное издание

**XXVI МЕЖДУНАРОДНАЯ НАУЧНО-ТЕХНИЧЕСКАЯ
КОНФЕРЕНЦИЯ
"НЕЙРОИНФОРМАТИКА-2024"**

Редакторы и корректоры:

В. А. Дружинина, И. А. Волкова, Н. Е. Кобзева

Компьютерная верстка: *З.Б. Сохова, Г.А. Бесхлебнова,*

Подписано в печать 18.10.2024. Формат $60 \times 84 \frac{1}{16}$. Усл. печ. л. 16,12.

Уч.-изд. л. 15,6. Тираж 25 экз. Заказ № 110.

Федеральное государственное автономное образовательное учреждение
высшего образования «Московский физико-технический институт
(национальный исследовательский университет)»

141700, Московская обл., г. Долгопрудный, Институтский пер., 9

Тел. (495) 408-58-22, e-mail: rio@mipt.ru

Отпечатано в полном соответствии с предоставленным оригиналом-макетом
ООО «Печатный салон ШАНС» 127412, г. Москва, ул. Ижорская, д. 13, стр. 2