

МИНИСТЕРСТВО ОБРАЗОВАНИЯ РОССИЙСКОЙ ФЕДЕРАЦИИ
МИНИСТЕРСТВО РОССИЙСКОЙ ФЕДЕРАЦИИ ПО АТОМНОЙ ЭНЕРГИИ
МИНИСТЕРСТВО ПРОМЫШЛЕННОСТИ, НАУКИ И ТЕХНОЛОГИЙ
РОССИЙСКОЙ ФЕДЕРАЦИИ
РОССИЙСКАЯ АССОЦИАЦИЯ НЕЙРОИНФОРМАТИКИ
МОСКОВСКИЙ ИНЖЕНЕРНО-ФИЗИЧЕСКИЙ ИНСТИТУТ
(ГОСУДАРСТВЕННЫЙ УНИВЕРСИТЕТ)

НАУЧНАЯ СЕССИЯ МИФИ–2004

НЕЙРОИНФОРМАТИКА–2004

**VI ВСЕРОССИЙСКАЯ
НАУЧНО-ТЕХНИЧЕСКАЯ
КОНФЕРЕНЦИЯ**

**ЛЕКЦИИ
ПО НЕЙРОИНФОРМАТИКЕ**

Часть 1

По материалам Школы-семинара
«Современные проблемы нейроинформатики»

Москва 2004

УДК 004.032.26 (06)

ББК 32.818я5

М82

НАУЧНАЯ СЕССИЯ МИФИ-2004. VI ВСЕРОССИЙСКАЯ НАУЧНО-ТЕХНИЧЕСКАЯ КОНФЕРЕНЦИЯ «НЕЙРОИНФОРМАТИКА-2004»: ЛЕКЦИИ ПО НЕЙРОИНФОРМАТИКЕ. Часть 1. – М.: МИФИ, 2004. – 199 с.

В книге публикуются тексты лекций, прочитанных на Школе-семинаре «Современные проблемы нейроинформатики», проходившей 28–30 января 2004 года в МИФИ в рамках VI Всероссийской научно-технической конференции «Нейроинформатика-2004».

Материалы лекций связаны с рядом проблем, актуальных для современного этапа развития нейроинформатики, включая ее взаимодействие с другими научно-техническими областями.

Ответственный редактор

Ю. В. Тюменцев, кандидат технических наук

ISBN 5-7262-0526-X

© *Московский инженерно-физический институт
(государственный университет), 2004*

Содержание

Предисловие	6
 Я. Б. Казанович, В. В. Шматченко. Осцилляторные нейросетевые модели сегментации изображений и зрительного внимания	
модели сегментации изображений и зрительного внимания	15
Введение. Биологические предпосылки моделирования	16
Сегментация зрительных сцен	19
Модели сегментации: Параллельные процедуры	20
Модели сегментации: Последовательные процедуры	29
Модели зрительного внимания	41
Модель нейронных механизмов селективного внимания, основанная на временной корреляции нейронов	42
Модели формирования фокуса внимания	47
Модель селективного внимания и задачи зрительного поиска	54
Заклучение. Перспективы использования осцилляторных нейронных сетей при моделировании когнитивных функций мозга	58
Приложение. Типы нейронных сетей	63
Литература	64
 А. Ю. Дорогов. Быстрые нейронные сети: Проектирование, настройка, приложения	
Быстрые нейронные сети: Проектирование, настройка, приложения	69
Введение	70
Структурный анализ алгоритмов БПФ	72
Структурный синтез быстрых нейронных сетей	76
Формальный язык регулярных сетей	76
Грамматика предложений и морфология сети	78
Семантическая интерпретация предложений	80
Топологическое проектирование быстрых нейронных сетей	81
Топологии нейронных слоев	82
Графическая интерпретация топологий	83
Граничные условия и топологические матрицы	86
 УДК 004.032.26 (06) Нейронные сети	 3

Алгоритм обработки данных в БНС	89
Слабосвязанные сети	91
Системные характеристики БНС	94
Вычислительная эффективность	94
Пластичность	94
Способность к обучению	95
Перестраиваемые спектральные преобразования	98
Настройка перестраиваемых преобразований	99
Настройка на базис Уолша	100
Настройка на базис Фурье	103
Быстрое вейвлет-преобразование	104
Нейросетевая аппроксимация регулярных фракталов	106
Аналитическая форма регулярного фрактала	107
Дискретная аппроксимация фракталов	108
Фрактальная фильтрация сигналов	111
Фильтрация дискретных сигналов	112
Типы фрактальных фильтров	113
Фильтрация непрерывных сигналов	115
Приспособленные быстрые преобразования	116
Алгоритм приспособления	116
Приспособленные преобразования в нечетком пространстве	120
Спектральные приспособленные преобразования	122
Быстрые нейронные сети в квантовых вычислениях	126
Многомерные БНС	131
Литература	133

В. Г. Яхно. Нейроподобные модели описания динамических процессов преобразования информации	136
Введение	136
Модели с адаптацией параметров	138
Модели адаптивных распознающих систем	139
Модели взаимодействующих распознающих систем	144
Выводы	147
Литература	148

Н. Г. Ярушкина. Нечеткие нейронные сети с генетической настрой-

кой	151
Введение	152
Нечеткие системы	154
Нечеткие множества	154
Функции принадлежности: Определение и смысл	154
Нечеткие числа	156
Нечеткие интервалы	156
Операции с нечеткими множествами	156
Нечеткие отношения	161
Функции нечетких переменных	163
Нечеткие системы	166
Генетические вычисления	172
Генетические алгоритмы	172
Применение генетических алгоритмов	173
Стандартный ГА	174
Гибридные системы	176
Нечеткие нейронные сети	176
Структуры гибридных систем (ГС)	180
Заключение	196
Литература	196

ПРЕДИСЛОВИЕ

1. В этой книге (она выходит в двух частях) содержатся тексты лекций, прочитанных на Школе-семинаре «Современные проблемы нейроинформатики», проходившей 28–30 января 2004 года в МИФИ в рамках VI Всероссийской научно-технической конференции «Нейроинформатика–2004».

При отборе и подготовке материалов для лекций авторы и редактор следовали принципам и подходам, сложившимся при проведении трех предыдущих Школ (см. [1–5]). А именно, основной целью Школы было рассказать слушателям о современном состоянии и перспективах развития важнейших направлений в теории и практике нейроинформатики, о ее приложениях. При этом особенно приветствовались лекции *междисциплинарные*, лежащие по охватываемой тематике «на стыке наук», рассказывающие о проблемах не только собственно нейроинформатики (т. е. о проблемах, связанных с нейронными сетями, как естественными, так и искусственными), но и о взаимосвязях нейроинформатики с другими областями мягких вычислений (нечеткие системы, генетические и другие эволюционные алгоритмы и т. п.), с системами, основанными на знаниях, с традиционными разделами математики, биологии, психологии, инженерной теории и практики.

Основной задачей лекторов, приглашаемых из числа ведущих специалистов в области нейроинформатики и ее приложений, смежных областей науки, было дать живую картину современного состояния исследований и разработок, обрисовать перспективы развития нейроинформатики в ее взаимодействии с другими областями науки.

Помимо междисциплинарности, приветствовалась также и *дискуссионность* излагаемого материала. Как следствие, не со всеми положениями, выдвигаемыми авторами, можно безоговорочно согласиться, но это только повышает ценность лекций — они стимулируют возникновение дискуссии, выявление пределов применимости рассматриваемых подходов, поиск альтернативных ответов на поставленные вопросы, альтернативных решений сформулированных задач.

2. В программу Школы-семинара «Современные проблемы нейроинформатики» на конференции «Нейроинформатика–2003» вошли следующие восемь лекций¹:

1. Я. Б. Казанович, В. В. Шматченко. Осцилляторные нейросетевые модели сегментации изображений и зрительного внимания.

¹Первые четыре из перечисленных лекций публикуются в части 1, а оставшиеся четыре — в части 2 сборника «Лекции по нейроинформатике».

2. *В. Г. Яхно*. Нейроноподобные модели описания динамических процессов преобразования информации.
3. *А. Ю. Дорогов*. Быстрые нейронные сети: Проектирование, настройка, приложения.
4. *Н. Г. Ярушкина*. Нечеткие нейронные сети с генетической настройкой.
5. *А. А. Жданов*. О методе автономного адаптивного управления.
6. *Л. А. Станкевич*. Нейрологические средства систем управления интеллектуальных роботов.
7. *С. А. Терехов*. Нейро-динамическое программирование автономных агентов.
8. *Н. Г. Макаренко*. Как получить временные ряды из геометрии и топологии пространственных паттернов?

Основные темы, рассматриваемые в этих лекциях — нетрадиционные² нейросетевые модели и их возможные применения.

3. Лекция **Я. Б. Казановича** и **В. В. Шматченко** «Осцилляторные нейросетевые модели сегментации изображений и зрительного внимания» посвящена теме, рассмотрение которой было начато в лекции *игумена Феофана (Крюкова)* на Школе 2002 года [6]. Эта тема относится к направлению в теории нейронных сетей, активно развиваемому в настоящее время. Оно ориентируется на изучение динамических и осцилляторных аспектов функционирования мозга. В рамках данного направления предложен ряд гипотез и моделей нейронных сетей, позволяющих объяснить возникновение пространственно-временных паттернов нейронной активности, а также их значение для обработки информации. Один из основных элементов этих моделей — *принцип синхронизации*, введение которого позволяет надеяться на решение ряда задач биологии и психологии, трудных для традиционных подходов. Лекция *Я. Б. Казановича* и *В. В. Шматченко* посвящена анализу моделей, в которых принцип синхронизации используется для решения задач сегментации объектов на изображении и формирования фокуса внимания.

4. К первой лекции по своему подходу и прикладной ориентации в определенной степени примыкает лекция **В. Г. Яхно** «Нейроноподобные модели описания динамических процессов преобразования информации», которая

²Нетрадиционные в том смысле, что отличаются от мультиперсептронов, сетей Хопфилда и еще двух-трех общеизвестных и повсеместно используемых видов моделей.

продолжает развитие темы, начатой автором в лекции на Школе 2001 года [7]. Здесь под *нейроподобными* понимаются распределенные системы, состоящие из активных элементов с несколькими устойчивыми или квазиустойчивыми состояниями, при этом взаимодействие между указанными неравновесными элементами осуществляется за счет нелокальных пространственных связей. Модели подобного рода успешно применяются для описания процессов в однородных нейронных сетях сетчатки глаза живых объектов, коры некоторых отделов головного мозга и т. п. Одна из важных областей применения предлагаемых моделей — адаптивные распознающие системы, а также совокупности таких систем, которые открывают возможность анализировать на количественном уровне реакции животных при восприятии и осознания действующих на них информационных сигналов.

5. Лекция А. Ю. Дорогова «Быстрые нейронные сети: Проектирование, настройка, приложения» базируется на том параллелизме, который существует между алгоритмом *быстрого преобразования Фурье* (БПФ), играющим огромную роль в обработке сигналов и многослойными нейронными сетями. Этот параллелизм заключается в том, что алгоритмы БПФ имеют выраженную многослойную структуру, которая подобна структуре многослойных персептронов. По этой причине представляется вполне естественным использовать потенциал, накопленный в области БПФ, для построения нейронных сетей соответствующей архитектуры, которые предлагаются именовать быстрыми нейронными сетями (БНС). В лекции подробно рассматривается алгоритмическая сторона формирования БНС, которая иллюстрируется на многочисленных примерах, в первую очередь из области адаптивной фильтрации. Одна из интересных возможностей использования БНС обусловлена тем фактом, что структура БНС идентична структуре тензорных произведений векторных пространств. Из этого вытекает возможность использовать БНС для построения алгоритмов *квантовых вычислений*³. В лекции показано также, что парадигма БНС допускает многомерное обобщение, что открывает возможность создания быстродействующих классификаторов зрительных сцен.

³Тематика квантовых вычислений и квантовых нейронных сетей уже затрагивалась в лекциях Школы (см., в частности, лекцию А. А. Ежова на Школе 2003 года [8]), а также в материалах Рабочего совещания «Квантовые нейронные сети» на конференции «Нейроинформатика–2000» [8].

6. Лекция **Н. Г. Ярушкиной** «Нечеткие нейронные сети с генетической настройкой» посвящена рассмотрению еще одной плодотворной парадигмы, которая позволяет сочетать методы обучения, характерные для нейронных сетей, с вербализацией правил вывода, типичной для нечетких систем. *Нечеткие нейронные сети* (ННС), изучаемые в рамках данной парадигмы, представляют собой реализацию систем нечеткого логического вывода методами нейронных сетей. Подобного рода сети включают в себя слои, которые состоят из специальных *И* и *ИЛИ* нейронов. Настройку нечетких нейронных сетей предлагается выполнять с помощью методов генетической оптимизации. В числе приложений для ННС, упоминаемых в лекции, задачи управления объектами (динамическими системами). Подобного рода комбинированное использование нескольких элементов, традиционно объединяемых под общим наименованием «мягкие вычисления» (искусственные нейронные сети, нечеткие системы, эволюционные вычисления) ранее в рамках Школы рассматривалось, в частности в лекциях **Ю. И. Нечаева** [10, 11] применительно к различным аспектам задачи управления динамическими объектами.

7. В той или иной степени с решением задач управления связаны следующие три лекции (**А. А. Жданова**, **Л. А. Станкевича** и **С. А. Терехова**). Общей чертой подходов, предлагаемых в этих лекциях, является то, что все они нацелены на решение задачи управления поведением объекта в условиях значительной неопределенности. В частности, в лекции **А. А. Жданова** «О методе автономного адаптивного управления» предлагается концептуальная модель нервных систем, названная методом «*автономного адаптивного управления*» (ААУ). Считается, что рассматриваемая система представляет собой совокупность из объекта управления, управляющей системы (УС) (моделируемой нервной системы) и среды. Данная УС является интеллектуальной, поскольку обладает такими свойствами, как наличие аппарата эмоций, который мотивирует, определяет, направляет и оценивает поведение УС; внутренняя активность, направленная на расширение знаний, повышающих вероятность выживания; адаптивность и саморазвитие; индивидуальность. На этой концептуальной основе предлагаются конкретные решения, позволяющие строить практически действующие управляющие системы, описаны примеры нескольких практических приложений.

8. Еще одна работа, связанная с проблемами интеллектуального управления, это лекция **Л. А. Станкевича** «Нейрологические средства систем управления интеллектуальных роботов». В ней вводится класс средств,

именуемых автором *нейрологическими*, которые способны обучаться в реальном времени отображению сложных функций и процессов. База для построения таких средств — нейронные сети и логических системы на правилах. Показано, что на основе нейрологических средств можно строить когнитивные и актуаторные структуры, способные обучаться, формировать и реализовывать сложное рациональное поведение динамических объектов в среде. Одна из очевидных сфер применения предлагаемых средств — создание интеллектуальных роботов. В качестве примеров такого применения указываются система управления для антропоморфного робота, а также перспективный проект системы интеллектуального управления гуманоидного робота. Лекторы Школы не впервые обращаются к тематике управления роботами, в том числе и антропоморфными роботами. В частности, эти проблемы рассматривались в лекциях А. А. Фролова и Р. А. Прокопенко [12], а также А. И. Самарина [13] на Школе 2001 года.

9. Еще один из подходов к управлению сложными динамическими системами представлен в лекции **С. А. Терехова** «Нейро-динамическое программирование автономных агентов», которая, в определенной степени, продолжает тематику, рассматривавшуюся в лекциях С. А. Терехова на Школах 2002 года [14] и 2003 года [15]. Здесь речь идет о многошаговых процессах принятия решений и об оптимизации таких процессов. Среди многообразия подходов, существующих в настоящее время для решения задач подобного рода, можно выделить приближенные методы поиска адаптивной стратегии, которые основаны на аппроксимации функций оптимального поведения искусственными нейронными сетями. Этот подход, использующий кроме нейросетей также еще и методологию динамического программирования, принято именовать *нейро-динамическим программированием*. Важную роль в нейро-динамическом программировании играет обучение нейронной сети на основе локальных подкреплений (reinforcement learning), используемое вместо традиционного обучения с учителем. Показана связь задачи целевой адаптации автономного агента при обучении с подкреплением и задачи прогностического оптимального управления. Приведены примеры приложений методов нейро-динамического программирования.

10. Лекция **Н. Г. Макаренко** «Как получить временные ряды из геометрии и топологии пространственных паттернов? Математические аспекты анализа распределенных динамических систем» продолжает серию его выступлений междисциплинарного характера, которые состоялись на Школах

2002 года [16] и 2003 года [17]. Объектом рассмотрения являются *распределенные динамические системы*, которые не только демонстрируют сложное поведение во времени, но и имеют нетривиальную пространственную структуру. И если в лекции [17] проекциями динамики рассматриваемых систем в «Мир экспериментатора» являлись одномерные временные ряды, то теперь это — «мгновенные снимки», изображения или «сцены». Таким образом для распределенных динамических систем приходится иметь дело с двумя видами сложности: временной, которая отслеживается временными рядами каких-либо интегральных параметров, и пространственной, которая «кодируется» геометрией и топологией «сцен». Применительно к данной ситуации рассматриваются пространственные паттерны, а также некоторые современные математические инструменты для «арифметизации» и анализа таких объектов, основанные на методах таких разделов математики, как вычислительная геометрия, геометрические вероятности, интегральная геометрия, математическая морфология, стохастическая геометрия, вычислительная топология. Показано, что аппарат искусственных нейронных сетей вполне естественным образом взаимодействует со многими из этих элементов.

* * *

Для того, чтобы продолжить изучение вопросов, затронутых в лекциях, можно порекомендовать такой уникальный источник научных и научно-технических публикаций, как цифровая библиотека **ResearchIndex** (ее называют также **CiteSeer**, см. позицию [19] в списке литературы в конце предисловия). Эта библиотека, созданная и развиваемая отделением фирмы NEC в США, на конец 2002 года содержала около миллиона публикаций, причем это число постоянно и быстро увеличивается за счет круглосуточной работы поисковой машины.

Каждый из хранимых источников (статьи, препринты, отчеты, диссертации и т. п.) доступен в полном объеме в нескольких форматах (PDF, PostScript и др.) и сопровождается очень подробным библиографическим описанием, включающим, помимо данных традиционного характера (авторы, заглавие, место публикации и/или хранения и др.), также и большое число ссылок-ассоциаций, позволяющих перейти из текущего библиографического описания к другим публикациям, «похожим» по теме на текущую просматриваемую работу. Это обстоятельство, в сочетании с весьма эффективным полнотекстовым поиском в базе документов по сформулированному пользователем поисковому запросу, делает библиоте-

ку **ResearchIndex** незаменимым средством подбора материалов по требуемой теме.

Помимо библиотеки **ResearchIndex**, можно рекомендовать также богатый электронный архив публикаций [20], а также портал научных вычислений [18].

Перечень проблем нейроинформатики и смежных с ней областей, требующих привлечения внимания специалистов из нейросетевого и родственных с ним сообществ, далеко не исчерпывается, конечно, вопросами, рассмотренными в предлагаемом сборнике, а также в сборниках [1–5].

В дальнейшем предполагается расширение данного списка за счет рассмотрения насущных проблем собственно нейроинформатики, проблем «пограничного» характера, особенно относящихся к взаимодействию нейросетевой парадигмы с другими парадигмами, развиваемыми в рамках концепции мягких вычислений, проблем использования методов и средств нейроинформатики для решения различных классов прикладных задач. Не будут забыты и взаимодействия нейроинформатики с такими важнейшими ее «соседями», как нейробиология, нелинейная динамика (синергетика — в первую очередь), численный анализ (вейвлет-анализ и др.) и т. п.

Замечания, пожелания и предложения по содержанию и форме лекций, перечню рассматриваемых тем и т. п. просьба направлять электронной почтой по адресу tium@mai.ru Тюменцеву Юрию Владимировичу.

Литература

1. Лекции по нейроинформатике: По материалам Школы-семинара «Современные проблемы нейроинформатики» // III Всероссийская научно-техническая конференция «Нейроинформатика-2001», 23–26 января 2001 г. / Отв. ред. Ю. В. Тюменцев. – М.: Изд-во МИФИ, 2001. – 212 с.
2. Лекции по нейроинформатике: По материалам Школы-семинара «Современные проблемы нейроинформатики» // IV Всероссийская научно-техническая конференция «Нейроинформатика-2002», 23–25 января 2002 г. / Отв. ред. Ю. В. Тюменцев. Часть 1. – М.: Изд-во МИФИ, 2002. – 164 с.
3. Лекции по нейроинформатике: По материалам Школы-семинара «Современные проблемы нейроинформатики» // IV Всероссийская научно-техническая конференция «Нейроинформатика-2002», 23–25 января 2002 г. / Отв. ред. Ю. В. Тюменцев. Часть 2. – М.: Изд-во МИФИ, 2002. – 172 с.
4. Лекции по нейроинформатике: По материалам Школы-семинара «Современные проблемы нейроинформатики» // V Всероссийская научно-техническая

- конференция «Нейроинформатика-2003», 29–31 января 2003 г. / Отв. ред. Ю. В. Тюменцев. Часть 1. – М.: Изд-во МИФИ, 2003. – 188 с.
5. Лекции по нейроинформатике: По материалам Школы-семинара «Современные проблемы нейроинформатики» // V Всероссийская научно-техническая конференция «Нейроинформатика-2003», 29–31 января 2003 г. / Отв. ред. Ю. В. Тюменцев. Часть 2. – М.: Изд-во МИФИ, 2003. – 180 с.
 6. *Игумен Феофан (Крюков)*. Модель внимания и памяти, основанная на принципе доминанты // В сб.: «Лекции по нейроинформатике». Часть 1. – М.: Изд-во МИФИ, 2002. – с. 66–113.
 7. *Яхно В. Г.* Процессы самоорганизации в распределенных нейроноподобных системах: Примеры возможных применений // В сб.: «Лекции по нейроинформатике». – М.: Изд-во МИФИ, 2001. – с. 103–141.
 8. *Ежов А. А.* Некоторые проблемы квантовой нейротехнологии // В сб.: «Лекции по нейроинформатике». Часть 2. – М.: Изд-во МИФИ, 2003. – с. 29–79.
 9. Квантовые нейронные сети: Материалы рабочего совещания «Современные проблемы нейроинформатики» // 2-я Всероссийская научно-техническая конференция «Нейроинформатика-2000», 19–21 января 2000 г.; III Всероссийская научно-техническая конференция «Нейроинформатика-2001», 23–26 января 2001 г. / Отв. ред. А. А. Ежов. – М.: Изд-во МИФИ, 2001. – 104 с.
 10. *Нечаев Ю. И.* Нейросетевые технологии в бортовых интеллектуальных системах реального времени // В сб.: «Лекции по нейроинформатике». Часть 1. – М.: Изд-во МИФИ, 2002. – с. 114–163.
 11. *Нечаев Ю. И.* Математическое моделирование в бортовых интеллектуальных системах реального времени // В сб.: «Лекции по нейроинформатике». Часть 2. – М.: Изд-во МИФИ, 2003. – с. 119–179.
 12. *Фролов А. А., Прокопенко Р. А.* Адаптивное нейросетевое управление антропоморфными роботами и манипуляторами // В сб.: «Лекции по нейроинформатике». – М.: Изд-во МИФИ, 2001. – с. 18–59.
 13. *Самарин А. И.* Нейросетевые модели в задачах управления поведением робота // В сб.: «Лекции по нейроинформатике». – М.: Изд-во МИФИ, 2001. – с. 60–102.
 14. *Терехов С. А.* Нейросетевые аппроксимации плотности распределения вероятности в задачах информационного моделирования // В сб.: «Лекции по нейроинформатике». Часть 2. – М.: Изд-во МИФИ, 2002. – с. 94–120.
 15. *Терехов С. А.* Введение в байесовы сети // В сб.: «Лекции по нейроинформатике». Часть 1. – М.: Изд-во МИФИ, 2003. – с. 149–187.
 16. *Макаренко Н. Г.* Фракталы, аттракторы, нейронные сети и все такое // В сб.: «Лекции по нейроинформатике». Часть 2. – М.: Изд-во МИФИ, 2002. – с. 121–169.

17. Макаренко Н. Г. Эмбедология и нейропрогноз // В сб.: *«Лекции по нейроинформатике»*. Часть 1. – М.: Изд-во МИФИ, 2003. – с. 86–148.
18. Портал научных вычислений (Matlab, Fortran, C++ и т. п.)
URL: <http://www.mathtools.net/>
19. NEC Research Institute CiteSeer (also known as ResearchIndex) — Scientific Literature Digital Library.
URL: <http://citeseer.nj.nec.com/>
20. The Archive arXiv.org e-Print archive — Physics, Mathematics, Nonlinear Sciences, Computer Science.
URL: <http://arxiv.org/>

Редактор материалов выпуска,
кандидат технических наук *Ю. В. Тюменцев*

E-mail: tium@mai.ru

Я. Б. КАЗАНОВИЧ ¹⁾, В. В. ШМАТЧЕНКО ²⁾

¹⁾ Институт математических проблем биологии РАН,
г. Пущино, Московской обл.

E-mail: kasanovich@impb.psn.ru

²⁾ Пущинский государственный университет, г. Пущино

E-mail: vadschmatchenko@mail.ru

ОСЦИЛЛЯТОРНЫЕ НЕЙРОСЕТЕВЫЕ МОДЕЛИ СЕГМЕНТАЦИИ ИЗОБРАЖЕНИЙ И ЗРИТЕЛЬНОГО ВНИМАНИЯ

Аннотация

Описываются принципы построения и функционирования осцилляторных нейронных сетей для решения задач сегментации изображений и зрительного внимания. Рассмотрены параллельные и последовательные процедуры сегментации, а также модели внимания, ориентированные на селекцию объектов в соответствии с заданными приоритетами. Обсуждаются биологические и технические перспективы осцилляторных моделей когнитивных функций мозга.

Y. B. KAZANOVICH ¹⁾, V. V. SHMATCHENKO ²⁾

¹⁾Institute of Mathematical Problems in Biology,
Russian Academy of Sciences, Pushchino

E-mail: kasanovich@impb.psn.ru

²⁾State University of Pushchino, Pushchino

E-mail: vadschmatchenko@mail.ru

OSCILLATORY NEURAL NETWORK MODELS OF BINDING AND ATTENTION

Abstract

The principles of oscillatory neural network design and functioning are described in application to the problems of image segmentation and visual attention. Parallel and sequential procedures of segmentation and models of attention implementing different priorities in object selection are considered. Biological and technical perspective of oscillatory models of brain cognitive functions is discussed.

Введение. Биологические предпосылки моделирования

В теории нейронных сетей сформировалось и активно развивается отдельное направление, ориентирующееся на исследование динамических и осцилляторных аспектов функционирования мозга. Появился ряд гипотез и моделей нейронных сетей, объясняющих возникновение определенных пространственно-временных паттернов нейронной активности и их значение для обработки информации. Особая роль в этих исследованиях отводится синхронизации как основному принципу, с помощью которого можно надеяться решить ряд сложных задач нейробиологии и психологии. В данной работе будут рассмотрены модели, в которых принцип синхронизации используется для решения задач сегментации объектов на изображении и формирования фокуса внимания.

Хорошо известно, что в ЭЭГ мозга присутствует ритмическая активность в различных частотных диапазонах (дельта 1–3 Гц, тета 4–10 Гц, альфа 8–13 Гц, бета 14–30 Гц, гамма 30–70 Гц, высокочастотные колебания свыше 70 Гц). При этом характер колебаний в значительной мере коррелирован с психологическими состояниями исследуемого организма. Устойчивые паттерны ритмической активности были обнаружены в различных структурах мозга на уровне отдельных нейронов и нейронных популяций. Такие экспериментальные данные получены в первичных зонах зрительной и обонятельной коры, сенсомоторной коре, в таламусе, в гиппокампе и в других структурах [1]. Поскольку колебательная активность является типичной для многих нейронных структур, можно предположить, что информация в мозге кодируется как частотой колебаний и их фазовыми соотношениями, так и пространственным распределением осциллирующих нейронов. Признаки, используемые при кодировании информации о стимулах в первичных зонах коры, имеют различную природу. Например, в случае зрительных стимулов это могут быть геометрические, спектральные или динамические характеристики изображения. Кроме того, признаки могут отличаться модальностью в зависимости от сенсорного источника (например, зрительного или слухового), от которого поступает информация. Известно, что первичная обработка различных по своей природе или модальности признаков идет в специализированных нейронных структурах коры мозга. Только в ассоциативных зонах коры возникает представление о целостных объектах, вызвавших стимуляцию сенсорных систем. В связи с этим возникает вопрос: каковы нейронные механизмы, позволяющие сохранить информацию о принадлежности признаков к отдельным объектам,

и как осуществляется объединение нужных признаков в цельный образ объекта?

Традиционный подход к решению этой проблемы связан с локальным или распределенным представлением объектов с помощью специфичных клеток (“grandmother cells” — «клетки бабушки»), которые обеспечивают конвергентную интеграцию реакций первичных нейронных систем на стимулы. Такие клетки действительно обнаружены: например, у обезьян есть клетки, реагирующие на определенные лица. Тем не менее, вряд ли можно ожидать, что такое решение задач запоминания и распознавания является универсальным. Одна из трудностей в реализации гипотезы о grandmother cells состоит в том, что потенциально эти клетки должны быть чувствительны ко всем возможным сочетаниям признаков, что приводит к комбинаторному взрыву числа связей [2].

Другая трудность возникает при одновременном предъявлении на изображении нескольких объектов. Дело в том, что информация о разных типах признаков обрабатывается независимо в специализированных структурах коры мозга, и только на уровне ассоциативной коры эта информация должна быть объединена в представление каждого объекта. Таким образом, весь набор признаков, извлеченных из изображения, должен быть сегментирован с указанием того, какой признак какому объекту принадлежит, и «метка» принадлежности должна сохраняться на всем пути продвижения зрительной информации по коре. Это так называемая проблема интеграции признаков (binding problem) [2, 3].

Решение данной проблемы было предложено в работах [4, 5], где была сформулирована гипотеза, что согласование во времени активностей нейронных ансамблей играет существенную роль в информационных процессах в мозге. Основная идея состоит в том, что признаки объекта кодируются синфазной (когерентной) активностью нейронов в различных областях коры мозга. Когерентность служит в качестве метки, помечающей информацию, относящуюся к одному объекту.

Эта гипотеза получила определенное экспериментальное подтверждение [6, 7]. В экспериментах на первичных зонах зрительной коры была обнаружена как стимулоспецифичная колебательная активность нейронов, так и когерентность этой активности при соблюдении определенных условий. В частности, проводились эксперименты на кошках, которым предъявлялась светлая полоска, движущаяся вверх или вниз по темному экрану. Активность на уровне одиночных нейронов (multiunit activity) и нейронных ансамблей (local field potential) регистрировалась в зоне V1 первичной зри-

тельной коры двумя электродами. Были получены следующие результаты.

1. Нейроны отвечают периодической активностью с частотой в гамма-диапазоне при прохождении полоски через их рецептивное поле.
2. При одновременном пересечении полоской двух рецептивных полей активность соответствующих нейронов синфазна, даже если эти нейроны находятся друг от друга на расстоянии нескольких миллиметров.
3. Если полоску разорвать на две изолированных части, то синхронность ответа падает тем больше, чем больше величина разрыва.
4. Если на экране движутся две полоски в разных направлениях (одна вверх, другая вниз), то синхронизованного ответа не возникает даже в тот момент, когда эти полоски оказываются на одной прямой.

Таким образом, в случае 2 имеет место синхронный ответ на единичный объект, а в случаях 3 и 4 — несинхронный ответ на два разных объекта.

Эти эксперименты были восприняты нейробиологическим сообществом неоднозначно. Одна часть ученых увидела в них радужные перспективы для понимания основных механизмов работы мозга и включилась в работу по проведению аналогичных экспериментов на разных животных, в разных условиях опыта и с регистрацией синхронной активности в других структурах. Другая часть продолжает относиться к этим результатам скептически, объясняя их специфическими условиями экспериментов. Очевидно, важную роль в этом вопросе должно сыграть математическое моделирование, поскольку с его помощью с гораздо большей полнотой могут быть изучены как сами осцилляторные механизмы интеграции признаков, так и их возможная роль в решении когнитивных задач.

Предположение о том, что принципы группировки информации должны быть сходными для различных информационных процессов, происходящих в мозге, привело к попытке использовать принципы синхронизации нейронной активности для объяснения механизма формирования и переключения фокуса внимания. К этому подталкивают как данные об участии внимания в группировке признаков [8–10], так и недавно полученное прямое экспериментальное подтверждение связи внимания с синхронизацией нейронной активности [11–13].

В упрощенном виде обработку информации в мозге можно разделить на два относительно независимых уровня [9]. На *нижнем уровне* (называемом *предвниманием*) происходит выделение признаков и распознавание простых стимулов (представленных однотипными признаками). Осо-

бенностью данного уровня является параллельная обработка информации, а синхронизация может быть следствием непосредственного взаимодействия между нейронными ансамблями первичной коры. На *верхнем уровне* (реализуемом с помощью внимания) формируется общее представление о реальности. Характерной особенностью этого уровня является последовательная обработка информации — в каждый момент времени обрабатывается та информация, которая находится в фокусе внимания. Синхронизация на этом уровне может использоваться для интеграции всей информации от органов чувств, находящейся в фокусе внимания. Есть основание предполагать, что работа верхнего уровня идет под управлением специальных координирующих структур (таких как септо-гиппокампальная система и лобные доли коры).

В данной работе будут рассмотрены осцилляторные модели сегментации и зрительного внимания, разработанные, в основном, в последнее десятилетие. Модели сегментации обычно соответствуют нижнему уровню, а модели внимания — верхнему. Но это правило не является абсолютным. Некоторые модели сегментации используют последовательные процедуры, а модели внимания наоборот строятся на принципах параллельной обработки информации и без использования специальных управляющих элементов или структур. Так что граница между этими моделями не всегда резкая, а используемые идеи бывают весьма сходными, что и обуславливает целесообразность их совместного рассмотрения. Сводка общих принципов построения осцилляторных нейронных сетей дана в приложении.

Сегментация зрительных сцен

Анализ и понимание зрительной сцены является одной из трудных задач как для моделирования в нейробиологии, так и при разработке систем искусственного интеллекта. Во многих случаях первым шагом к такому анализу является сегментация — выделение из всего изображения фона и осмысленных объектов. Как правило, невозможно распознать тот или иной объект без предварительной (или проходящей одновременно) сегментации. Сегментация также необходима как условие эффективного обучения, позволяющего идентифицировать паттерны, возникновения которых можно ожидать в последующих сценах.

Отметим, что в терминах осцилляторных нейронных сетей сегментацию зрительной сцены на принципах синхронизации одновременно можно

рассматривать как решение задачи интеграции признаков объекта в цельный образ. Результатом сегментации должна явиться такая активность нейронной сети, при которой разным объектам входного изображения соответствуют ансамбли элементов сети, имеющие различную динамику: элементы, кодирующие определенный объект, работают синфазно, а активность элементов, кодирующих разные объекты, не коррелирована. Таким образом, синфазная активность оказывается той «меткой», с помощью которой идентифицируются потоки информации, относящиеся к тому или иному объекту.

Предположим, что сенсорный вход нейронной сети активирован набором одновременно предъявленных объектов. Формирование синфазно работающих нейронных ансамблей может происходить параллельно (одновременно для всех объектов) или последовательно (синфазные ансамбли возникают в определенной последовательности). Оба эти подхода представлены в рассматриваемых ниже моделях сегментации.

Модели сегментации: Параллельные процедуры

При разработке модели сегментации на основе синхронизации достаточно очевидной представляется идея, согласно которой элементы нейронной сети синхронизируются с помощью подходящим образом подобранных локальных связей. К сожалению, попытка реализации этой идеи приводит к ряду трудностей. Если связи между элементами малы, то их синхронизирующей силы недостаточно для быстрой и надежной синхронизации всех осцилляторов, кодирующих признаки одного объекта. Если же связи сделать сильными, то это может привести к синхронизации групп осцилляторов, кодирующих разные объекты.

Для преодоления этих трудностей было выдвинуто два предложения. Согласно первому предложению нужно использовать адаптирующиеся связи, а именно: сила связи между осцилляторами, работающими синхронно, должна увеличиваться, а между несинхронно работающими — убывать. Второе предложение состояло в том, чтобы сделать локальные связи между близко расположенными осцилляторами синхронизирующими, а связи между удаленными друг от друга осцилляторами — десинхронизирующими. Обе идеи приводят к сходному результату: если ансамбли когерентно работающих осцилляторов образуют изолированные кластеры, то когерентности между этими ансамблями не будет.

Использование адаптирующихся связей. В ряде работ используется эффект модификации эффективности синаптической связи за короткие интервалы времени, хотя возможный механизм такой модификации остается недостаточно изученным, в частности, неизвестно, является ли она следствием биофизических изменений в синапсах или отражает влияние других популяций нейронов. Модификация эффективности синаптических связей была использована в работах [14, 15].

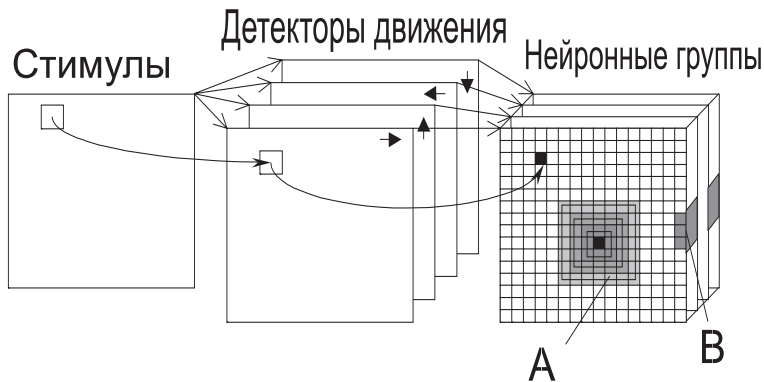


Рис. 1. Архитектура модели для работы с движущимися объектами
 А — область распределения возбуждающих связей в группе, В — возбуждающие связи между нейронами в разных группах с пересекающимися рецептивными полями.

Модель [14] рассчитана на работу с движущимися объектами, аппроксимированными прямолинейными отрезками. Такой подход обусловлен стремлением достаточно точно имитировать экспериментальные данные [6,7]. Архитектура сети представлена на рис. 1. Модель состоит из сенсорного поля (32×32 пикселей), четырех слоев детекторов элементарных признаков (32×32 клеток) и четырех слоев групп нейронов (16×16 групп). Каждая группа нейронов состоит из 40 возбуждающих и 20 тормозных клеток. Детекторы элементарных признаков, чувствительные к направлению движения пикселей на сенсорном поле в четырех направлениях (вверх, вниз, вправо, влево), распределены по разным слоям. Возбуждающие ней-

роны в группах получают топографически соответствующий им вход от 9 детекторов движения, так что активность нейронов в группе оказывается специфичной как по отношению к направлению движения, так и по отношению к ориентации стимулов. Кроме того, каждая возбуждающая клетка в группе получает 10 связей от других возбуждающих и 10 связей от тормозных клеток, а каждая тормозная клетка получает 10 связей от возбуждающих клеток. Что касается адаптирующихся связей, то каждая возбуждающая клетка получает два типа таких связей. Во-первых, это 30 связей от клеток в том же слое, распределенных по концентрическим областям из 7×7 групп (см. позицию *A* на рис. 1). Плотность связей падает экспоненциально с ростом расстояния от данной группы (это изображено в виде градации серого, центральная группа — черная). Во-вторых, существует еще 10 связей от возбуждающих клеток в группах с пересекающимися рецептивными полями и со специфичностью к направлению, отличающемуся на 90 градусов (см. позицию *B* на рис. 1).

В сети использовались импульсные нейроны с активностями, вычисляемыми как

$$s_i(t) = \left[\sum_j (c_{ij} + e_{ij}) s_j(t), \theta \right] + \eta + \alpha s_i(t-1),$$

где s_i, s_j — состояния нейронов i, j ($0 \leq s_{i,j} \leq 1$), t — время (одна итерация равна 1 мс реального времени), c_{ij} — базовая синаптическая эффективность и e_{ij} — быстро изменяющаяся эффективность ($0 \leq c_{ij} + e_{ij} \leq 1$), суммирование по j — суммирование по всем афферентным связям, η — шум, α — коэффициент релаксации. Квадратные скобки означают полулинейную функцию с порогом θ :

$$[x, \theta] = \begin{cases} x - \theta, & x > \theta, \\ 0, & x \leq 0. \end{cases}$$

Перед началом работы сети все синапсы имеют базовые значения c_{ij} , которые не равны нулю и не изменяются во время всего процесса; e_{ij} равны нулю. Для адаптирующихся связей величина e_{ij} изменяется в зависимости от значений пост- и пресинаптической активности в соответствии с правилом

$$e_{ij}(t+1) = (1 - \gamma) e_{ij}(t) + \delta \cdot F(e_{ij}) \cdot h(\bar{s}_i - \theta_i) \cdot h(\bar{s}_j - \theta_j),$$

где $\bar{s}_{i,j} = \lambda s_{i,j}(t) + (1 - \lambda) \bar{s}_{i,j}(t-1)$ — усредненные по времени активности клеток i и j ; δ — коэффициент усиления (параметр, регулирующий скорость

изменения), γ — константа затухания синаптической эффективности, $\theta_{i,j}$ — пороги срабатывания для пост- и пресинаптических нейронов, $F(t)$ — сигмоидальная функция, $h(t)$ — функция Хевисайда.

Стимулы представляли собой совокупности коротких вертикальных и горизонтальных полосок, движущихся в одном или разных направлениях с одинаковой скоростью. Перед началом движения сеть иницировалась случайным набором полосок. Это приводило к возникновению осцилляций со случайным распределением начальных фаз. В процессе экспериментов модель могла решать следующие задачи.

1. Выделять однородно движущийся объект на неоднородном фоне. Использовался протяженный объект, составленный из нескольких вертикальных полосок, движущихся вправо на фоне, составленном из по-разному ориентированных полосок, движущихся в случайных направлениях. При этом синхронизация групп нейронов, отвечающих частям объекта, наступала достаточно быстро — в течение пяти осцилляторных циклов (100 мс).
2. Выделять однородный объект, движущийся на однородном фоне.
3. Разделять два однородных объекта. В этом случае по реакции нейронов разделялись на два когерентных, но не коррелирующих множества.
4. Группировать элементарные признаки различных типов, принадлежащие разным объектам. Объекты представляют собой контуры (допустим, квадратов), составленные из вертикальных и горизонтальных полосок и движущиеся в разных направлениях. В данном случае корреляция реакций нейронов требует связей между слоями групп нейронов с разной специфичностью и пересекающимися рецептивными полями.

Усовершенствованная модель [15] позволяет включить в рассмотрение также цветные изображения. В сети используются группы осцилляторов, чувствительные к наличию в их рецептивных полях соответствующих признаков изображения (цвет, форма, движение). Архитектура этой модели достаточно сложна и воспроизводит во многих известных деталях структуру связей в мозге между различными областями, участвующими в обработке зрительной информации, включая область фронтальной коры, управляющую движением глаз.

Использование синхронизирующих и десинхронизирующих связей.

В работе [16] представлена модель, которая решает проблему сегментации, используя механизмы синхронизации и десинхронизации между нейронами, распределенными по различным признак-селективным модулям. Для реализации этой идеи авторы используют осцилляторы Вилсона-Коуэна. Каждый из осцилляторов состоит из возбуждающих (X_i) и тормозных (Y_i) элементов, имитирующих, соответственно, ансамбли возбуждающих и тормозных нейронов. Динамика осциллятора определяется системой дифференциальных уравнений

$$\begin{aligned}\tau_0 \frac{dX_i}{dt} &= -\alpha_X X_i + w_X F(Y_i(t - \tau_X)) + \eta_X(t) + I_i(t), \\ \tau_0 \frac{dY_i}{dt} &= -\alpha_Y Y_i + w_Y F(X_i(t - \tau_Y)) + \eta_Y(t),\end{aligned}$$

где $w_{X,Y} > 0$ — силы связи, $\eta_{X,Y}(t)$ — белый шум, $\alpha_{X,Y}$ — константа затухания, $I_i(t)$ — внешний сигнал от сенсорного поля, F — сигмообразная функция. Введение в модель задержек $\tau_{X,Y}$ обусловлено, во-первых, их присутствием в нервной системе (ввиду конечной скорости проводимости импульсов по аксонам и дендритам и из-за запаздывания, связанного с синаптической передачей), а также существенным влиянием задержек на фазовые соотношения между осцилляторами в нейронных сетях [17,18].

Связи между осцилляторами также организованы с временным запаздыванием. Связь от возбуждающего нейрона одного осциллятора к тормозному нейрону другого подобрана так, что синхронизует работу этих осцилляторов. Связи между возбуждающими нейронами осцилляторов являются десинхронизирующими. Как показано на рис. 2, дальное действие десинхронизирующих связей больше, чем синхронизирующих (рисунок отражает только общий принцип организации связей).

В признак-селективном модуле, специфичном к определенному признаку, каждый осциллятор чувствителен к стимулу, представляющему соответствующий признак в ассоциированном с данным осциллятором пикселе рецептивного поля. Модуль можно представить как трехмерную структуру: двумерное топографическое отображение зрительного поля и одномерное представление градаций признака (рис. 3). Значения признаков определяется для каждого пикселя изображения. Такое задание признаков отходит от экспериментальных условий [6,7], в которых в качестве признаков фигурировали свойства движущихся отрезков, но предоставляет более удобную возможность оперирования с яркостными и спектральными характери-

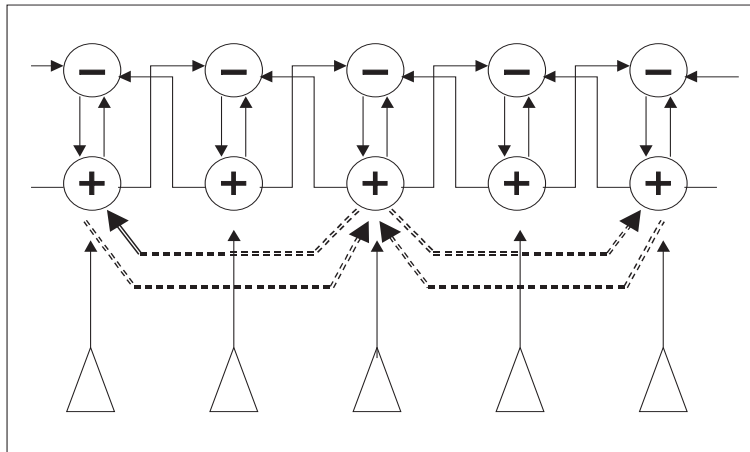


Рис. 2. Элементы осцилляторов (возбуждающие (+) и тормозные (-)), объединенные двумя типами связей — локальными синхронизирующими (сплошные линия) и удаленными десинхронизирующими (штрих-линии). Треугольниками обозначены входы от внешних стимулов.

ками объектов.

Рассмотрим архитектуру одного модуля. Каждый осциллятор имеет синхронизирующие связи с соседними осцилляторами, находящимися на расстоянии ближайшего соседства. Кроме того, осциллятор имеет десинхронизирующие связи с осцилляторами из слоев, лежащих на несколько уровней выше или ниже слоя, в котором находится данный осциллятор. В результате такого распределения синхронизирующих и десинхронизирующих связей осцилляторы, представляющие один топографически связанный объект, будут работать синфазно, так как стимулы, идущие от сенсорного поля и представляющие этот объект, непрерывны по зрительному пространству и по характеристике признака.

Следующим этапом организации архитектуры сети является введение модулей, отвечающих за различные признаки (цвет, яркость, движение и т.д.). Между модулями организованы синхронизирующие связи, которые объединяют каждый осциллятор в данном модуле со всеми осцилляторами



Рис. 3. Сеть с трехмерной структурой: двумерное топографическое отображение зрительного поля и одномерное представление градаций признака (ось, направленная вниз)

в любом другом модуле, имеющими такое же топографическое положение, как и данный осциллятор. Это согласуется с экспериментальными данными: активности нейронов с одинаковыми рецептивными полями имеют тенденцию к синхронизации. Такие соединения позволяют синхронизироваться ансамблям осцилляторов, распределенным по различным модулям.

Представленная архитектура сети позволяет корректно выполнять интеграцию признаков объектов, имеющих пересечение по какому-либо признаку. В случае двух пересекающихся объектов, представленных различными характеристиками какого-то признака (например, объекты имеют разный цвет), в соответствующем модуле (в данном случае модуле детекции цвета) происходит выделение двух непересекающихся синхронно работающих ансамблей осцилляторов. При этом десинхронизирующие связи делают активность в этих ансамблях некоррелированной, тем самым позволяя разделять объекты, изображения которых наложены одно на другое в зрительном поле.

Влияние биофизических свойств нейронов на синхронизирующие и десинхронизирующие связи. Значительно усовершенствованный вариант рассмотренной в предыдущем модели представлен в работе [19]. Ав-

торы использовали тот факт, что микроскопические биофизические свойства нейронов оказывают выраженный эффект на пространственно-временное взаимодействие нейронов, в частности, на синхронизацию нейронной активности. Так, калиевый ток утечки, проводимость и емкость мембраны нейрона не являются постоянными, а изменяются под влиянием множества факторов, таких как воздействие нейромодуляторов (в частности, ацетилхолина), распространение постсинаптического сигнала к соме, динамические свойства дендрита.

Основные отличия новой модели состоят в следующем. Сеть по-прежнему имеет модульный характер, но теперь каждый модуль сформирован не из осцилляторов Вилсона–Коуэна, а из четырех типов импульсных нейронов, различающихся типом нейротрансммиттера и типом рецепторов. Это возбуждающие нейроны, ГАМК-А и ГАМК-В тормозные нейроны и управляющие глутамат-эргические нейроны. Активность нейрона i в момент времени t задается как

$$s_i(t) = h(V_i(t) - \theta),$$

где $h(t)$ — функция Хевисайда, $V_i(t)$ — мембранный потенциал нейрона i в момент t , θ — порог, s_i — бинарная переменная, демонстрирующая наличие либо отсутствие потенциала действия в момент t .

Когда потенциал действия сгенерирован, мембранный потенциал уменьшается на фиксированное значение β — величину гиперполяризации, определенной для каждого типа нейронов.

Мембранный потенциал нейрона определяется как

$$V_i(t+1) = -\varepsilon V_i(t) + \sum_{j=1}^{N_i} S_j(t - \tau_{ij}) \cdot W_{ij} \cdot \exp(-A_i(t) D_{ij}),$$

где ε — скорость пассивного затухания мембранного потенциала; N_i — общее число возбуждающих и тормозных входов нейрона i . Поляризация сомы определяется интегрированием по соответствующим запаздывающим на время τ_{ij} активностям S_j нейронов j , взвешенным с соответствующими синаптическими эффективностями W_{ij} , ослабленными в соответствии с расстоянием до сомы D_{ij} и значением коэффициента затухания $A_i(t)$.

Основной особенностью модели является наличие управляющих нейронов, которые не оказывают прямого воздействия на мембранные потенциалы других популяций, но влияют на дендриты других нейронов. Так,

коэффициент затухания постсинаптического сигнала, распространяющегося в соме нейрона i , определяется модулирующим воздействием управляющих нейронов

$$A_i(t) = A^C - \sum_{j=1}^{M_i} S_j(t - \tau_{ij}) \cdot W_{ij},$$

где A^C — базовое значение коэффициента для данной популяции C , M_i — число модуляторных входов i -го нейрона.

Организация связей между популяциями выглядит следующим образом. Входные сигналы приходят на слой управляющих нейронов. Управляющие нейроны оказывают воздействие на другие популяции нейронов, отвечающих различным градациям признака и разному топографическому положению. В свою очередь, нейроны из других популяций воздействуют на мембранные потенциалы управляющих нейронов. Такое взаимодействие отличается размерами рецептивных полей и силами синаптической связи. Основную роль в организации связей играет размещение синапсов на дендритном дереве. Расстояния от синапсов до сомы в популяциях управляющих (У) и ГАМК-А тормозных нейронов зависят от расположения пре- и постсинаптических нейронов. Синапсы нейронов с близкими рецептивными полями и чувствительностью к одному и тому же признаку расположены близко к соме. Синапсы нейронов с удаленными рецептивными полями и/или чувствительностью к разным признакам расположены по периферии дендритного дерева. Такое размещение позволяет преобразование эффективности синаптических связей в зависимости от состояния модуляторной системы. Выделяются четыре режима взаимодействия нейронов:

- **«разобщенный»:** если $A^У$ и $A^{\text{ГАМК-А}}$ остаются при своих начальных значениях, большинство постсинаптических потенциалов оказываются подавлены;
- **«локальное взаимодействие»:** низкие значения модулирующих параметров $A^У$ и $A^{\text{ГАМК-А}}$ позволяют эффективно взаимодействовать ближайшим соседям;
- **«взаимодействие в колонке»:** при средних значениях, когда $A^У \approx 1$ и $A^{\text{ГАМК-А}}$ достигает 0, поведение управляющих клеток будет определяться возбуждающим внешним сигналом и тормозным постсинаптическим сигналом клеток ГАМК-А. При этом латеральное взаимодействие управляющих клеток ограничено;

- **«глобальное взаимодействие»:** при дальнейшем уменьшении модулирующих параметров все возбуждающие сигналы клеток, взаимодействующих с управляющими нейронами, будут вносить вклад в деполяризацию этих нейронов.

Кроме обычной для модели сегментации способности к синхронизации активностей нейронов, кодирующих различные объекты, в работе продемонстрированы две новые возможности обработки изображений. Во-первых, данная модель позволяет работать с движущимися объектами без специальных детекторов движения. При этом синхронизация активности наступает сразу при попадании объекта в рецептивные поля соответствующих нейронов. Во-вторых, модель чувствительна к контексту, в котором показаны объекты. Для иллюстрации этой способности приведем два примера.

В первом примере на изображении имеется два разнесенных в пространстве объекта, каждый из которых состоит из четырех полосок. При этом цвета полосок, составляющих один объект, различаются незначительно, в то время как разброс значений цвета по объектам значителен (рис. 4А). Для данной зрительной сцены, включающей широкий спектр цветов, выбирается режим «взаимодействия в колонке». Это приводит к правильной дискриминации двух объектов, составленных из одинаковых или близких (рис. 4А, пунктирная и сплошная кривые) цветов. Части, представляющие различные объекты, демонстрируют фазовое несоответствие (точечная кривая). Во втором примере имеется одно изображение, составленное из чередующихся полосок с близкими значениями цветов (рис. 4В). Сеть функционирует в условиях режима «локального взаимодействия». В этом случае сеть воспринимает сцену, как два несвязанных объекта, независимо от пространственного распределения частей этих объектов (рис. 4В, пунктирная и точечная кривые). Нейроны, представляющие разные цвета, не синхронизованы (рис. 4В, сплошная линия).

Модели сегментации: Последовательные процедуры

Рассмотрим теперь другой класс моделей, в которых сегментация одновременно предъявленных объектов осуществляется последовательно, один объект за другим. При этом последовательность объектов может быть случайной или детерминированной.

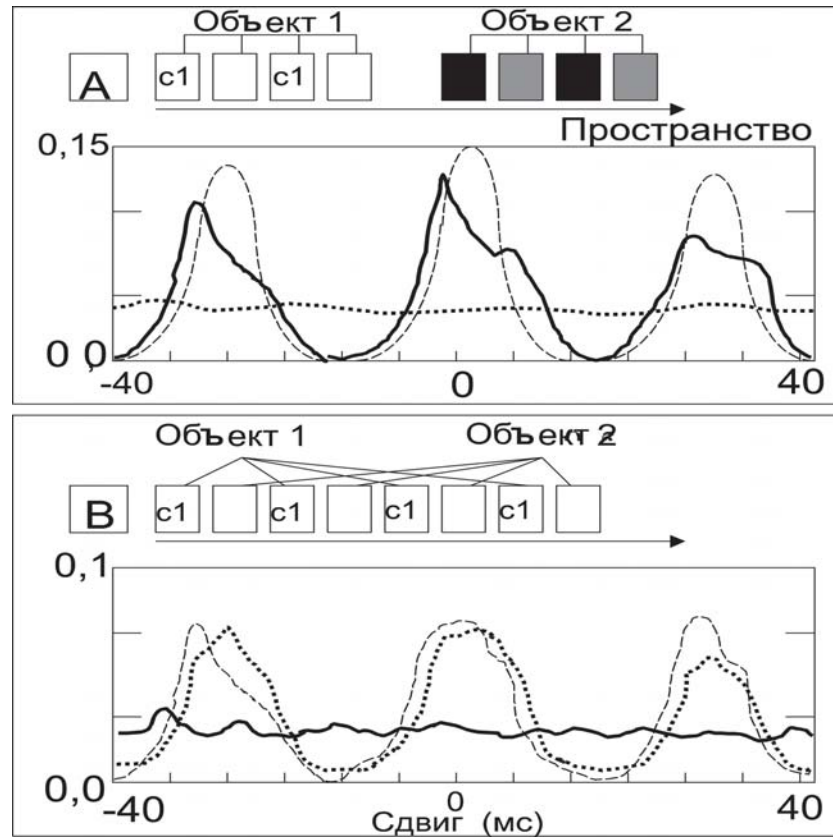


Рис. 4. Сегментация, чувствительная к контексту

А. Два разнесенных в пространстве объекта (вверху). Разброс значений цвета по объектам значителен, цвета в самих объектах изменяются незначительно. Кросс-корреляционные нормированные функции (внизу) для пар нейронов, представляющих идентичные цвета (штрихи), похожие цвета (сплошная линия) и сильно отличающиеся цвета (точки). **В.** Объекты, имеющие небольшие вариации в цветовом диапазоне. Корреляционные паттерны для пар нейронов, представляющих близкие (штрихи), удаленные (точки) и похожие цвета (сплошная линия).

Модель ассоциативной интеграции признаков. Пример модели с последовательным выбором объектов содержится в работах [20,21]. Эту модель можно рассматривать как адаптацию идей модели Хопфилда ассоциативной памяти применительно к осцилляторным нейронным сетям. Сеть составлена из импульсных нейронов: пирамидальных клеток и (тормозных) интернейронов. Осцилляторы сети сформированы из пары пирамидальная клетка – интернейрон. Архитектура связей в сети отражает тот факт, что пирамидальные клетки имеют как дальние, так и локальные связи, в то время как интернейроны имеют преимущественно локальные связи.

Генерация спайка i -м нейроном моделируется формальной переменной $S(t) = \{0, 1\}$ и имеет типичную для нейронного импульса ширину 1 мс. Временной шаг модели $dt = 1$ мс. Динамика нейрона определяется вероятностью генерации спайка в течение временного интервала dt и задается величиной мембранного потенциала h_i

$$P[S_i(t+dt) = 1 | h_i(t)] = \frac{1}{2} \left\{ 1 + \tanh[\beta(h_i(t) - \theta)] \right\},$$

где θ – порог и β – параметр, отвечающий за внешний шум. Мембранный потенциал включает три компонента

$$h_i(t) = h_i^{syn}(t) + h_i^{ext}(t) + h_i^{ref}(t),$$

где h_i^{syn} – сумма синаптических входов от всех других нейронов (см. ниже), h_i^{ext} – внешний сигнал, h_i^{ref} – формальное описание рефрактерности нейрона

$$h_i^{ref}(t) = \begin{cases} -R, & t_F \leq t \leq t_F + \tau_{ref}, \\ 0, & \text{в противном случае,} \end{cases}$$

где $R \gg 1$: генерация спайка не допускается в течение интервала τ_{ref} после последнего спайка, имевшего место в момент t_F .

Схема связей в сети из N осцилляторов выглядит следующим образом. Каждая пирамидальная клетка связана со своим тормозным партнером прямой и обратной связью со значением J^{inh} . Пирамидальные клетки, связанные между собой по принципу «все на все», формируют сеть, способную запомнить q объектов с помощью настройки весов J_{ij} связей по правилу Хебба

$$J_{ij} = \frac{2}{(1-a^2)N} \sum_{\mu}^q \xi_i^{\mu} (\xi_j^{\mu} - a), \quad (1)$$

где каждый объект состоит из независимо распределенных величин ξ_i^μ , принимающих значения $\xi_i^\mu = \pm 1$ с вероятностью $(1 \pm a)/2$, где $a = \langle \xi_i^\mu \rangle$ — средняя активность по всем объектам. Согласно (1), осцилляторы, принадлежащие одному объекту, связываются синхронизирующими связями, а осцилляторы, принадлежащие разным объектам — десинхронизирующими.

Каждая клетка i получает сигнал от клетки j с некоторым запаздыванием Δ_i^{ax} . Величины таких задержек случайно распределены в интервале от $\Delta_{\min}^{ax}$ до $\Delta_{\max}^{ax}$. После этого формируется постсинаптический потенциал с амплитудой, зависящей от J_{ij} , и откликом $\varepsilon(\tau)$. Суммарный вклад в формирование потенциала пирамидальной клетки i определяется следующим образом:

$$h_i^{syn}(t) = \sum_{j=1}^N J_{ij} \sum_{\tau=0}^{\infty} \varepsilon(\tau) S_j(t - \tau - \Delta_i^{ax}) + h_i^{inh}(t),$$

где $\varepsilon(t) = (t/\tau_{ex}^2) \exp(t/\tau_{ex})$ — функция отклика возбуждающего синапса с параметром $\tau_{ex} = 2$ мс, $h_i^{inh}(t)$ — вклад тормозного партнера нейрона i .

Динамика тормозного нейрона не моделируется детально, а лишь учитывается его вклад в формирование постсинаптического потенциала пирамидальной клетки. Этот вклад представляет собой ответ тормозного нейрона на его возбуждение пирамидальным нейроном, приходящее с задержкой Δ_i^{inh} ($\Delta_{\min}^{inh} \leq \Delta_i^{inh} \leq \Delta_{\max}^{inh}$),

$$h_i^{inh}(t) = \sum_{\tau=0}^{\tau_{\max}} \eta(t) S_i(t - \tau - \Delta_i^{inh}),$$

где $\eta(t) = J^{inh} \exp(t/\tau_{inh})$ — функция отклика с $\tau_{inh} = 6$ мс. Граница суммирования $\tau_{\max}$ подвижна и выбирается так, чтобы оборвать суммирование с момента, когда компоненты суммы становятся слишком малы.

Кроме обычной способности сети — восстанавливать из памяти образы сохраненных объектов при подаче на вход зашумленных образов, сеть демонстрирует также возможность сегментации зрительной сцены и ассоциативной интеграции признаков. Если на вход такой сети подать одновременно несколько паттернов, принадлежащих разным объектам, которые были использованы при формировании связей, то в сети установится следующая динамика. Синфазная активность осцилляторов, представляющих один объект, через некоторое время сменится синфазной активностью осцилляторов, представляющих другой объект. Это происходит из-за того,

что активный в данный момент времени паттерн подавляет активность осцилляторов в других паттернах, но активность самого этого паттерна через некоторое время подавляется возрастающей активностью входящих в него тормозных нейронов. Порядок активизации разных паттернов случаен и определяется шумом.

Продemonстрировать ассоциативную интеграцию признаков позволяет сеть с иерархической (двухурвневой) архитектурой. Рассмотрим K групп сетей с архитектурой, описанной выше. Будем называть их «колонками» (использовались колонки с 4000 осцилляторов в каждой). Колонки на нижнем уровне иерархии имеют прямые и обратные связи с колонкой высшего уровня, но при этом отсутствуют какие-либо связи между самими колонками нижнего уровня (рис. 5А). Каждая колонка запоминает паттерны объектов в соответствии с правилом (1). Вместе с этим, прямые и обратные связи, объединяющие два уровня, запоминают комбинации паттернов на нижнем уровне, ассоциируя их со специфичным паттерном на высшем уровне. Это приводит к прямым

$$J_{0i,kj} = g_0 \frac{2}{(1-a^2)N} \sum_{\langle \mu, \mu' \rangle} \xi_{0i}^{\mu} (\xi_{kj}^{\mu'} - a)$$

и обратным связям

$$J_{ki,0j} = g_k \frac{2}{(1-a^2)N} \sum_{\langle \mu, \mu' \rangle} \xi_{ki}^{\mu'} (\xi_{0j}^{\mu} - a).$$

Здесь k — номер колонки (номер 0 имеет колонка верхнего уровня), $\langle \mu, \mu' \rangle$ означает комбинацию паттернов μ' в различных колонках нижнего уровня, которые ассоциируются с паттерном μ на высшем уровне. Нормировка g_k выбрана так: $g_0 = (1+K)^{-1}$; $g_k = 1$ для прямой и $g_k = 0.5$ — для обратной связи.

Рассмотрим следующий пример. В качестве объектов, запоминаемых сетью, будут выступать два «слова» — «САТ» и «ТНЕ», состоящих из трех букв (признаков) каждое. Теперь колонке с номером 2 предъявляется зашумленная суперпозиция паттернов «А» и «Н» (рис. 5В). В зависимости от стимулов в других колонках в этой колонке может быть восстановлен образ «Н», синхронизированный с «Т» и «Е», в то время как образ «А» восстанавливается с задержкой, обусловленной сдвигом фаз. Если буквы в колонках 1 и 3 поменять на «С» и «Т», то сеть синхронизирует с ними «А» и восстанавливает слово «САТ».

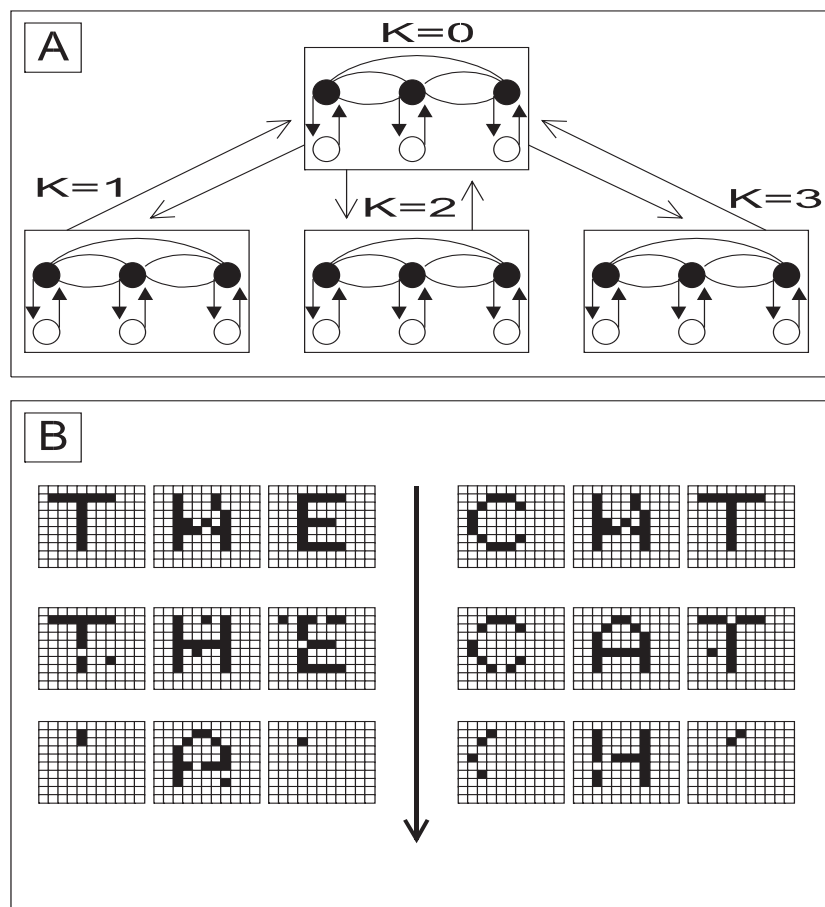


Рис. 5. Ассоциативное связывание признаков
А. Иерархическая структура сети. **В.** Демонстрация интеграции признаков, чувствительной к контексту.

Нейронная сеть с общим ингибитором. Десинхронизирующие связи в той или иной форме используются во многих моделях сегментации. Рассмотренная в предыдущем подразделе сеть содержит многочисленные дальние десинхронизирующие связи, которые могут связать любую пару осцилляторов. Такой сложной и биологически мало правдоподобной архитектуры связей можно избежать, если использовать в качестве источника десинхронизации небольшое число тормозных нейронов, каждый из которых соединен со своей группой возбуждающих нейронов.

В работе [22] был предложен механизм, позволяющий различать на зрительном поле несколько объектов и фон, на котором они представлены. В рассматриваемой сети осцилляторы чувствительны к локализованным входам от зрительного поля и кодируют локальный тип признака объекта. Все нейроны, представляющие один объект, работают с нулевым сдвигом фаз. Нейроны, принадлежащие разным объектам, работают либо с некоррелированным уровнем активности, либо не в фазе, что зависит от уровня шума в сети.

Взаимное ингибирование между двумя блоками сети, которые чувствительны к различным типам признаков в одном местоположении, и торможение между блоками сети, которые чувствительны к одному свойству, но имеют разное местоположение, а также глобальное ингибирование используются для поддержания необходимого уровня конкуренции между осцилляторами и для предотвращения глобальной синхронизации всей сети.

Эта модель была усовершенствована в работе [23], в которой разработана двумерная сеть LEGION (Locally Excitable, Globally Inhibitory Oscillator Network), работающая с использованием одного тормозного нейрона (рис. 6). Рассмотрим организацию и работу этой сети подробнее.

Сеть состоит из релаксационных осцилляторов Термана-Ванга, которые являются модифицированной версией осциллятора Ван-дер-Поля. Каждый осциллятор

$$\begin{aligned}\frac{dx_i}{dt} &= 3x - x^3 + 2 - y_i + I_i + S_i + \eta(t), \\ \frac{dy_i}{dt} &= \varepsilon \left[\alpha \left(1 + \tanh\left(\frac{x_i}{\beta}\right) \right) - y_i \right].\end{aligned}$$

представляет собой петлю прямой и обратной связи между возбуждающим x_i и тормозным y_i элементами. Здесь I_i — внешняя стимуляция, η — шум, $\varepsilon, \alpha, \beta$ — параметры ($\varepsilon \ll 0$ — малый параметр), S_i — сигнал, приходящий от других осцилляторов сети.

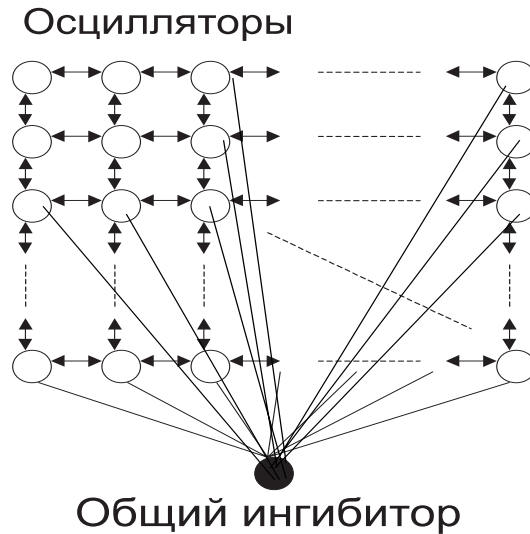


Рис. 6. Двумерная осцилляторная сеть LEGION

Рассмотрим изолированный осциллятор без шума и с постоянным внешним сигналом I . В фазовом пространстве (x, y) изоклина $dx/dt = 0$ представляет собой кубическую функцию, а изоклина $dy/dt = 0$ — сигмоидальную функцию. Если $I_i > 0$, изоклины пересекаются в единственной неустойчивой стационарной точке. В результате система порождает предельный цикл, который представляет собой смену периодов с относительно высокими и относительно низкими значениями x . Эти периоды называются фазой покоя и активной фазой, соответственно. В случае $I_i < 0$ изоклины имеют пересечение в устойчивой стационарной точке. В этом случае периодическое решение отсутствует.

При включении осцилляторов в сеть сигналы, приходящие на осцилляторы, имеют вид

$$S_i = \sum_{k \in N(i)} W_{ik} h(x_k - \theta_x) - W_z h(z - \theta_z),$$

где h — функция Хевисайда, $N(i)$ — осцилляторы, соседние с осциллятором i (в простейшем случае $N(i)$ включает четырех ближайших соседей, как это

показано на рис. 6), W_{ik} — параметры связи с соседними осцилляторами, θ_x, θ_z — пороги, W_z — вес тормозной связи от общего ингибитора z , чья активность определяется как

$$\frac{dz}{dt} = \mu(\sigma_\infty - z),$$

где параметр μ задает скорость реакции тормозного нейрона; $\sigma_\infty = 1$, если $x_i > \theta_z$ хотя бы для одного осциллятора, и $\sigma_\infty = 0$ в любом другом случае.

Особенностью сети является зависимость величины возбуждающих связей W_{ik} от внешнего сигнала. Динамические веса W_{ik} формируются на основе перманентных весов T_{ik} согласно следующей динамике:

$$\begin{aligned} \frac{du_i}{dt} &= \lambda(1 - u_i)I_i - \nu u_i, \\ \frac{dW_{ik}}{dt} &= T_{ik}u_iu_k \sum_{k \in N(i)} W_{ik} - W_{ik} \sum_{j \in N(i)} T_{ij}u_iu_j. \end{aligned}$$

Параметры этой системы подобраны так, что когда $I_i > 0$, то $u_i \rightarrow 1$. В случае, если осциллятор не стимулирован ($I_i < 0$), то $u_i \rightarrow 0$. Таким образом, если активность осциллятора достаточно велика, то его воздействие на окружающие осцилляторы возрастает. Благодаря этому достигается надежная синхронизация в группе осцилляторов, представляющих один объект.

В результате конкуренции, вызванной тормозным нейроном, «побеждает» тот ансамбль осцилляторов, активность которого растет быстрее, в то время как активность других осцилляторов будет подавлена. Через какое-то время активность ансамбля подавляется тормозным нейроном, после чего сеть активирует другой ансамбль осцилляторов. Какой именно ансамбль осцилляторов активируется в тот или иной интервал времени, зависит от уровня шума, вследствие чего последовательность, в которой активируются различные ансамбли, имеет случайный характер.

В работе [24] сеть LEGION применялась для решения задач сегментации объектов на реальных медицинских изображениях и показала достаточно хорошие результаты. В качестве объектов выступали медицинские снимки, полученные с помощью ЯМР. Каждый осциллятор сети соответствовал одному пикселю изображения, представленному в градациях серого цвета. В результате сегментации изображение разделялось на несколько основных фрагментов и фон. Серьезной проблемой, связанной с работой сети, является ее ограниченная сегментирующая способность. При фиксированных параметрах сеть может различить только от 5 до 11 объектов.

Это число зависит от отношения времени, в течение которого отдельный осциллятор находится в фазе покоя, и времени, в течение которого он находится в активной фазе. Если это отношение становится слишком большим, правильный ход сегментации нарушается.

Сегментация с помощью коррелирующих шумов. Этот подход к сегментации с использованием небольшого числа тормозных нейронов был реализован в работе [25]. Идея модели состоит в том, что признаки одного объекта помечаются одним и тем же шумом, в то время как шумы, помечающие признаки разных объектов, не коррелированы. Теоретическая основа для такого использования шума состоит в предположении, что отдельный объект кодируется в локально ограниченной области первичной зоны зрительной коры, которая подвергается воздействию сильно коррелированного шума. В то же время, шум в удаленных друг от друга областях должен быть некоррелированным.

Как и в рассмотренных выше моделях, признаки объектов (форма, цвет и т. п.) воспринимаются различными кортикальными модулями (рис. 7А). Каждый модуль состоит из нескольких популяций возбуждающих нейронов и популяции тормозных нейронов. Паттерны воспринимаемых признаков (например, для модуля, отвечающего за форму, это могут быть квадрат, треугольник и т. д.) хранятся в популяциях возбуждающих нейронов. Взаимодействие между различными модулями (которое обуславливает правильную интеграцию формы объекта с его цветом) происходит только через тормозные нейроны. Сеть построена из модифицированных элементов типа Вилсона–Коуэна (напомним, что каждый элемент Вилсона–Коуэна заменяет собой популяцию нейронов, тормозных или возбуждающих). Модель состоит из модулей, каждый из которых содержит набор возбуждающих элементов и один тормозный элемент. Внутри модуля возбуждающие элементы связаны только с тормозным элементом. Модули взаимодействуют между собой через тормозные элементы.

Рассмотрим для определенности два модуля, каждый из которых содержит три возбуждающих элемента. Пусть первый модуль отвечает за форму объекта, а составляющие его элементы кодируют, соответственно, такие признаки формы, как «круг», «треугольник», «квадрат» (можно представлять себе, что возбуждающий элемент — это колонка неокортекса). Пусть второй модуль отвечает за спектральные характеристики, а содержащиеся в нем возбуждающие элементы кодируют такие признаки, как «красный»,

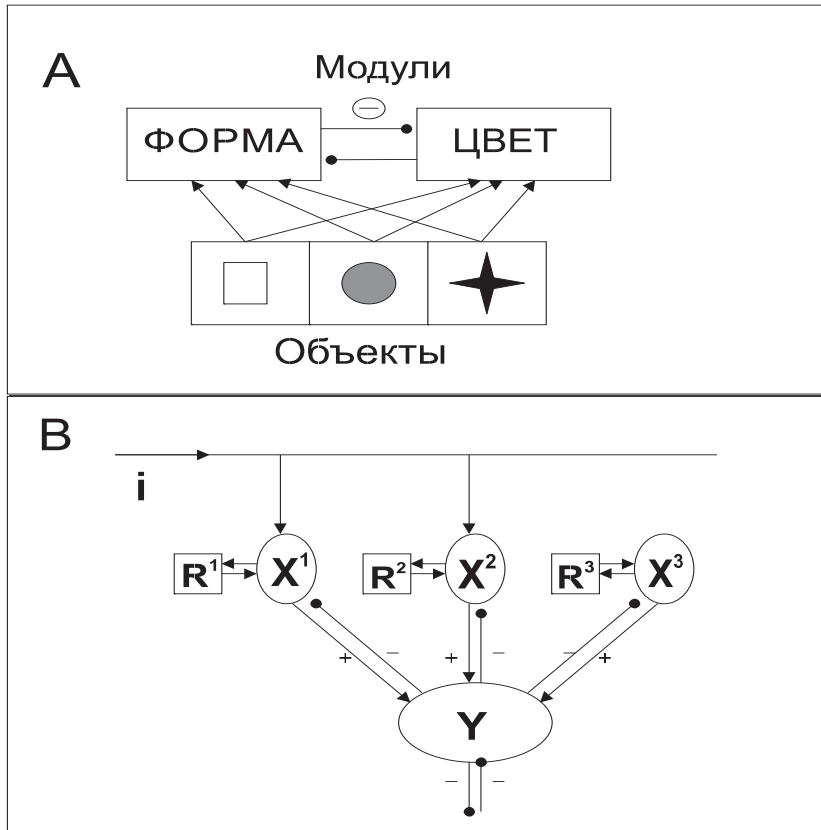


Рис. 7. Сегментация с помощью коррелирующих шумов

А. Схематическое изображение задачи объединения сегментов. Примеры объектов. Знаком (–) обозначены тормозные связи между двумя модулями.
В. Блок-схема модели.

«зеленый», «синий». Внешние сигналы приходят на возбуждающие элементы. Например, одновременное предъявление зеленого круга и красного квадрата приведет к подаче внешних сигналов на первый и третий элементы первого модуля и первый и второй элементы второго модуля.

Формально динамика модели описывается следующим образом. Пусть x_i^m — активность возбуждающих элементов в i -ом модуле; y_i — активность тормозного элемента в i -ом модуле. Особенность элементов состоит в том, что пороги возбуждающих элементов являются динамическими переменными r^m . Система уравнений для случая двух модулей имеет вид

$$\begin{aligned}\frac{dx_{1,2}^m}{dt} &= -x_{1,2}^m + F(Ax_{1,2}^m - By_{1,2} - \theta^E - br_{1,2}^m + i_{1,2}^m), \\ \frac{dy_{1,2}}{dt} &= -y_{1,2} + F(C \cdot X_{1,2} - Dy_{1,2} - \theta^I - \lambda y_{2,1}), \\ \frac{dr_{1,2}^m}{dt} &= (c - 1)r_{1,2}^m + x_{1,2}^m, \\ X_{1,2} &= \sum_m x_{1,2}^m.\end{aligned}$$

Здесь $c < 1$, θ^E и θ^I — фиксированные значения порогов для возбуждающих и тормозных нейронов, $i_{1,2}^m$ — внешний сигнал. Нижние индексы указывают на два различных модуля, которые объединяются через член $\lambda y_{2,1}$. На рис. 7В схематически представлены популяции и связи, отвечающие переменным, входящим в указанные дифференциальные уравнения. На рисунке представлен один из модулей с тремя кортикальными колонками. При $\lambda = 0$ каждый модуль демонстрирует феномен сегментации: чередование активностей популяций, кодирующих разные признаки, обеспечивается тормозными элементами.

Для интеграции сегментов информации от двух объединенных модулей ($\lambda \neq 0$) на них вместе с постоянной составляющей внешнего входного сигнала подается шум. Используемый входной сигнал приобретает вид

$$i^m(t) = 0.1 + 0.1[\rho^m(t) - 0.5],$$

где ρ^m — случайная величина, распределенная между 0 и 1.

Динамика активности элементов сети имеет колебательный характер и организована следующим образом. Допустим, что на вход сети подаются изображения красного квадрата и синего треугольник. Тогда, как показывают результаты имитационного моделирования, активности популяций

$x_1^{\text{КРАСНЫЙ}}$ и $x_2^{\text{КВАДРАТ}}$ синхронизированы с нулевым сдвигом фаз, активности $x_1^{\text{СИНИЙ}}$ и $x_2^{\text{ТРЕУГОЛЬНИК}}$ также синхронизированы с нулевым сдвигом фаз. В то же время пары $(x_1^{\text{КРАСНЫЙ}}, x_1^{\text{СИНИЙ}})$ и $(x_1^{\text{КВАДРАТ}}, x_1^{\text{ТРЕУГОЛЬНИК}})$ работают в противофазе.

Если модули кодируют по три признака, а сети предъявляются, соответственно, три объекта, то правильная сегментация также возможна, хотя это и занимает несколько больше времени. В случае же, если количество кодируемых признаков превышает три, возникает вероятность ложной интеграции признаков, что интерпретируется в терминах иллюзий, известных из психологических экспериментов. Типичным примером являются эксперименты, когда испытуемому предъявляется экран, на который проецируются (с коротким временем экспозиции) образы объектов с различными цветом и формой. За счет этого достигается рассеивание внимания по всему экрану. Тогда возможен случай, когда испытуемый сообщает о появлении на экране зеленого квадрата, в то время как представлены красный квадрат и зеленый круг. Таким образом, правильная интеграция признаков требует фокусировки внимания.

Модели зрительного внимания

Ограниченная работоспособность зрительной системы не позволяет ей одновременно анализировать несколько различных объектов, представленных в сложной зрительной сцене. Под селективным вниманием подразумевается процесс, с помощью которого из зрительного поля отбираются сегменты информации для более детальной дальнейшей переработки. Зрительное внимание иногда ассоциируют с ограниченной областью, которая может варьировать в размере и сканирует объекты в зрительном поле со скоростью 30–50 мс. Предполагается, что внимание фильтрует информацию, отдавая предпочтение объектам в фокусе внимания и подавляя информацию об объектах за пределами фокуса внимания. На самом деле, столь высокая скорость, с которой внимание как будто бы переходит с одного объекта на другой, вызывает сомнения, поэтому вопрос о том, в каких случаях внимание работает путем последовательного перебора объектов, а в каких путем параллельного селектирования входной информации, является предметом исследования и незавершенных дискуссий [26].

Поведение, соответствующее зрительному вниманию, можно наблюдать и количественно оценивать на нейронном уровне. В работе [27] пред-

ставлены результаты экспериментов, проводившихся на обученных макаках. Было обнаружено, что средняя скорость генерации спайков нейронами зависит от того, на чем животное сосредотачивает внимание: на стимуле, расположенном в рецептивном поле нейрона, или на стимулах, лежащих вне рецептивного поля нейрона. В первичных зонах зрительной коры V1 и V2 спайковая активность под действием внимания модифицируется — либо возрастает, либо подавляется. Данный эффект проявлялся только для стимулов, близких к ориентации, на которую реагирует клетка, и только в присутствии конкурирующего стимула с другой ориентацией. С другой стороны, в области V4 реакция клеток не зависела от ориентации стимулов. Более того, модулирующий эффект внимания наблюдался для изолированных стимулов.

Другие исследования строились на экспериментах, в которых более чем один стимул размещался в рецептивном поле нейрона из области V4 [28, 29]. Когда два различных объекта (скажем, красный и зеленый), размещались в рецептивном поле нейрона, селективного к красному цвету, нейрон реагировал интенсивно, если внимание животного сосредотачивалось на красном объекте. Если же внимание фокусировалось на зеленом объекте, то нейрон реагировал слабее.

Модель нейронных механизмов селективного внимания, основанная на временной корреляции нейронов

В работе [30] представлена модель зрительного внимания, предназначенная для объяснения результатов экспериментов [28, 29] и основанная на временной корреляции групп нейронов. Модель базируется на гипотезе о том, что под воздействием фокуса внимания в первичной зоне зрительной коры (V2) происходит временная модуляция серий спайков, генерируемых нейронами, являющимися пресинаптическими к нейронам из области V4 (внутри фокуса внимания). При этом, в согласии с экспериментом, не происходит изменения средней скорости генерации спайков нейронами V2. Внимание проявляется лишь в том, что нейроны V2, лежащие в фокусе внимания, начинают одновременно генерировать спайки чаще, чем нейроны вне фокуса внимания. Синхронизация спайков на уровне V2 и является меткой, позволяющей распознать признаки объекта в фокусе внимания на уровне V4. На высших уровнях (V4 и далее) сигналы от выделенных вниманием нейронов начинают конкурировать с сигналами от других групп

нейронов, что приводит к подавлению активности нейронов в V4 за пределами фокуса внимания.

Общая структура модели показана на рис. 8А. Сигналы от сетчатки (10×10 пикселей) подаются через латеральные геникулярные ядра в область V1 (не показаны на рис. 8А), откуда они поступают в область V2, где происходит их модуляция. Предполагается, что клетки в V2 чувствительны к одному из двух разных признаков, красному или зеленому цветам. Далее сигналы проецируются на нейронные колонки (рис. 8В) в области V4, состоящие из возбуждающих пирамидальных и тормозных интернейронов. Рецептивные поля нейронов V4 больше, чем у нейронов V2 (это отражает тот экспериментальный факт, что рецептивные поля нейронов увеличиваются по мере продвижения к более высоким уровням зрительной коры).

Конкуренция в V4 возникает за счет того, что тормозные нейроны ингибируют активности пирамидальных клеток с селективностью к противоположному свойству. Поступление на вход «зеленого» тормозного нейрона в V4 синхронизованных сигналов от «красных» нейронов слоя V2 повышает его активность, что в свою очередь влечет подавление активности «зеленых» возбуждающих нейронов в V4. Важно отметить, что эта схема торможения применяется только к возбуждающим нейронам в V4, имеющим сильно пересекающиеся рецептивные поля в V2.

Импульсы от нейронов из V2 генерируются с помощью пуассоновского процесса со средней скоростью λ , которая является суммой скоростей спонтанной активности λ_{SPONT} и стимул-зависимой активности λ_0 . Если представлен стимул, предпочитаемый клеткой (красный или зеленый), то λ_0 выбирается пропорционально степени пространственного перекрытия рецептивного поля клетки со стимулом

$$\lambda_0 = \lambda_{\text{max}} \cdot \text{overlap},$$

где $\lambda_{\text{max}} = 200$ Гц и overlap варьируется между 0 (стимула нет) и 1 (стимул полностью покрывает рецептивное поле). Для выполнения модуляции используется неоднородный пуассоновский процесс, скорость которого зависит от времени. События этого процесса не соответствуют потенциалам действия нейронов V2, но соответствуют последовательности синхронного нарастания и подавления скоростей спайков всех нейронов, включенных во внимание. Таким образом, средняя скорость генерации спайков нейроном

$$\lambda(t) = \lambda_0 P(t) + \lambda_{\text{SPONT}},$$

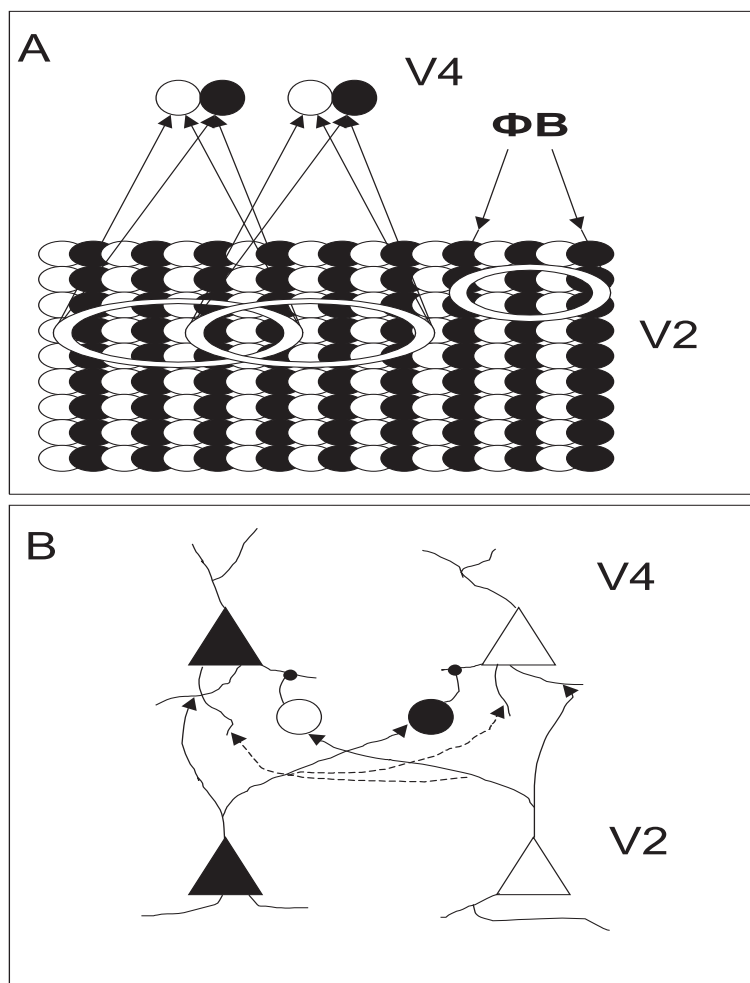


Рис. 8. Модель селективного внимания

А. Рецептивные поля клеток области V4, ФВ — фокус внимания. **В.** Синаптические связи между нейронами областей V2 и V4.

где $P(t) = 1$, исключая период модуляции. Если период модуляции начинается в момент t_0 , тогда

$$P(t) = \begin{cases} 8, & t_0 \leq t < t_0 + 2.5 \text{ мс}, \\ 0, & t_0 + 2.5 \leq t \leq t_0 + 17.5 \text{ мс}. \end{cases}$$

Такой вид функции $P(t)$ гарантирует, что средняя скорость спайков всегда постоянна (не зависит от наличия или отсутствия модуляции). Нейронные колонки области V4 состоят из двух типов нейронов. Тормозные нейроны работают как детекторы совпадений и определяют корреляции на входе. Возбуждающие нейроны являются выходными, они передают сигнал высшим кортикальным областям. Нейроны обоих типов получают входные сигналы от всех 100 нейронов V2. Основное отличие между типами нейронов — поведение порога срабатывания. Возбуждающие нейроны имеют не зависящий от времени порог $\theta = 15$ мВ, в то время как порог тормозных нейронов представляет собой гладкое усреднение мембранного потенциала клетки с временной константой τ_θ . Предполагается, что порог в среднем превышает мембранный потенциал на величину $\theta_0 = 7$ мВ. Поэтому для тормозных нейронов порог задается дифференциальным уравнением

$$\tau_\theta \frac{d\theta}{dt} = -\theta + \theta_0 + V + D(V_{Na} - V)s(t),$$

где D — константа и $V_{Na} = 100$ мВ — натриевый обратный потенциал, $s(t)$ — потенциал действия (равен 1 или 0). Уравнение для мембранного потенциала

$$\tau \frac{dV}{dt} = -V + I + g_K(V_K - V),$$

где τ — временная константа, I — нормированный синаптический ток. Последний член в уравнении представляет деполяризацию мембраны, g_K — проводимость калиевых каналов.

Синаптический ток является результатом синаптической проводимости и разности между величинами $V(t)$ и соответствующего обратного потенциала. Натриевые каналы активируются возбуждающими синапсами, а хлорные — ингибирующими синапсами. Поэтому общий синаптический ток

$$I = g_{exc}(V - V_{Na}) + g_{inh}(V - V_{Cl}).$$

Изменения синаптических проводимостей постсинаптических нейронов V4 (g_{exc} и g_{inh}) вызываются потенциалами действия нейронов V2.

Пусть t_{ij}^S — момент j -го спайка от i -го пресинаптического нейрона V2 с той же избирательностью к признаку, что и нейрон в V4, и пусть W_{exc}^S — вес связи между этими нейронами. Аналогично t_{ij}^P и W_{exc}^P — соответствующие момент спайка и вес связи для нейронов из V2 с другой избирательностью (предполагается $W_{exc}^S > W_{exc}^P$). Кроме этого вводятся две стохастические серии спайков от других областей, соответствующие возбуждающим (спайки в моменты t_i^E) и тормозным (спайки в моменты t_i^I) входам. Изменения в проводимости возбуждающих синапсов

$$\tau_{exc} \frac{dg_{exc}}{dt} = -W_{exc}^S \sum_{i,j} \delta(t - t_{ij}^S) + W_{exc}^P \sum_{i,j} \delta(t - t_{ij}^P) + W_{exc}^E \sum_j \delta(t - t_j^E).$$

Для тормозных синапсов выражение такое же, только отсутствуют входы от V4 и нейронов со сходной избирательностью к признаку

$$\tau_{inh} \frac{dg_{inh}}{dt} = -W_{inh}^P \sum_{i,j} \delta(t - t_{ij}^P) + W_{exc}^I \sum_j \delta(t - t_j^I).$$

Когда $V > \theta$ и последний спайк пришел до периода рефрактерности (период рефрактерности выбирается случайно из интервала (2 мс, 5 мс)), тогда генерируется потенциал действия $s(t) = 1$. Клетка реполяризуется, открывая кальций-зависимые калиевые каналы, проводимость которых задана как

$$\tau_K \frac{dg_K}{dt} = -g_K + B \cdot s(t).$$

Поведение сети изучалось путем регистрации уровня синхронизации спайковой активности нейронов в V2 и V4 и построения постстимульных гистограмм импульсной активности для возбуждающих (выходных) нейронов в V4. Результаты имитационного моделирования показали качественное совпадение динамики активности в полях V2 и V4 с экспериментальными данными. При одновременной подаче двух стимулов разного цвета в рецептивное поле, общее для двух нейронов в V4, результирующая активность этих нейронов определялась тем, какой из них был в фокусе внимания. Активность нейрона вне фокуса внимания в значительной мере подавлялась. В то же время, если конкурирующий стимул подавался в другое рецептивное поле, то уровень активности соответствующего ему нейрона в V4 оставался высоким.

Модели формирования фокуса внимания

Одной из плодотворных идей в теории внимания является идея центрального исполнителя (central executive). Она была предложена Бэддели [31, 32] и оформилась в виде экспериментально обоснованной функциональной схемы в работе Коуэна [33]. Согласно этой концепции, в мозге имеется распределенная нейронная структура, которая организует процесс обмена между долговременной и кратковременной памятью и фокусом внимания.

В неявном виде центральный исполнитель уже появлялся в работах [23, 24], где его роль фактически играл общий тормозный элемент. Работа [34] модифицирует сеть из [23, 24], превращая ее в модель внимания, способную выделить из изображения наиболее проявленный (по размеру или контрастности) объект. При выборе этого объекта используется механизм «победитель получает все» (winner takes all), реализованный на осцилляторных принципах. Главным отличием данной модели является использование в ней в качестве центрального исполнителя двух тормозных элементов, быстрого и медленного. Быстрый элемент организует попеременную активность разных ансамблей осцилляторов, соответствующих объектам на изображении, как это имело место в [23, 24]. Медленный элемент позволяет идентифицировать ансамбль осцилляторов, соответствующий наиболее проявленному объекту, и подавить активность других ансамблей.

Последовательность просмотра объектов такова: сначала в фокус внимания выбираются несколько объектов (этот набор может включать, а может и не включать искомый объект), потом постепенно ненужные (менее проявленные) объекты исключаются из фокуса внимания, а более проявленные включаются в фокус внимания. Эта процедура заканчивается, когда в фокусе внимания останется один наиболее проявленный объект.

Другой подход к моделированию внимания с помощью центрального исполнителя (на наш взгляд биологически более оправданный) предложен в работах Крюкова. Опираясь на результаты Виноградовой с сотрудниками [35, 36], Крюков предположил, что роль центрального исполнителя играет септо-гиппокампальная система, и предложил использовать фазовую синхронизацию в качестве механизма взаимодействия между центральным исполнителем и структурами новой коры [37, 38].

В модели Крюкова центральный исполнитель — это такой же осциллятор, как и другие компоненты сети. Модель описывается следующей системой уравнений для так называемых фазовых осцилляторов с непрерывной

динамикой.

$$\frac{d\theta_0}{dt} = \omega_0 + \frac{1}{n} \sum_{i=1}^n A_i f(\theta_i - \theta_0), \quad (2)$$

$$\frac{d\theta_i}{dt} = \omega_i + B_i g(\theta_0 - \theta_i), \quad i = 1, \dots, n, \quad (3)$$

где θ_k ($k = 0, 1, \dots, n$) — фазы осцилляторов, ω_k — собственные частоты осцилляторов, A и B — параметры, задающие силу взаимодействия между осцилляторами, f и g — функции взаимодействия (периодические), производная $d\theta_k/dt$ описывает текущую частоту осциллятора. Уравнение (2) задает динамику центрального осциллятора (ЦО), который играет в модели роль центрального исполнителя. Уравнение (3) описывает динамику кортикальных осцилляторов (колонок неокортекса — популяций взаимодействующих возбуждающих и тормозных нейронов), которые мы будем называть периферическими осцилляторами (ПО). Каждый ПО кодирует определенный признак и является активным, если этот признак присутствует на входном изображении.

С технической точки зрения достоинством архитектуры системы (2)–(3) является то, что в ней глобальное взаимодействие между осцилляторами осуществляется через центральный элемент. Это позволяет избежать связей типа «все на все», ограничившись числом связей, имеющим тот же порядок, что и число элементов в сети. Отметим, что в отличие от модели внимания с центральным элементом, рассматриваемой в следующем подразделе, связи между ЦО и ПО имеются как в прямом, так и в обратном направлении. Это значительно увеличивает возможности выбора нужного режима синхронизации между ЦО и ансамблями ПО.

При подходящих значениях параметров системы возникает синхронизация между ЦО и всеми или частью ПО. Под синхронизацией здесь понимается работа осцилляторов с одной и той же текущей частотой (при этом между синхронизованными осцилляторами может быть некоторая фиксированная разность фаз). Считается, что те ПО, которые работают синхронно с ЦО, формируют фокус внимания, а остальные осцилляторы представляют собой признаки отвлекающих стимулов (дистракторов). Какие именно ПО войдут в фокус внимания, зависит от параметров системы и может регулироваться путем подходящей модификации этих параметров, в частности, в процессе обучения.

Математическое исследование, проведенное в [18, 39] для случая синусоидальной функции взаимодействия между осцилляторами, позволило описать различные режимы синхронизации (полной или частичной) в системе (2)–(3). В терминах параметров системы были получены условия, при которых может формироваться тот или иной фокус внимания, и описаны возможные промежуточные стадии, возникающие при смене одного фокуса внимания другим.

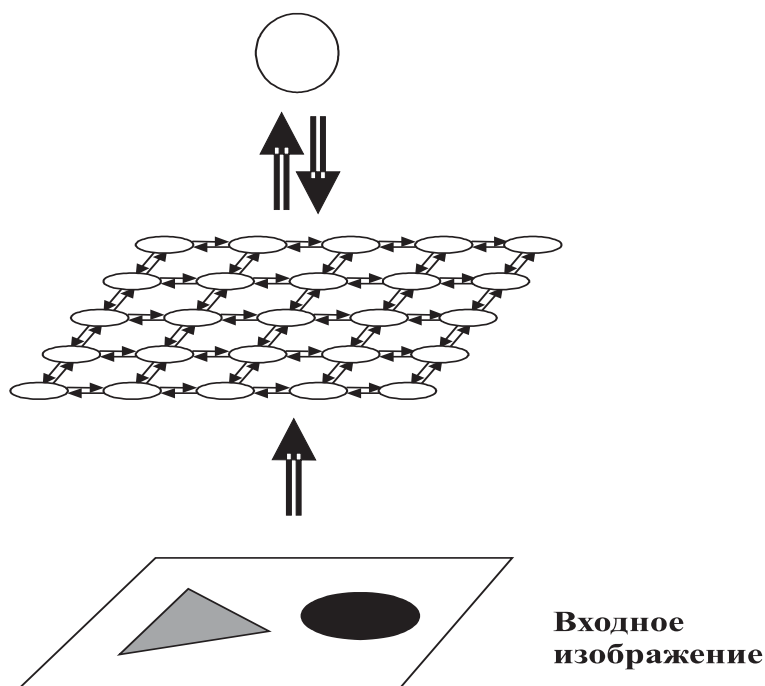


Рис. 9. Архитектура связей в модели формирования фокуса внимания

Хотя модель (2)–(3) позволяет при подходящем выборе параметров включить в фокус внимания любой набор ПО, вопрос об автоматическом формировании требуемых параметров в рамках этой модели остается открытым. При неудачном выборе параметров возможна ситуация,

когда в фокус внимания попадут ПО, кодирующие признаки разных объектов (появление «химеры»). Эта проблема становится особенно острой, если собственные частоты ПО, кодирующих разные объекты, оказываются близки друг другу. В силу сравнительной узости используемого диапазона собственных частот (например, в случае тета-ритма это 4–10 Гц), вероятность совпадения этих частот достаточно велика. В рассматриваемой ниже модели внимания [40] эти недостатки преодолены за счет использования (дополнительно к фазовой синхронизации) принципа резонанса — резкого возрастания амплитуды ПО, работающего синфазно с ЦО. Эта модель позволяет выбирать из имеющихся на изображении объектов один связный объект и включать его в фокус внимания. При этом предпочтение отдается объектам с большей площадью и большей контрастностью по отношению к фону.

Архитектура связей модели представлена на рис. 9. Как и в модели (2)–(3), сеть состоит из центрального осциллятора, имеющего прямые и обратные связи с периферическими осцилляторами. ПО расположены в узлах плоской решетки и имеют локальные связи с ближайшими соседями.

Входное изображение задано на плоской решетке того же размера, что и решетка ПО; таким образом, каждый ПО получает входной сигнал от своего пикселя на изображении. Этот сигнал задает собственную частоту ПО, которая тем выше, чем выше контрастность пикселя по отношению к фону. ПО, получающие сигнал от пикселей фона, не активны и не участвуют в функционировании сети.

Каждый осциллятор описывается тремя переменными: фазой колебаний, амплитудой колебаний и собственной частотой осциллятора. Динамика сети определяется уравнениями

$$\frac{d\theta_0}{dt} = \omega_0 + \frac{w}{n} \sum_{i=1}^n a_i v(\omega_i) g(\theta_i - \theta_0), \quad (4)$$

$$\frac{d\theta_i}{dt} = \omega_i - a_0 w_0 h(\theta_0 - \theta_i) + w_1 \sum_{j \in N_i} a_j p(\theta_j - \theta_i) + \rho, \quad i = 1, \dots, n, \quad (5)$$

$$\frac{da_i}{dt} = \beta_1 [-a_i + \gamma f(\theta_0 - \theta_i)]^+ + \beta_2 [-a_i + \gamma f(\theta_0 - \theta_i)]^-, \quad (6)$$

$$\frac{d\omega_0}{dt} = \alpha \frac{w}{n} \sum_{i=1}^n a_i v(\omega_i) g(\theta_i - \theta_0) = -\alpha \left[\omega_0 - \frac{d\theta_0}{dt} \right]. \quad (7)$$

В этих уравнениях θ_0 — фаза ЦО, θ_i — фазы ПО, ω_0 — собственная частота ЦО, ω_i — собственные частоты ПО, a_0 — амплитуда ЦО (можно считать, что $a_0 = 1$), a_i — амплитуды ПО, w, w_0, w_1 — параметры (положительные константы), задающие силу взаимодействия между осцилляторами, n — число активных ПО (только активные ПО включены в уравнения (4)–(6)), N_i — набор активных осцилляторов в ближайшей окрестности осциллятора i , ρ — гауссовский шум со средним 0 и дисперсией σ^2 , функции g, h, p задают взаимодействие между осцилляторами (эти функции 2π -периодические, нечетные, одномодальные в интервале периода), функция v положительная и монотонно возрастающая, она обеспечивает усиление влияния на ЦО тех ПО, которые получают больший внешний входной сигнал, функция f управляет динамикой амплитуды ПО и переходом в резонансное состояние (f — периодическая, четная, положительная, одномодальная в интервале периодичности, с максимумами в точках $2\pi k$), $\alpha, \beta_1, \beta_2, \gamma$ — параметры (положительные константы). Значения ω_i определяются внешним входным сигналом, а переменные $\theta_0, \theta_i, \omega_0, a_i$ характеризуют текущее состояние системы. По определению,

$$(x)^+ = \begin{cases} x, & \text{для } x \geq 0, \\ 0, & \text{для } x < 0, \end{cases} \quad (x)^- = \begin{cases} x, & \text{для } x \leq 0, \\ 0, & \text{для } x > 0. \end{cases}$$

Уравнения (4)–(5) являются уравнениями фазовой синхронизации. Связи от ПО к ЦО и локальные связи между ПО синхронизирующие. С помощью локальных связей между ПО формируются синфазные ансамбли осцилляторов, соответствующие связным объектам на изображении. С помощью связей от ПО к ЦО происходит синхронизация ЦО с одним из ансамблей синхронно работающих ПО. Связи от ЦО к ПО десинхронизирующие. С помощью этих связей десинхронизируется активность разных ансамблей ПО. Это препятствует вовлечению в синхронизацию с ЦО сразу нескольких ансамблей осцилляторов, что соответствовало бы включению в фокус внимания нескольких объектов. В соответствии с (4)–(5) различные ансамбли осцилляторов конкурируют за синхронизацию с ЦО. Чем больше ансамбль и чем больше контрастность объекта, соответствующего этому ансамблю, тем сильнее синхронизирующее влияние ансамбля на ЦО, и тем скорее данный ансамбль будет включен в фокус внимания.

Шум ρ в (5) используется как дополнительный источник десинхронизации между различными ансамблями ПО. Благодаря шуму происходит рандомизация положения в фазово-частотном пространстве различных ансамблей осцилляторов, что делает их «различимыми» для ЦО.

В качестве функций взаимодействия в (4)–(5) используются

$$h(\phi) = p(\phi) = \sin \phi,$$

$$g(\phi) = \begin{cases} 10\phi, & \text{для } 0 \leq \phi < 0.1, \\ -4\phi + 1, & \text{для } 0.1 \leq \phi < 0.2, \\ -0,1\phi + 0,62, & \text{для } 0.2 \leq \phi \leq \pi, \\ -g(-\phi), & \text{для } -\pi < \phi < 0. \end{cases}$$

Вне интервала $(-\pi, \pi)$ функция $g(\phi)$ продолжается как периодическая.

Уравнение (6) описывает динамику амплитуд ПО и обеспечивает резонансное возрастание амплитуды при долговременном совпадении фаз ЦО и ПО. Функция $f(x)$ имеет вид

$$f(x) = F((\cos x)^+),$$

где $F(y)$ — сигмообразная функция вида

$$F(y) = \zeta + \frac{\exp((y - \xi)/\eta)}{1 + \exp((y - \xi)/\eta)}.$$

Параметры ξ и η выбраны так, чтобы $F(y)$ приближалась к максимальному значению $1 + \zeta$ при $y \rightarrow 1$; при убывании y функция $F(y)$ быстро убывает до уровня ζ . В результате оказывается, что амплитуда ПО возрастает до максимального значения $a_{\max} = \gamma(1 + \zeta)$ при синфазной работе ПО и ЦО; амплитуда принимает минимальное значение $a_{\min} = \gamma\zeta$, если фазы ПО и ЦО существенно различаются (ζ на порядок меньше 1).

Считается, что ПО перешел в резонансное состояние, если его амплитуда превысила порог $R = 0,8a_{\max}$. Параметры β_1 и β_2 задают скорость возрастания и убывания амплитуды. Те ПО, которые находятся в резонансном состоянии, считаются включенными в фокус внимания. Это согласуется с экспериментальными данными, согласно которым активность нейронов в фокусе внимания повышена по сравнению с активностью нейронов за пределами фокуса внимания [28, 29, 41, 42].

Уравнение (7) описывает механизм адаптации собственной частоты ЦО. В соответствии с этим уравнением ω_0 стремится к текущему значению частоты ЦО. Такая адаптация позволяет ЦО «настраиваться» на тот ансамбль ПО, с которым ЦО собирается синхронизоваться. Параметр α задает скорость адаптации.

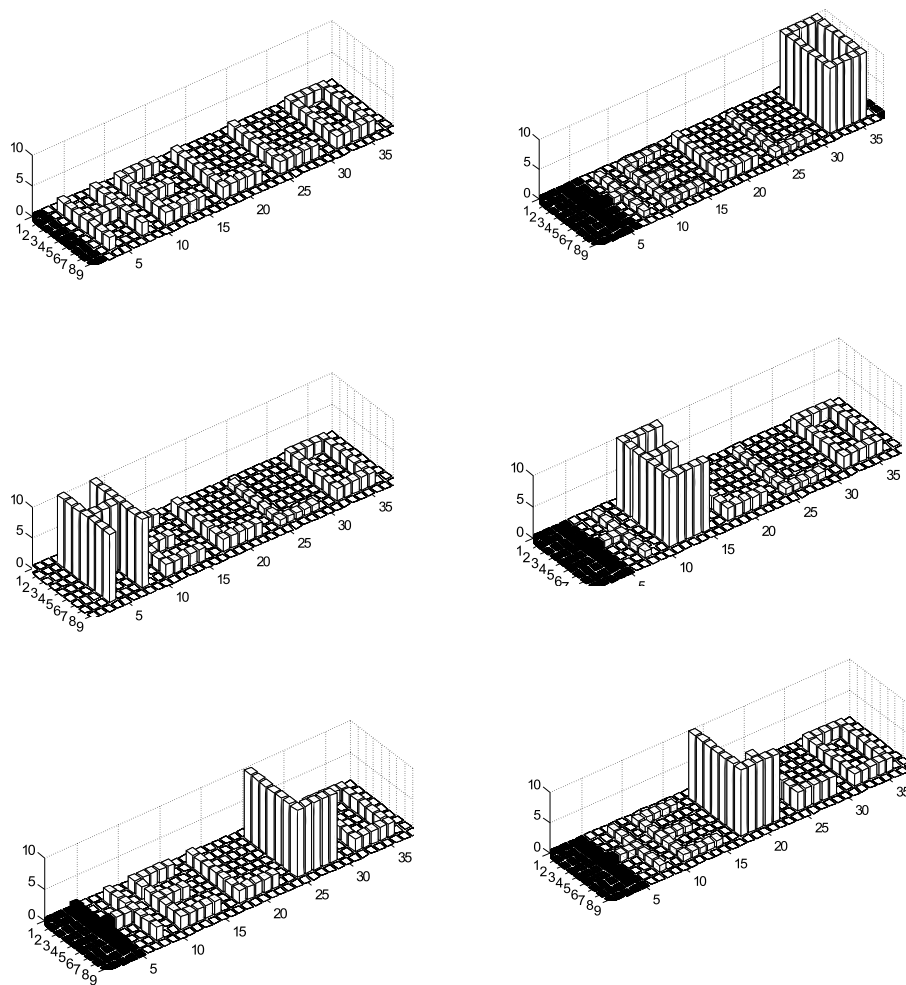


Рис. 10. Амплитуды периферических осцилляторов в различные моменты времени при последовательном выборе объектов

Помимо выбора наиболее заметного объекта в фокус внимания, модель (4)–(7) может работать и в режиме последовательного выбора в фокус внимания всех имеющихся на изображении объектов. Для этого ее достаточно пополнить условием, что объект после определенного времени нахождения в фокусе внимания блокируется, так что соответствующие ему ПО не могут взаимодействовать с ЦО. В этом случае фокус внимания освобождается от уже просмотренных объектов, и в него включаются другие объекты, расположенные на изображении, в порядке, соответствующем их яркости и размеру. На рис. 10 показан пример работы модели с изображением, представляющим собой напечатанное слово “HELLO”, буквы которого рассматриваются как отдельные объекты. Система последовательно выбирает буквы и переводит осцилляторы, соответствующие этим буквам, в резонансное состояние. Поскольку яркость всех черных пикселей предполагается одинаковой, порядок выбора букв определяется их размером.

Модель селективного внимания и задачи зрительного поиска

Одной из классических задач, связываемых с использованием системы внимания, является задача зрительного поиска. При зрительном поиске испытуемый осматривает дисплей, по которому случайным образом распределены объекты, с целью найти заранее заданный объект. Заданный объект является целью поиска, все другие объекты в данном случае играют роль отвлекающих стимулов (дистракторов). Обычно решение (правильное или ошибочное) должно быть принято за короткое время и при неподвижных зрачках, фиксированных на центр изображения.

Эксперименты по зрительному поиску можно разделить на два типа: поиск по одному признаку и поиск по нескольким признакам. В первом случае объект поиска отличается от дистракторов только одним признаком. Например, это может быть поиск зеленой буквы Т среди красных Т. Во втором случае цель отличается от дистракторов более чем одним признаком. Например, поиск зеленой буквы Т среди красных букв Т и зеленых букв Х. Важной измеряемой величиной является время поиска как функция от числа объектов. Как показывают эксперименты [8–10], в случае поиска по одному признаку цель определяется быстро за время, не зависящее от числа дистракторов на экране. В случае поиска по нескольким признакам время поиска линейно возрастает с ростом числа дистракторов. В связи с этими результатами принято считать (несколько упрощенно, в действительности ситуация более сложна [43]), что поиск по одному признаку выполняется

без использования внимания, а поиск по нескольким признакам требует привлечения внимания, с чем и связаны дополнительные затраты времени.

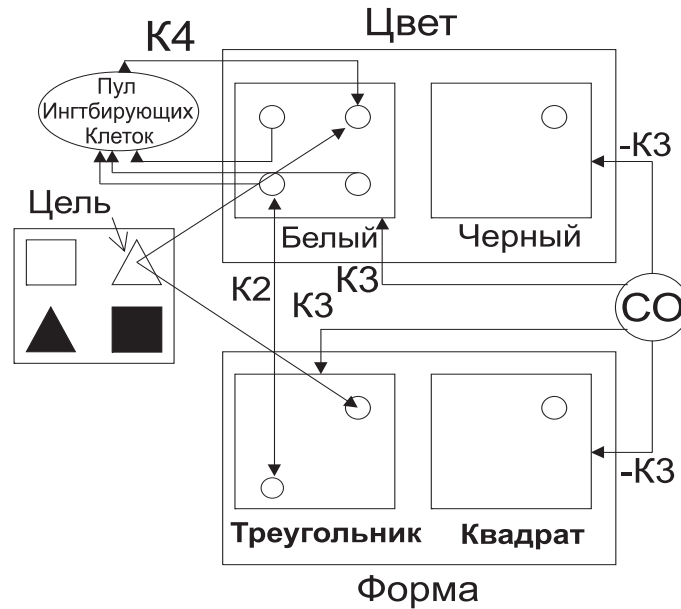


Рис. 11. Схема функционирования модели, решающей задачу зрительного поиска

На рис. 11 представлена иерархическая структура модели [44], решающая задачу поиска по одному и нескольким признакам. Основными структурными элементами модели являются колонки нейронов. Каждая колонка представляет собой сеть, построенную из импульсных фазовых осцилляторов. Динамика каждого осциллятора задается его фазой φ , которая подчиняется уравнению

$$\tau \frac{d\varphi}{dt} = -\varphi + A(t). \quad (8)$$

После того как φ достигает значения 2π (это событие соответствует моменту генерации спайка), фаза сбрасывается до нуля и вновь начинает расти

согласно (8). Переменная $A(t)$ представляет внешний стимул. В дальнейшем будет предполагаться, что $A(t) = A$. Для того чтобы сделать понятным принцип взаимодействия между осцилляторами, приведем в качестве примера описание динамики в сети из N осцилляторов, связанных по принципу «все на все»:

$$\tau \frac{d\varphi_i}{dt} = -\varphi_i + A + \eta_i(t) + \frac{1}{n} K J_i(t), \quad i = 1, \dots, N, \quad (9)$$

где K — константа связи и η — шум. Последний член в уравнении (9) представляет плотность спайков, получаемых осциллятором,

$$J_i(t) = \sum_{j \neq i} \sum_k^N \delta(t - t_j^k), \quad (10)$$

где t_j^k представляет k -й момент генерации спайка нейроном j .

Вход модели — это матрица из $S \times S$ зрительных образов (на рис. 11 $S = 2$). Положение каждого объекта определяется парой индексов (i, j) . Каждый объект имеет M признаков, а каждый признак m может принимать $L(m)$ значений (например, $L(\text{цвет}) = 2$, а именно $l = 1$ для белого и $l = 2$ — для черного). Таким образом, для каждого признака m существует $L(m)$ слоев нейронных колонок, характеризующих представленные свойства признаков. Каждая колонка идентифицируется индексами $ijml$. Для простоты будем считать, что $L(m) = L$.

Нейроны в колонке взаимодействуют с константой связи K_1 . Связи между колонками, имеющими одинаковое топографическое положение, но кодирующими свойства разных признаков, взаимодействуют с константой K_2 . Предполагается, что селективное внимание приводит к конкуренции между колонками, отвечающими за разные объекты. Поэтому в каждом нейронном слое нейроны каждой колонки взаимодействуют через пул тормозных нейронов с отрицательной силой связи $-K_4$.

Для осуществления задачи поиска используется заранее известная информация о цели. Функцию запоминающего механизма выполняет изолированная колонка — центральный осциллятор (ЦО). Цель поиска кодируется с помощью связей между ЦО и колонками, кодирующими признаки: если колонки принадлежат уровням, кодирующим признаки, соответствующие объекту поиска, то их связи с ЦО равны K_3 , для колонок, кодирующих другие признаки, связи равны $-K_3$. Если, например, нужно найти белый

треугольник, то положительные связи будут между ЦО и уровнями с индексами $l = \text{«белый»}$ для $m = \text{«цвет»}$ и $l = \text{«треугольник»}$ для $m = \text{«форма»}$.

Эволюция всей системы задана уравнениями

$$\begin{aligned} \tau \frac{d\varphi_{nijml}}{dt} &= -\varphi_{nijml} + A_{nijml} + \eta_{ij}(t) + \\ &+ \frac{1}{N} \left[K_1 \cdot J_{ijml}^1(t) + K_2 \cdot J_{ijm}^2(t) + \right. \\ &\left. + K_3^{ml} \cdot J^{CO}(t) - K_4 \cdot J_{ijm}^4(t) \right], \\ n &= 1, \dots, N; \quad m = 1, \dots, M; \quad l = 1, \dots, L, \\ \tau \frac{d\varphi_n^{CO}}{dt} &= -\varphi_n^{CO} + A^{CO} + \eta(t) + \frac{1}{N} K_1 \cdot J^{CO}(t), \end{aligned}$$

где A_{nijml} — внешний сигнал (в вычислениях $A_{nijml} = 2.2\pi$), показывающий присутствие свойства l признака m в позиции ij ; J^1 и J^{CO} — плотности спайков, соответствующие признаковым колонкам и ЦО, и определяемые уравнением (10). Другие плотности спайков вычисляются как

$$J_{ijm}^2(t) = \sum_{m' \neq m}^M \sum_{l'}^L J_{ijm'l'}^1(t), \quad J_{ijm}^4(t) = \sum_{i' \neq i}^S \sum_{j' \neq j}^S \sum_{l'}^L J_{i'j'ml'}^1(t).$$

Коллективные осцилляции нейронов в колонке появляются при достаточно большом K_1 . Благодаря связям между колонками с силой K_2 , осуществляется поддержание высокого уровня активности в колонках, кодирующих один объект. Нейронные уровни, отвечающие свойствам цели поиска, получают возбуждающие сигналы от ЦО через связи с силой K_3 . При этом некоторые дистракторы также получают возбуждающие сигналы, но не на всех уровнях. Только один объект из представленных на экране получает такую «помощь» на всех уровнях и выигрывает в конкуренции. Другой механизм конкуренции, включенный через пул тормозных нейронов, имеет место на каждом уровне и приводит к подавлению активностей дистракторов. Те колонки, которые осциллируют синхронно с центральным элементом, определяют фокус внимания. То есть, как и в предыдущем подразделе, внимание реализуется в виде синхронизации колонок, характеризующих свойства объекта, с центральным элементом. Но в отличие от этой модели, данная модель внимания не перебирает последовательно имеющиеся на изображении объекты, а путем параллельной обработки

находит нужный объект (заданный определенным набором признаков) и включает его в фокус внимания.

Модель тестировалась на различных изображениях, содержащих 4, 9 и 16 объектов при числе признаков $M = 2$ и числе градаций признаков $L = 2$. Результаты вычислений показали, что модель правильно отражает экспериментальные данные. Время поиска объекта по одному признаку не меняется при изменении числа дистракторов, а время поиска объекта по двум признакам линейно возрастает в зависимости от числа дистракторов. Увеличение времени в последнем случае обусловлено тем, что некоторые дистракторы получают синхронизирующие связи от центрального элемента, приходящие на один из двух имеющихся уровней. Это в свою очередь приводит к тому, что остальные дистракторы увеличивают свою активность, благодаря связям с коэффициентом K_2 . Из-за этого целевому объекту становится труднее выиграть конкуренцию за синхронизацию с ЦО, причем тем труднее, чем больше число дистракторов. Конечно, надо иметь в виду, что входная информация в модели задается довольно условно, поэтому требуются дополнительные усилия, чтобы довести эту модель до уровня обработки реальных изображений. Но идея использования признаков как источников синхронизации при поиске нужного объекта представляется нам плодотворной.

Заключение. Перспективы использования осцилляторных нейронных сетей при моделировании когнитивных функций мозга

Теория нейронных сетей выросла из попыток описать в формально-математических и вычислительных терминах способность мозга к решению сложных интеллектуальных задач. Нейронные структуры и алгоритмы, разработанные к концу 80-х годов, оказались весьма успешными в решении различных прикладных задач (формирования ассоциативной памяти, распознавания, прогноза, синтеза временных рядов и т. д.), но их функционирование лишь очень отдаленно напоминает реальные процессы, происходящие в мозге. Как правило, они не отражают структурную организацию мозга, не учитывают многие существенные детали функционирования физиологических нейронов и зачастую оперируют лишь усредненными характеристиками динамики активности нейронных ансамблей и структур. В этих рамках невозможно ни моделирование метастабильных состояний

и фазовых переходов, ни изучение пространственно-временных соотношений нейронной активности, что значительно сужает возможности кодирования и обработки информации.

Преодоление этих недостатков традиционных нейронных сетей идет по пути использования биологически более адекватных моделей элементов (импульсных нейронов, нейронов типа Ходжкина–Хаксли, многосегментных моделей нейронов), а также попыток воспроизведения динамических режимов (в частности, различных режимов колебательной активности), возникающих в реальных нейронных сетях. Наиболее биологически обоснованные нейронные сети были разработаны для моделирования движений, где сложная динамика возникает естественным путем и ее биологическая роль достаточно очевидна [45]. По-другому обстоит дело с моделированием когнитивных функций мозга. Несмотря на значительные усилия, приложенные в этом направлении за последние 15 лет, динамический подход не смог реализоваться в моделях, которые бы превзошли по своим способностям традиционные коннекционистские сети. Главная причина этого состоит в том, что роль сложной динамики нейронной активности в когнитивных процессах остается недостаточно понятой. Выдвинутые в связи с этим гипотезы еще не дошли до стадии, когда их экспериментальное подтверждение не вызывало бы сомнений. Кроме того, эти гипотезы, несмотря на свою исключительную важность (например, проблема интеграции признаков считается одной из центральных в современной нейрофизиологии), дают лишь фрагментарную картину обработки информации в мозге. Соединение этих фрагментов в цельную теорию — важнейшая задача, в решении которой нейросетевое моделирование может сыграть существенную роль.

До настоящего времени большинство динамических нейросетевых моделей было направлено на воспроизведение отдельных аспектов работы мозга. С нашей точки зрения, успех может быть достигнут тогда, когда будет разработана полномасштабная модель, в которой такие когнитивные функции как сегментация информации, внимание, детекция новизны, распознавание, кратковременная и долговременная память будут реализованы в виде системы взаимодействующих модулей. Имеющиеся в настоящее время разработки позволяют уже сейчас начать реализацию такой системы на осцилляторных принципах.

Упрощенная блок-схема системы приведена на рис. 12. Предполагается, что исходная зрительная информация должна пройти фильтрацию с помощью системы внимания, что позволило бы выделить ту часть информации, которая относится к одному объекту. Эта задача может быть решена

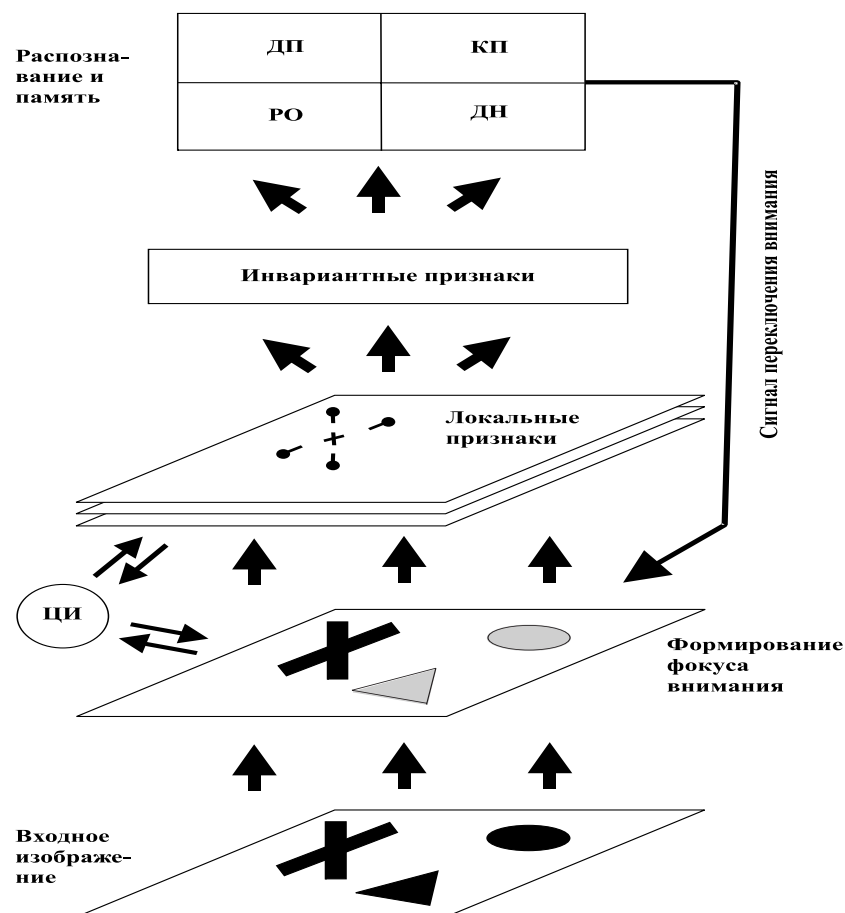


Рис. 12. Блок-схема модели когнитивных функций мозга
 ЦИ — центральный исполнитель системы внимания, ДП — долговременная память, КП — кратковременная память, РО — распознавание образов, ДН — детекция новизны. В модуле формирования фокуса внимания объект в фокусе внимания обозначен черным цветом. В модуле локальных признаков показаны геометрические признаки свободных концов (разной ориентации) и перекрещивающихся линий.

в результате взаимодействия центрального исполнителя с модулями, где формируются локальные признаки. К таким признакам относятся характеристики отдельных пикселей (яркость, контрастность, цвет), характеристики линий и их пересечений (с учетом ориентации), локальные комбинации геометрических и спектральных признаков. Эти модули соответствуют первичной и вторичной зонам зрительной коры. Синхронно работающий ансамбль осцилляторов в этих модулях может играть роль источника синхронизации для всего восходящего потока информации, относящейся к объекту в фокусе внимания. На наш взгляд, основные принципы функционирования этого модуля достаточно хорошо разработаны в существующих моделях, но практические примеры их применения к реальным изображениям весьма немногочисленны [24] и ограничиваются черно-белыми статичными изображениями. Очевидно, существующие подходы должны быть усовершенствованы и применены к анализу реальных цветных динамически меняющихся трехмерных сцен.

В следующем модуле происходит преобразование признаков в форму инвариантную к расположению объекта и масштабу. Это соответствует работе височных зон коры. Проблема инвариантности представления информации — одна из наиболее обсуждаемых в теории нейронных сетей. С нашей точки зрения, перспективным является подход, предложенный в работе [46] и основанный на иерархической «сборке» сложного объекта из более простых компонент. К сожалению, осцилляторная реализация этого подхода в настоящее время отсутствует.

Последний модуль высших когнитивных функций представляет собой наибольшую трудность. В мозге его роль играют несколько структур, такие как лимбическая кора, лобные доли коры, ассоциативные зоны. Существенным компонентом этого модуля является детекция новизны. Эта операция более простая и быстрая, чем распознавание, в то же время она позволяет избежать бесполезных затрат времени на анализ известных или несущественных для данной ситуации объектов. Осцилляторная реализация детекции новизны в гиппокампе предложена в [47, 48]. Она построена на тех же принципах, что и описанная выше модель внимания [40] и может быть легко с ней объединена. Однако это достаточно грубая модель, в которой не использованы ни управляющая активностью в гиппокампе обратная связь от гиппокампа к ретикулярной формации ствола мозга [49], ни долговременная память. Необходимо построить более совершенную модель, в которой эти пробелы были бы ликвидированы.

Наиболее серьезные проблемы в данной схеме связаны с распознаванием образов. Ничего сравнимого по эффективности с сетями обратного распространения ошибок и сетями Кохонена в осцилляторных сетях пока нет. Определенные надежды здесь могут быть связаны с разработкой осцилляторного аналога сетей типа ART (Artificial Resonance Theory) [50], поскольку принципы функционирования этих сетей наиболее близки биологическим, а используемый в них «искусственный» резонанс естественным образом может быть преобразован в реальный резонанс колебательной активности.

Существующие осцилляторные модели памяти пока что идут в основном по пути перевода на осцилляторный язык идей ассоциативной памяти Хопфилда. Один из примеров этого подхода приводился выше [20, 21]. Другой пример с использованием фазовых осцилляторов можно найти в работе [51]. К сожалению, сеть Хопфилда обладает известными недостатками (низкая емкость и большое число связей), которые переносятся и на ее осцилляторные аналоги. Поэтому перед исследователями стоит важная задача найти более эффективные и биологически более оправданные варианты ассоциативной памяти на осцилляторах.

Важную роль в реализации когнитивных функций играет обратный поток информации от верхних модулей вниз. На схеме показан простейший пример такого сигнала: после окончания процесса распознавания и запоминания объекта или в случае его детекции как известного на систему внимания должен быть послан сигнал, который позволит переместить фокус внимания на другой объект. На самом деле таких управляющих сигналов, идущих сверху вниз, достаточно много. Пример использования таких сигналов для выбора объекта с заданными признаками рассматривался выше [44].

Подводя итоги, можно утверждать, что использование осцилляторных нейронных сетей для моделирования когнитивных функций мозга имеет хорошие перспективы и может принести не только лучшее понимание принципов работы мозга, но приведет к построению более эффективных систем искусственного интеллекта. Однако, для того чтобы эти перспективы стали реальностью, специалистам по моделированию и нейробиологам нужно проделать еще достаточно сложную и интересную работу.

Участие в данной публикации *Я. Б. Казановича* было поддержано Российским фондом фундаментальных исследований (грант 03-04-48482 и грант для научных школ НШ 1872.2003.4).

Приложение. Типы нейронных сетей

Обычно осцилляторная динамика в нейронных ансамблях возникает в результате взаимодействия популяций возбуждающих и тормозных нейронов или за счет использования пейсмекерных нейронов, задающих эндогенный ритм. Различают четыре типа нейронных сетей, соответствующих типам используемых элементов [45, 52, 53].

1. *Динамика нейрона описывается системой дифференциальных уравнений.* Это могут быть уравнения ионного транспорта через мембрану, как в модели Ходжкина–Хаксли, или сегментные (multicompartment) модели нейрона, описывающие динамику токов в сегментах нейрона (дендритах, соме и аксоне). Сети из таких нейронов позволяют наиболее полно учитывать динамику реальных нейронов, но анализ динамики самой сети обычно сопряжен с трудностями.

2. *Импульсные (integrate-and-fire) нейроны.* Модель импульсного нейрона является сравнительно простым устройством, накапливающим поступающие сигналы и генерирующим импульс (спайк) при достижении порога. В этой модели могут учитываться такие аспекты функционирования реального нейрона, как абсолютная и относительная рефрактерность, синаптическая задержка при передаче сигнала, динамика постсинаптических потенциалов, шумовая компонента (имитирующая, например, далекие дендриты или синаптический шум) и т.д. До настоящего времени сети, построенные из таких нейронов, являются наиболее распространенным объектом в имитационном моделировании.

3. *Нейронные осцилляторы.* Динамика нейрона, входящего в осциллятор, описывается средним уровнем активности популяции, к которой он принадлежит. Типичным примером нейронного осциллятора является осциллятор Вилсона–Коуэна. Сети из таких осцилляторов исследуются методами теории бифуркаций, что позволяет численно и аналитически определять области параметров, при которых имеют место те или иные виды динамики сети.

4. *Фазовые осцилляторы.* В случае, когда активность нейронных ансамблей имеет осцилляторный характер, удобно описывать эту активность посредством фазового осциллятора, динамика которого характеризуется только одной переменной — фазой. Сети из фазовых осцилляторов полезны для аналитических и численных исследований условий синхронизации в сети.

С точки зрения архитектуры связей нейронные сети можно разделить на следующие типы:

1. *Сети с локальными связями, связями "все на все", а также случайными и разреженными связями.* Однородность архитектуры связей и идентичность элементов позволяют аналитически исследовать такие сети с помощью методов статистической механики.

2. *Многослойные сети с прямыми, обратными и рекуррентными связями между слоями.* Такая архитектура наиболее популярна в приложениях нейронных сетей к распознаванию образов, прогнозу, кодированию и фильтрации сигналов, запоминанию временных последовательностей.

3. *Сети с центральным элементом.* Эти сети позволяют избежать большого числа связей, характерного для полносвязных сетей, так как элементы сети взаимодействуют через центральный элемент, имеющий связи с остальными элементами.

4. *Сети со сложной архитектурой.* К этой категории относятся сети с иерархической структурой, состоящей из более простых сетей разных типов. Потребности в таких сетях возникают при решении задач многоэтапной обработки информации и при моделировании сложного поведения.

Литература

1. *Basar E.* Brain function and oscillations. – Springer-Verlag, New York, 1998.
2. *Singer W.* Neuronal synchrony: A versatile code for the definition of relations // *Neuron*, **24**, 49–65, 1999.
3. *Singer W., Gray C.M.* Visual feature integration and the temporal correlation hypothesis // *Ann. Rev. Neurosci.*, **18**, 555–586, 1995.
4. *Milner P.M.* A model for visual shape recognition // *Phys. Rev.*, **81**, 521–535, 1974.
5. *Malsburg C. von der.* The correlation theory of brain function. Internal report 81-2, Max-Planck Institute for Biophysical Chemistry, 1981 (reprinted in *Models of Neural Networks*, Eds. *E. Domany, J. L. van Hemmen, K. Schulten*. – Springer, New York, 1994, p. 95).
6. *Gray C.M., Konig P., Engel A.K., Singer W.* Oscillatory responses in cat visual cortex exhibit intercolumnar synchronization which reflects global stimulus properties // *Nature*, **388**, 334–337, 1989.
7. *Eckhorn R., Bauer R., Jordon W., Brosch M., Kruse W., Munk M., Reitboeck H.J.* Coherent oscillations: a mechanism of feature linking in the visual cortex // *Biol. Cybern.*, **60**, 121–130, 1988.

8. Treisman A., Gelade G. A feature-integration theory of attention // *Cognitive Psychol.*, **12**, 97–136, 1980.
9. Treisman A., Sato S. Conjunction search revisited // *J. Exp. Psychol.*, **16**, 459–478, 1990.
10. Treisman A. Feature binding, attention and object perception // *Philosophical Transactions. Biological Sciences*, **353**, 1295–1306, 1998.
11. Fries P., Schroder J.-H., Roefsma P.R., Singer W., Engel A.K. Oscillatory neural synchronization in primary visual cortex as a correlate of stimulus selection // *J. Neurosci.*, **22**, 3739–3754, 2002.
12. Niebur E. Electrophysiological correlates of synchronous neural activity and attention: A short review // *BioSystems*, **67**, 157–166, 2002.
13. Niebur E., Hsiao S.S., Johnson K.O. Synchrony: a neuronal mechanism for attentional selection? // *Curr. Opinion Neurobiol.*, **12**, 190–194, 2002.
14. Sporns O., Tononi G., Edelman G. Modeling perceptual grouping and figure-ground segregation by means of active reentrant connections // *Proc. Nat. Acad. Sci. (USA)*, **88**, 129–133, 1991.
15. Tononi G., Sporns O., Edelman G. Reentry and the problem of integrating multiple cortical areas. Simulation of dynamic integration in the visual system // *Cerebral Cortex*, **2**, 310–335, 1992.
16. Schillen T.B., Konig P. Binding by temporal structure in multiple feature domains of an oscillatory neural network // *Biol. Cybern.*, **70**, 397–405, 1994.
17. Neibur E., Schuster H., Kammen D. Collective frequencies and metastability in networks of limit cycles oscillators with time delay // *Phys. Rev. Lett.*, **67**, 2753–2756, 1991.
18. Borisyuk R.M., Kazanovich Y.B. Oscillatory neural network model of attention focus formation and control // *BioSystems*, **71**, 29–38, 2003.
19. Verschure P., Konig P. On the role of biophysical properties of cortical neurons in binding and segmentation of visual scenes // *Neural Comput.*, **11**, 1113–1138, 1999.
20. Ritz R., Gerstner W., van Hemmen J. Associative binding and segregation in neural network of spiking neurons // *In Models of Neural Networks*, Eds. E. Domany, J.L. van Hemmen, K. Schulten. – Springer, New York, 1994, pp.175–219.
21. Ritz R., Gerstner W., Fuentes U., van Hemmen J. A biologically motivated and analytically soluble model of collective oscillations in the cortex. II. Application to binding and pattern segmentation // *Biol. Cybern.*, **71**, 349–358, 1994.
22. Malsburg C. von der, Buhmann J. Sensory segmentation with coupled neural oscillators // *Biol. Cybern.*, **67**, 233–242, 1992.

23. Wang D. L., Terman D. Locally excitatory globally inhibitory oscillator network // *IEEE Trans. Neural Networks*, **6**, 283–286, 1995.
24. Wang D. L., Terman D. Image segmentation based on oscillatory correlation // *Neural Comput.*, **9**, 805–836, 1997.
25. Horn D., Sagi D., Usher M. Segmentation, binding, and illusory conjunctions // *Neural Comput.*, **3**, 510–525, 1991.
26. Moore C., Wolfe J. Getting beyond the serial/parallel debate in visual search: A hybrid approach // In *The Limits of Attention*, Ed. K. Shapiro. Oxford Univ. Press, Oxford, 2001, pp. 178–198.
27. Motter B. C. Focal attention produces spatially selective processing in visual cortical areas V1, V2 and V4 in the presence of competing stimuli // *J. Neurophysiology*, **70**, 909–919, 1993.
28. Moran J., Desimone R. Selective attention gates visual processing in the extrastriate cortex // *Science*, **229**, 782–784, 1985.
29. Desimone R. Neural circuits for visual attention in the primate brain // In *Neural Networks for Vision and Image Processing*, Eds. G. A. Carpenter, S. Grossberg. Bradford Book, MIT Press, Cambridge, MA, 1992, pp. 343–364.
30. Niebur E., Koch C. A model for the neural implementation of selective visual attention based on temporal correlation among neurons // *J. Comput. Neurosci.*, **1**, 141–158, 1994.
31. Baddeley A. Working memory. – Oxford Univ. Press, London, 1986.
32. Baddeley A. Exploring the central executive // *Quarterly J. Experimental Psychol.*, **40A**, 5–28, 1996.
33. Cowan N. Evolving conceptions of memory storage, selective attention and their mutual constraints within the human information-processing system // *Psychol. Bull.*, **104**, 163–191, 1988.
34. Wang D. L. Object selection based on oscillatory correlation // *Neural Netw.*, **12**, 579–592, 1999.
35. Виноградова О. С. Гиппокамп и память. – М: Наука, 1975.
36. Vinogradova O. S., Brazhnik E. S., Stafekhina V. S., Belousov A. B. Septo-hippocampal system: rhythmic oscillations, and information selection // In *Neurocomputers and Attention I: Neurobiology, Synchronization and Chaos*, Eds. A. V. Holden, V. I. Kryukov. Manchester University Press, Manchester, 1991, pp. 129–148.
37. Kryukov V. I. An attention model based on the principle of dominantia // In *Neurocomputers and Attention I: Neurobiology, Synchronization and Chaos*, Eds. A. V. Holden, V. I. Kryukov. Manchester University Press, Manchester, 1991, pp. 319–352.

38. Крюков В. И. Модель внимания и памяти, основанная на принципе доминанты и компараторной функции гиппокампа // *Журнал высшей нервной деятельности*, №1, 2004 (в печати).
39. Kazanovich Y. B., Borisyuk R. M. Dynamics of neural networks with a central element // *Neural Netw.*, **12**, pp. 441–454, 1999.
40. Kazanovich Y., Borisyuk R. Object selection by an oscillatory neural network // *BioSystems*, **67**, 103–111, 2002.
41. Treue S., Maunsell J. H. R. Attentional modulation of visual motion processing in cortical areas MT and MST // *Nature*, **382**, 539–541, 1996.
42. Kastner S., De Weerd P., Desimone R., Ungerleider L. G. Mechanisms of directed attention in the human extrastriate cortex as revealed by functional MRI // *Science*, **282**, 108–111, 1998.
43. Chun M. M., Wolfe J. M. Visual attention // In: *Blackwell Handbook of Perception*, Ch.9, Ed. E. B. Goldstein. Blackwell, Oxford, UK, 2001, pp. 272–310.
44. Corch S., Deco G. A neorodynamical model for selective visual attention using oscillators // *Neural Netw.*, **14**, 981–990, 2001.
45. Борисюк Г. Н., Борисюк Р. М., Казанович Я. Б., Иваницкий Г. Р. Модели динамики нейронной активности при обработке информации мозгом — итоги «десятилетия» // *УФН*, **172**(10), 1189–1214, 2002.
46. Hummel J. E. Complementary solutions to the binding problem in vision: Implications for shape perception and object recognition // *Visual cognition*, **8**, 489–517, 2001.
47. Борисюк Р. М., Виноградова О. С., Денэм М., Казанович Я. Б., Хоппенштедт Ф. Модель детекции новизны на основе осцилляторной нейронной сети с разреженной памятью // *III Всероссийская конференция «Нейроинформатика-2001»*, Москва, МИФИ, 2001, т. 1, с. 183–190.
48. Borisyuk R., Denham M., Kazanovich Y., Hoppensteadt F. Vinogradova O. Oscillatory model of novelty detection // *Network*, **12**, pp. 1–20, 2001.
49. Vinogradova O. S. Hippocampus as comparator: the role of two input and two output systems of the hippocampus in selection and registration of information // *Hippocampus*, **11**, pp. 578–598, 2001.
50. Grossberg S. How does the cerebral cortex work? Learning, attention and grouping by laminar circuits of visual cortex // *Spatial Vision*, **12**, 163–186, 1999.
51. Kazanovich Y. B., Kryukov V. I., Luzyanina T. B. Synchronization via phase-locking in oscillatory models of neural networks // In *Neurocomputers and Attention I: Neurobiology, Synchronization and Chaos*, Eds. A. V. Holden, V. I. Kryukov. Manchester University Press, Manchester, 1991, pp. 269–284.

52. Абарбанель Г.Д.И., Рабинович М.И., Селверстон А., Рубчинский Л.Л. и др. Синхронизация в нейронных ансамблях // *УФН*, **166**(4), 363–390, 1996.
53. Sturm A. K., Konig P. Mechanisms to synchronize neural activity // *Biol. Cybern.*, **84**, 153–182, 2001.

Яков Борисович КАЗАНОВИЧ, старший научный сотрудник, кандидат физико–математических наук, сотрудник Лаборатории нейронных сетей Института математических проблем биологии РАН (Пущино, Московская область). Область научных интересов: нейросетевые модели в нейрофизиологии и психологии, теория динамических систем. Автор (соавтор) более 70 публикаций.

Вадим Вадимович ШМАТЧЕНКО, аспирант Пущинского государственного университета; Лаборатории нейронных сетей Института математических проблем биологии РАН (Пущино, Московская область); сотрудник отдела электронной микроскопии ООО «Карл Цейсс». Область научных интересов: математическое моделирование в нейрофизиологии, психологии и микробиологии, теория динамических систем.

А. Ю. ДОРОГОВ

Санкт-Петербургский государственный электротехнический университет
«ЛЭТИ»

E-Mail: dorv@lens.spb.ru

**БЫСТРЫЕ НЕЙРОННЫЕ СЕТИ: ПРОЕКТИРОВАНИЕ,
НАСТРОЙКА, ПРИЛОЖЕНИЯ¹**

Аннотация

Обсуждается парадигма быстрых нейронных сетей (БНС). Представлены лингвистические методы проектирования структуры и топологии БНС. На уровне морфологии установлены общие принципы построения быстрых преобразований. Найдены количественные оценки быстродействия и способности нейронной сети к обучению. Рассмотрены методы настройки БНС при реализации спектральных преобразований, регулярных фракталов, оптимальных фильтров. Исследованы вопросы применения БНС в квантовых вычислениях. Рассмотрены методы построения многомерных БНС.

A. Yu. DOROGOV

Saint-Petersburg state electrotechnical university "LETI"

E-Mail: dorv@lens.spb.ru

**FAST NEURAL NETWORKS: DESIGN, TUNING AND
APPLICATIONS**

Abstract

A paradigm of fast neural networks (FNN) is discussed. Formal linguistic methods are presented to design structure and topology for FNN. A general principle of fast transformations have been founded on morphological level. Some quantitative estimations are obtained for processing speed and learning ability of FNN. Tuning techniques are described to implement FNN-based spectral transformations, regular fractals and optimum filters. Capabilities of FNN for quantum computations are investigated. A design approach is suggested to build multi-dimensional FNN.

¹Работа выполнена при поддержке гранта Минобразования РФ.

Введение

Хорошо известно, какую огромную роль в обработке сигналов сыграло открытие алгоритма *быстрого преобразования Фурье* (БПФ). Использование БПФ позволило кардинальным образом уменьшить количество вычислительных операций при выполнении спектральных преобразований, что значительно расширило сферу использования спектрального анализа для обработки данных. Реализация БПФ в технологии больших интегральных схем приводит к существенному уменьшению площади кристаллов спектральных анализаторов и, соответственно, к радикальному уменьшению их энергопотребления.

Для нейронных сетей задача сокращения числа вычислительных операций не менее актуальна. При больших размерностях обрабатываемых данных нейронные сети с многослойной структурой требуют значительного объема вычислительных операций как при обработке данных, так и при обучении сетей. Это ограничивает применение больших нейронных сетей данного класса в системах реального времени и существенно усложняет аппаратную реализацию.

Алгоритмы БПФ имеют выраженную многослойную структуру, подобную структуре многослойных персептронов, поэтому напрашивается естественное решение — использовать идеологию БПФ для построения нейронных сетей [1]. Для этого достаточно в операциях «бабочка» БПФ заменить коэффициенты преобразования перестраиваемыми синаптическими весами и добавить нелинейные функции активации. Следует отметить, что первый шаг в этом направлении был сделан достаточно давно. Еще в исторической работе Гуда [2] (1958 г.) было приведено аналитическое описание быстрых алгоритмов для обобщенных спектральных преобразований. Обобщенное спектральное преобразование с позиций сегодняшнего дня можно рассматривать как нейронную сеть с линейными функциями активации. В последующие годы тема обобщенных спектральных преобразований развивалась в работах Г. Эндрюса, А. И. Солодовникова, В. Г. Лабунца и других авторов [3–6].

Алгоритм БПФ обладает двумя системными ограничениями: во-первых, размерность вектора обрабатываемых данных для всех слоев должна быть одинакова и, во-вторых, значение этой размерности должно быть составным числом, т. е. разлагаться в произведение целых множителей. Первое ограничение обусловлено ортогональностью матрицы БПФ, а второе — регулярностью алгоритма что, по-видимому, является неизбежной платой

за быстроедействие. Парадигма быстрых нейронных сетей наследует регулярность алгоритма БПФ, но расширяет ее возможности за счет отказа от ортогональности и жесткой схемы топологической реализации. Для структурно-регулярной сети может быть построена как регулярная, так и нерегулярная топология. На выбор типа топологии могут влиять два обстоятельства: структура данных и простота технической реализации. Для больших и сверхбольших нейронных сетей (потенциальная область применения БНС) простота технической реализации является определяющим условием. Поэтому регулярность топологии далее будет рассматриваться как составная компонента парадигмы БНС.

БНС представляют собой вариант многослойных сетей, поэтому для их обучения могут быть использованы градиентные методы типа Еггс Backpropagation, однако структурные особенности БНС позволяют упростить и процедуру обучения. Частным случаем БНС являются быстрые перестраиваемые линейные преобразования. От нейронных сетей они отличаются линейными функциями активации и отсутствием смещений. Ортогональные перестраиваемые преобразования с различной топологией традиционно используются для построения спектральных преобразований. В терминах перестраиваемых преобразований процедура обучения обычно называется настройкой. Типичной задачей для перестраиваемого преобразования является настройка на заданную систему базисных функций, например базис Фурье, Адамара–Уолша и др. Классические системы базисных функций, как правило, имеют аналитическую форму, поэтому настройка преобразования также выполняется в аналитическом виде. Поскольку структура БНС обладает фрактальными свойствами, то регулярные фрактальные последовательности могут быть реализованы как многослойные БНС. Фрактальность порождает и особые методы настройки БНС для реализации адаптивных фильтров. Идентичность структуры БНС со структурой тензорных произведений векторных пространств позволяет использовать БНС для построения алгоритмов квантовых вычислений. Парадигма БНС допускает многомерное обобщение, что наряду с высоким быстрыедействие способствует использованию БНС для построения классификаторов зрительных сцен. Перечисленные выше вопросы составляют основное содержание настоящей лекции. Цель лекции состоит в том, чтобы дать представление о методах построения БНС, оценить их характеристики и показать потенциальные области применения.

Структурный анализ алгоритмов БПФ

Набор алгоритмов, называемых алгоритмами быстрого преобразования Фурье вошел в практику спектрального анализа в 60-х годах, начиная с работ Кули-Тьюки [7]. Существует множество различных схем построения алгоритмов БПФ, как правило, они основаны на способе разбиения входного вектора на подвекторы. Следуя [8], рассмотрим в качестве примера схему «с прореживанием по частоте». Для вектора $X = \{x(i)\}$, $0 \leq i \leq N-1$ прямое дискретное преобразование Фурье (ДПФ) определяется выражением:

$$y(m) = \sum_{i=0}^{N-1} x(i) e^{-j \frac{2\pi}{N} im},$$

где $m = 0, 1, 2, \dots, N-1$ — номер гармоники. Величину $\omega = e^{-j \frac{2\pi}{N}}$ обычно называют поворачивающим множителем. Число комплексных операций умножения при выполнении ДПФ в общем случае равно N^2 .

Для алгоритмов БПФ, как правило, выбирают значение $N = 2^r$, где r — целое число. В схеме «с прореживанием по частоте» вектор $X = \{x(i)\}$ разбивается на два подвектора по правилу:

$$\begin{aligned} x_1(i) &= x(i), & x_2(i) &= x(i + N/2), \\ i &= 0, 1, 2, \dots, (N/2) - 1. \end{aligned}$$

Этот прием позволяет свести вычисление N -точечного преобразования к вычислению двух $N/2$ -точечных преобразований, при этом число комплексных операций умножения будет значительно меньше N^2 . Математические детали этой идеи подробно представлены в главе 6 работы [8]. На рис. 1 в графической форме показан пример перехода от восьмиточечного ДПФ к двум четырехточечным ДПФ при «прореживании по частоте».

Графический образ в виде пары перекрещивающихся линий с кружком посередине (напоминающий «бабочку») обозначает базовую операцию над парой координат входного вектора. Для алгоритма БПФ «с прореживанием по частоте» базовая операция представляется унитарной матрицей вида:

$$\begin{pmatrix} 1 & \omega^i \\ 1 & -\omega^i \end{pmatrix}.$$

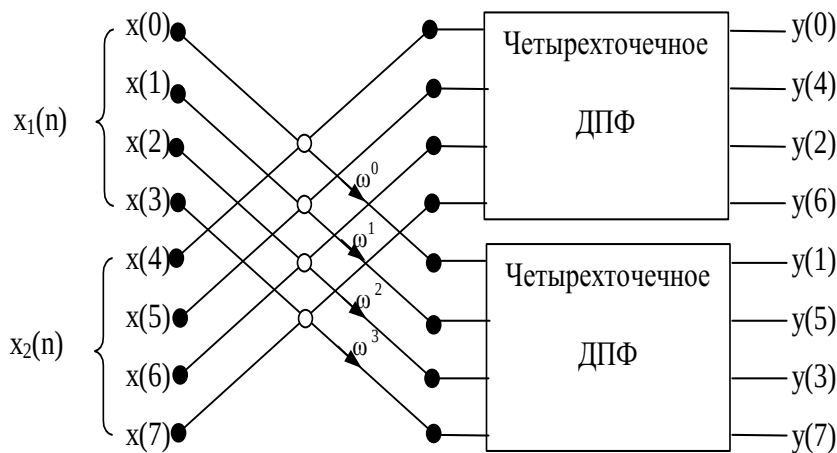


Рис. 1. Переход от восьмиточечного ДПФ к двум четырехточечным

В общем случае базовая операция определяется матрицей W размерности $p \times p$:

$$W = \begin{pmatrix} w_{00} & w_{01} & \cdots & w_{0,p-1} \\ w_{10} & w_{11} & \cdots & w_{1,p-1} \\ \cdots & \cdots & \cdots & \cdots \\ w_{p-1,0} & w_{p-1,1} & \cdots & w_{p-1,p-1} \end{pmatrix}$$

(здесь и далее принято, что нумерация строк и столбцов в матрицах начинается с нулевого индекса). Для базовой операции «бабочка» значение p равно двум, а для четырехточечного преобразования ДПФ будет $p = 4$.

Если сопоставить каждой базовой операции вершину графа, а дугам — операторы связи между базовыми операциями, то получим структурную модель алгоритма БПФ. На рис. 2 показана структурная модель для алгоритма, изображенного на рис. 1. На графе структурной модели приведены также размеры матриц базовых операций и ранги операторов связи (для алгоритмов БПФ ранги связей всегда равны единице). По структурной модели нетрудно подсчитать вычислительные затраты связанные с выполнением БПФ, для этого достаточно просуммировать число операций умножения по всем базовым операциям, например для структурной модели, показанной на рис. 2 получим:

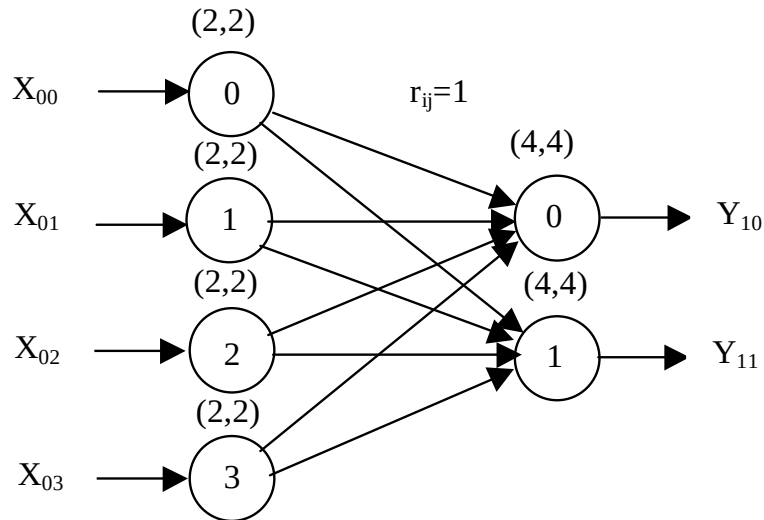


Рис. 2. Структурная модель БПФ первого шага

$$Z = 2 \times 2 + 2 \times 2 + 2 \times 2 + 2 \times 2 + 4 \times 4 + 4 \times 4 = 48.$$

Заметим для сравнения, что прямое восьмиточечное ДПФ имеет $8 \times 8 = 64$ операции умножения. Каждое дискретное преобразование размерности $N/2$ в свою очередь может быть сведено к двум $N/4$ точечным преобразованиям. На рис. 3 показан полный граф восьмиточечного БПФ с «прореживанием по частоте».

Для данного алгоритма спектральные коэффициенты на выходе преобразования оказываются переставленными в двоично-инверсном порядке по отношению к частоте гармоники. На рис. 4 приведена структурная модель для полного восьмиточечного БПФ. В общем случае размерность преобразования может быть составным числом: $N = n_0 n_1 \dots n_{\varkappa-1}$, где n_λ — целые числа. В этом случае структурная модель состоит из $\varkappa$ слоев с размерностью базовых операций (n_λ, n_λ) в слое λ и числом вершин в слое равным $k_\lambda = N/n_\lambda$.

Обозначим через i номер вершины во входном слое, а через j — номер вершины в выходном слое и представим эти числа в позиционной

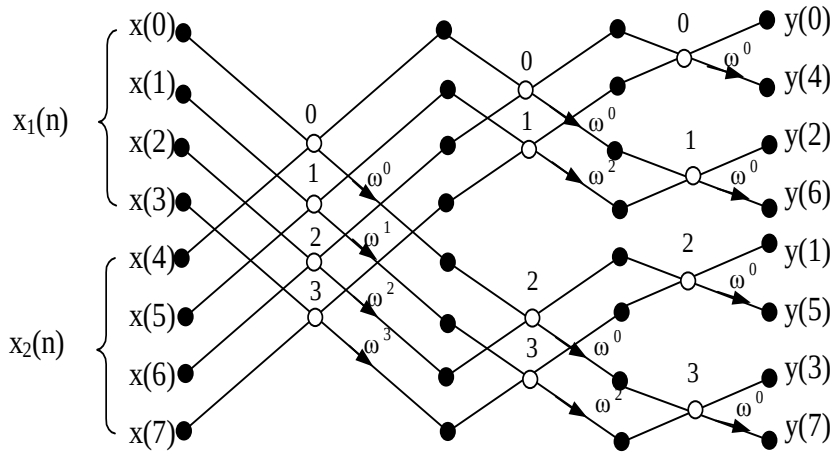


Рис. 3. Полный граф восьмиточечного БПФ с «прореживанием по частоте»

многоосновной системе счисления [9] следующим образом:

$$i = \langle i_{\lambda-2} i_{\lambda-3} \dots i_0 \rangle = i_{\lambda-2} n_{\lambda-3} n_{\lambda-4} \dots n_0 + \dots + i_1 n_0 + i_0,$$

$$j = \langle j_1 j_2 \dots j_{\lambda-1} \rangle = j_1 n_2 n_3 \dots n_{\lambda-1} + \dots + j_{\lambda-1} n_{\lambda-2} + j_{\lambda-1},$$

где $i_\lambda, j_\lambda \in [0, 1, \dots, (n_\lambda - 1)]$ — разрядные числа. Непосредственной проверкой можно убедиться, что нумерация вершин в графе структурной модели будет определяться выражением:

$$i^\lambda = \langle j_1 j_2 \dots j_\lambda i_{\lambda-2} i_{\lambda-3} \dots i_0 \rangle.$$

Из приведенного выше анализа можно сделать следующие выводы:

- На структурном уровне быстрый алгоритм представляется многослойной модульной сетью, где каждому модулю отвечает базовая операция.
- Размерность быстрого преобразования всегда выражается составным числом, где числовые множителями определяют размерности базовых операций.
- Базовые операции выражаются квадратными матрицами, которые в пределах слоя имеют одинаковый порядок.

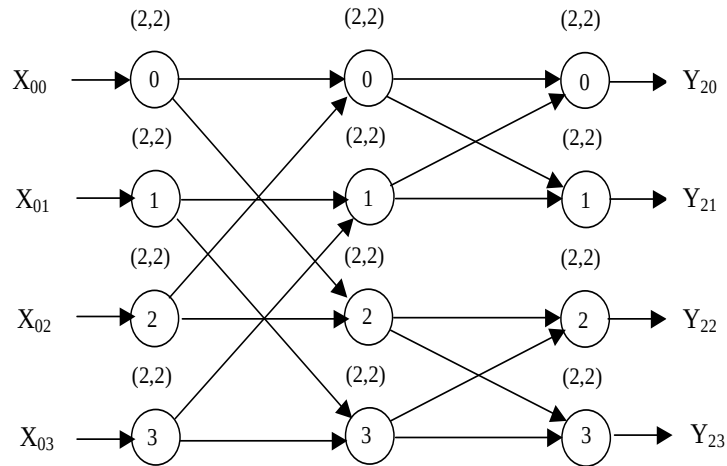


Рис. 4. Структурная модель восьмиточечного БПФ

- Ранги связей между базовыми операциями всегда равны единице.
- Алгоритм БПФ имеет регулярную структуру и регулярную топологию.

Структурный синтез быстрых нейронных сетей

В данном разделе рассматриваются лингвистические модели регулярных БНС на уровне структуры и морфологии. Лингвистическая модель основана на концепции формального языка. В контексте языка структура сети интерпретируется как графическая семантика допустимых предложений.

Формальный язык регулярных сетей

Любой формальный язык является средством для выражения инвариантных свойств предметной области знаний. Для построения языка вводится алфавит, грамматика и устанавливается семантика слов и предложений. Практическая ценность формального языка определяется его семантической интерпретацией, поэтому построение языка регулярных сетей целесообразно начать с его семантики.

Семантика слов. Пусть $\langle z_{n-1}z_{n-2} \dots z_0 \rangle$ — некоторое допустимое слово формального языка. Поставим в соответствие каждой букве z_i целое положительное число p_i , и будем рассматривать букву как символическое обозначение переменной принимающей целочисленные значения из интервала $Z_i = [0, 1, \dots, p_i - 1]$. Слово языка будем интерпретировать как целочисленную функцию, заданную правилом:

$$\langle z_{n-1}z_{n-2} \dots z_0 \rangle = z_{n-1}p_{n-2}p_{n-3} \dots p_0 + z_{n-2}p_{n-3}p_{n-4} \dots p_0 + \dots + z_1p_0 + z_0 \quad (1)$$

Порядок следования аргументов для данной функции имеет существенное значение, поэтому каждое слово языка представляет собой кортеж, т. е. множество, на котором зафиксирован линейный порядок. Для сокращенного обозначения кортежа все элементы которого имеют общее родовое имя (в данном случае — z), будем использовать следующую форму: $\langle z \rangle_J$, где J — отношение линейного порядка. Для определенности будем полагать, что собственные позиции кортежа всегда упорядочены по возрастанию слева направо, начиная с нулевой. В дальнейшем интерпретирующую функцию будем называть *конвергенцией*. Нетрудно видеть, что функция конвергенции представляет собой правило перехода от позиционного представления числа в многоосновной системе счисления к его количественному значению. В этом контексте переменную z_i уместно назвать разрядом, а константу p_i основанием разряда.

Алфавит языка и грамматика слов. Алфавитом языка будем считать кортеж неповторяющихся символов $A = \langle z_{i_0}z_{i_1} \dots z_{i_\kappa} \rangle$. Отношение линейного порядка на алфавитном множестве назовем *канон* языка и обозначим через K . В сокращенной форме кортеж алфавита, состоящий из букв с общим родовым именем, записывается в виде: $A = \langle z \rangle_K$. Канон языка может быть задан таблицей:

$$K = \begin{pmatrix} z_{i_0} & z_{i_1} & \dots & z_{i_\kappa} \\ 0 & 1 & \dots & \kappa \end{pmatrix}.$$

Нижняя строка таблицы определяет порядок символов в алфавите. Далее будем использовать алфавит с двумя различными типами букв. Каждый тип буквы определяется его родовым именем (например, u, v) в этом случае для сокращенной записи алфавита будем использовать обозначения вида:

$$A = \langle \langle u \rangle \oplus \langle v \rangle \rangle_K.$$

Пусть алфавит A состоит из $\varkappa$ букв. Допустимым словом языка будем считать любой кортеж длиной $n \leq \varkappa$, состоящий из неповторяющихся букв алфавита. Если последовательность букв в слове соответствует канону, то слово будем называть *каноническим*.

Операция перестановки. Пусть Q есть множество всевозможных перестановок букв, определенных на кортеже слова $\langle z \rangle_K$, и q — некоторая перестановка принадлежащая этому множеству. Операцию перестановки для слова $\langle z \rangle_K$ будем записывать в виде: $\langle z \rangle_K * q$, например,

$$\langle z_1 z_3 z_0 z_6 \rangle * \begin{pmatrix} 0 & 1 & 2 & 3 \\ 3 & 2 & 1 & 0 \end{pmatrix} = \langle z_6 z_0 z_3 z_1 \rangle.$$

На множестве слов языка определим отношение эквивалентности π следующим образом. Два слова $\langle z \rangle_{J_1}$ и $\langle z \rangle_{J_2}$ одинаковой длины будем считать эквивалентными, если существует перестановка q такая, что

$$\langle z \rangle_{J_1} * q = \langle z \rangle_{J_2}.$$

Отношение эквивалентности разбивает множество слов языка на классы. Понятно, что если $\langle z \rangle_{J_1}$ и $\langle z \rangle_{J_2}$ два различных канонических слова, то они всегда принадлежат разным классам, т.е. канонические слова являются представителями эквивалентных классов. Каждый класс может быть образован всевозможными перестановками букв канонического слова.

Смежные классы. Пусть задан алфавит длиной $\varkappa$. Рассмотрим множество эквивалентных классов с длиной слов, равной n . Если $n = \varkappa$, то множество состоит из одного класса. Для слов длиной $n = \varkappa - 1$ число возможных классов равно $\varkappa$. Два класса на множестве слов длиной n назовем *смежными*, если они отличаются только одной буквой. Смежные классы образуют упорядоченные цепочки по подмножествам входящих букв. Длина каждой цепочки не превышает $n + 1$. Цепочки смежных классов длиной $n + 1$ будем называть *максимальными*.

Грамматика предложений и морфология сети

Правила построения предложений (иначе эти правила называются порождающей грамматикой) целиком подчинены задачам предметной области. Порождающая грамматика составляет основу метода синтеза регулярных многослойных графов.

В линейных языках каждое предложение имеет одно начало и один конец. Возможны различные способы кодирования границ предложения. Будем использовать для этой цели вариант родовых имен. Определим алфавит языка как линейно упорядоченный набор неповторяющихся букв с двумя видами родовых имен:

$$A = \langle \langle i \rangle_I \oplus \langle j \rangle_J \rangle_K.$$

Число букв каждого вида будем считать одинаковыми и равными $\varkappa - 1$. Таким образом, длина алфавита равна $2(\varkappa - 1)$. Рассмотрим множество всевозможных слов длиной $(\varkappa - 1)$, из которых будем строить предложения. Как и прежде будем полагать, что на множестве слов операций перестановки определено отношение эквивалентности π , разбивающее это множество на непересекающиеся классы. Допустимым предложением будем считать упорядоченную последовательность слов, каждое из которых является представителем максимальной цепочки смежных классов. В допустимом предложении представлены все смежные классы цепочки, и каждый класс представлен только одним словом. Длина предложения всегда равна $\varkappa$. Последовательность слов в предложении определяется правилом вывода. Это правило устанавливает начальное слово предложения и описывает связи между смежными словами. Зафиксируем следующее правило вывода:

- слова в предложении линейно упорядочены по числу букв с родовым именем j ;
- первое слово предложения выбирается из класса $\langle i \rangle_I$ и имеет, поэтому нулевое число вхождений букв с родовым именем j .

Каждая цепочка смежных классов порождает множество однотипных предложений. Среди этого множества существует единственное предложение, составленное из канонических слов, такое предложение назовем *каноническим*. Все остальные предложения для данной цепочки могут быть получены произвольной перестановкой букв в словах канонического предложения. Поэтому достаточно рассматривать порождающую грамматику только для канонических предложений. Пусть числа λ и $\lambda + 1$ определяют порядковые номера двух смежных классов в цепочке. Обозначим через I_λ и J_λ канонически упорядоченные подмножества букв каждого типа для смежного класса λ , а через $\langle \langle i \rangle_{I_\lambda} \oplus \langle j \rangle_{J_\lambda} \rangle$ его каноническое слово. Тогда правило вывода для канонических предложений можно сформулировать

следующим образом:

$$\begin{aligned} I_\lambda \supset I_{\lambda+1}, & \quad J_\lambda \subset J_{\lambda+1}, \\ I_0 = I, & \quad I_{\kappa-1} = \emptyset, \quad J_{\kappa-1} = J, \quad J_0 = \emptyset. \end{aligned}$$

Построенное правило не связано ни с размерностью сети, ни с ее топологией, ни со структурными характеристиками вершин — это правило морфологического уровня, которое в аксиоматической форме раскрывает внутреннюю сущность быстрых преобразований.

Семантическая интерпретация предложений

Поставим в соответствие каноническому предложению n -слойный сетевой граф. Каждое слово с порядковым номером λ будем интерпретировать как функциональное правило вычисления номера вершины в слое λ , в соответствии с семантикой (1). Вершины двух смежных слоев будем считать связанными дугой, если в поразрядном представлении номеров вершин

$$i^\lambda = \langle \langle i \rangle_{I_\lambda} \oplus \langle j \rangle_{J_\lambda} \rangle, \quad i^{\lambda+1} = \langle \langle i \rangle_{I_{\lambda+1}} \oplus \langle j \rangle_{J_{\lambda+1}} \rangle$$

одноименные разрядные переменные имеют совпадающие значения. Рассмотрим пример построения сетевого графа. Пусть алфавит языка определен кортежем

$$A = \langle i_1 i_2 i_3 j_2 j_1 j_0 \rangle.$$

Нетрудно убедиться, что предложение

$$[\langle i_1 i_2 i_3 \rangle \langle i_1 i_2 j_0 \rangle \langle i_1 j_1 j_0 \rangle \langle j_2 j_1 j_0 \rangle] \quad (2)$$

является каноническим. Будем полагать, что основания разрядных переменных заданы таблицей

$$\begin{pmatrix} A \\ p \end{pmatrix} = \begin{pmatrix} i_1 & i_2 & i_3 & j_2 & j_1 & j_0 \\ 2 & 2 & 2 & 2 & 2 & 2 \end{pmatrix}.$$

В соответствии с семантикой (1) слова предложения интерпретируются как следующие алгебраические выражения:

$$\begin{aligned} i^0 &= \langle i_1 i_2 i_3 \rangle = 2^2 i_1 + 2 i_2 + i_3, \\ i^1 &= \langle i_1 i_2 j_0 \rangle = 2^2 i_1 + 2 i_2 + j_0, \\ i^2 &= \langle i_1 j_1 j_0 \rangle = 2^2 i_1 + 2 j_1 + j_0, \\ i^3 &= \langle j_2 j_1 j_0 \rangle = 2^2 j_2 + 2 j_1 + j_0. \end{aligned}$$

Эти формулы используются для вычисления номеров вершин сетевого графа. На рис. 5 приведены поразрядные представления номеров вершин для каждого слоя графа и частично показаны принципы образования межслойных связей. На рис. 6 показан полный граф, соответствующий каноническому предложению (2).

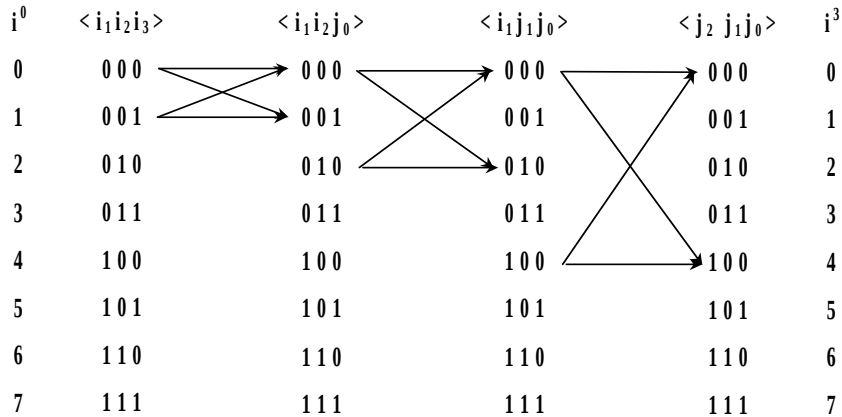


Рис. 5. Правило построения графической интерпретации предложения

Топологическое проектирование быстрых нейронных сетей

Рассмотрим регулярную сеть и будем полагать, что каждая вершина слоя λ представляет собой однослойный персептрон с размерностью рецепторного поля, равной p_λ и числом нейронов — g_λ . Персептрон, соответствующий вершине графа, будем называть *нейронным ядром*. Обозначим через $u_\lambda = [0, 1, \dots, p_\lambda - 1]$ номер рецептора, а через $v_\lambda = [0, 1, \dots, g_\lambda - 1]$ номер нейрона в нейронном ядре. Будем полагать, что любой нейрон имеет только один выходной аксон. Каждой дуге графа поставим в соответствие точную однозначную связь аксон-рецептор. Алфавит языка расширим буквами с родовыми именами u и v .

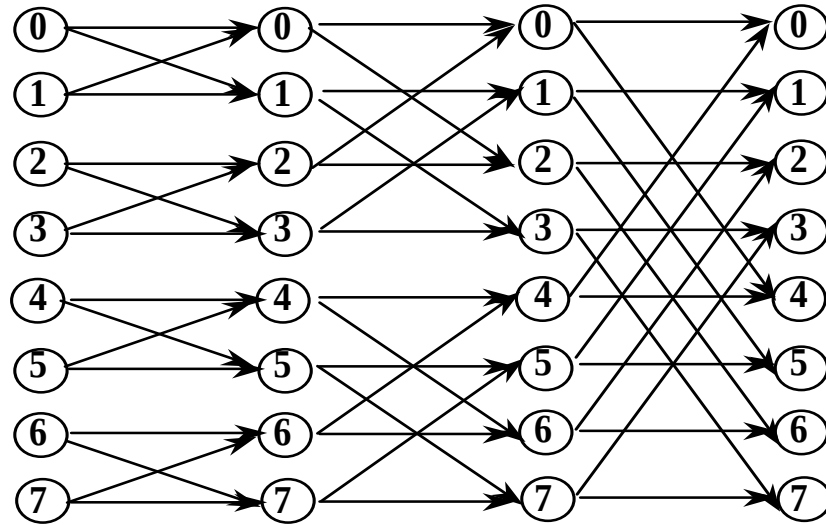


Рис. 6. Графическая интерпретация канонического предложения

Топологии нейронных слоев

Позиционный номер рецепторов и аксонов в пределах слоя может быть задан конвергенциями

$$U^\lambda = \langle \langle i^\lambda \rangle \oplus u_\lambda \rangle, \quad V^\lambda = \langle \langle i^\lambda \rangle \oplus v_\lambda \rangle. \quad (3)$$

Данные выражения включают в себя поразрядное представление номера вершины и дополнительный разряд, соответствующий локальному номеру рецептора/аксона в пределах нейронного ядра. Дополнительный разряд может занимать любую позицию в кортеже слова, что позволяет в широком диапазоне варьировать размещение терминальных полей нейронного ядра в нейронном слое. Слова языка, соответствующие конвергенциям (3), назовем *топологиями* нейронного слоя. Для смежного слоя с порядковым номером $\lambda + 1$ топологии определяются конвергенциями:

$$U^{\lambda+1} = \langle \langle i^{\lambda+1} \rangle \oplus u_{\lambda+1} \rangle, \quad V^{\lambda+1} = \langle \langle i^{\lambda+1} \rangle \oplus v_{\lambda+1} \rangle.$$

Потребуем, чтобы соответствие между аксонами и рецепторами двух смежных слоев удовлетворяло условиям регулярности. В контексте фор-

мального языка регулярное соответствие может быть выражено операцией перестановки

$$\langle \langle i^\lambda \rangle \oplus v_\lambda \rangle * q^\lambda = \langle \langle i^{\lambda+1} \rangle \oplus u_{\lambda+1} \rangle.$$

Это выражение справедливо для всех значений разрядных переменных, т. е. является тождеством. Топологию будем называть компактной, если все межслойные перестановки тождественны единичной. Сеть с компактной топологией наиболее проста в технической реализации. Учитывая, что $I_\lambda \supset I_{\lambda+1}$, $J_\lambda \subset J_{\lambda+1}$, можно сделать вывод о существовании взаимно однозначных соответствий

$$I_\lambda \setminus I_{\lambda+1} \leftrightarrow u_{\lambda+1}, \quad J_{\lambda+1} \setminus J_\lambda \leftrightarrow v_\lambda, \quad (4)$$

(символ “ $\setminus$ ” означает разность множеств). Соответствия (4) позволяют во всех словах допустимого предложения заменить буквы вида i, j на буквы с родовыми именами вида u, v , в результате получим:

$$V^\lambda = \langle \langle u \rangle_{S_\lambda} \oplus \langle v \rangle_{C_\lambda} \rangle, \quad U^{\lambda+1} = \langle \langle u \rangle_{S_\lambda} \oplus \langle v \rangle_{C_\lambda} \rangle * q^\lambda,$$

где $S_\lambda \supset S_{\lambda+1}$ и $C_\lambda \subset C_{\lambda+1}$ — подмножества смежных классов в алфавите u, v .

Терминальные поля нейронной сети занимают особое положение. Обозначим через U позиционный номер рецептора входного слоя, а через V позиционный номер аксона выходного слоя, тогда топологии терминальных полей можно записать в виде конвергенций:

$$U = \langle \langle i^0 \rangle \oplus u_0 \rangle, \quad V = \langle \langle i^{\varkappa-1} \rangle \oplus v_{\varkappa-1} \rangle.$$

В этих выражениях u_0 и $v_{\varkappa-1}$ — свободные переменные, которые не связаны каким либо соответствием вида (4) с переменными $\langle i \rangle, \langle j \rangle$. Число букв с родовыми именами u, v на единицу больше (по каждому виду), чем число букв с родовыми именами i, j .

Графическая интерпретация топологий

Пусть алфавит языка определен кортежем

$$A = \langle i_1 i_2 i_3 j_2 j_1 j_0 \rangle.$$

Рассмотрим каноническое предложение

$$[\langle i_1 i_2 i_3 \rangle \langle i_1 i_2 j_0 \rangle \langle i_1 j_1 j_0 \rangle \langle j_2 j_1 j_0 \rangle], \quad (5)$$

определяющее структурную модель регулярной сети. Используя соответствия (4), сделаем замены переменных

$$\begin{aligned} I_0 \setminus I_1 = i_3 \leftrightarrow u_1, & \quad I_1 \setminus I_2 = i_2 \leftrightarrow u_2, & \quad I_2 \setminus I_3 = i_1 \leftrightarrow u_3, \\ J_3 \setminus J_2 = j_2 \leftrightarrow v_2, & \quad J_2 \setminus J_1 = j_1 \leftrightarrow v_1, & \quad J_1 \setminus J_0 = j_0 \leftrightarrow v_0, \end{aligned} \quad (6)$$

В новых переменных каноническое предложение примет вид:

$$[\langle u_3 u_2 u_1 \rangle \langle u_3 u_2 v_0 \rangle \langle u_3 v_1 v_0 \rangle \langle v_2 v_1 v_0 \rangle].$$

Для построения топологий рецепторных полей необходимо к каждому слову с порядковым номером λ добавить переменную u_λ . Переменные могут быть добавлены в любые позиции, например, добавим их в конец слов, тогда множество рецепторных топологий будет описываться предложением

$$[\langle u_3 u_2 u_1 u_0 \rangle \langle u_3 u_2 v_0 u_1 \rangle \langle u_3 v_1 v_0 u_2 \rangle \langle v_2 v_1 v_0 u_3 \rangle]. \quad (7)$$

Аналогичным образом для построения топологий аксоновых полей необходимо к каждому слову канонического предложения с порядковым номером λ добавить переменную v_λ . Например, это можно сделать следующим образом:

$$[\langle u_3 u_2 v_0 u_1 \rangle \langle u_3 v_1 v_0 u_2 \rangle \langle v_2 v_1 v_0 u_3 \rangle \langle v_3 v_2 v_1 v_0 \rangle]. \quad (8)$$

Все другие топологии могут быть получены всевозможными перестановками букв в словах полученных предложений. Предложения (7) и (8) будем называть *топологическими*, а их совокупность — *топологической траекторией*. Топологическую траекторию можно интерпретировать как сетевой многослойный граф, каждой вершине которого соответствует один нейрон, а дуги определяют связи между нейронами смежных слоев. Каждое слово в предложении (8) определяет правило для вычисления номера нейрона в пределах слоя. Слова предложения (7) интерпретируются как номер рецептора в пределах слоя. Рецептор и нейрон смежных слоев связываются между собой дугой графа, если значения одноименных разрядов в соответствующих конвергенциях совпадают. Правило построения графа показано на рис. 7. В данном примере основания всех разрядных переменных приняты равными двум. На рис. 8 приведен полный топологический граф быстрой нейронной сети. Для построения компактной топологической траектории можно воспользоваться правилом скользящего интервала. Выберем канон алфавита в виде

$$A = \langle u_0 u_1 \dots u_{\kappa-1} v_0 v_1 \dots v_{\kappa-2} v_{\kappa-1} \rangle.$$

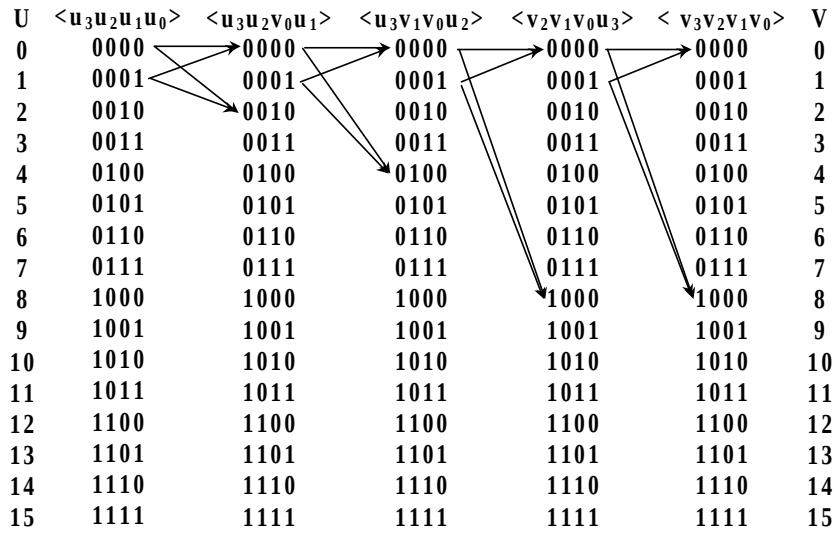


Рис. 7. Правило построения топологического интерпретирующего графа

Будем перемещать пустой интервал длиной $\varkappa$, ограниченный угловыми скобками вдоль алфавитного кортежа, начиная с крайней левой позиции. На каждом шаге подвижный интервал вычленяет из кортежа каноническое слово. В результате последовательность первых $\varkappa - 1$ порождаемых слов будет определять топологию рецепторных полей нейронной сети, а последнее $\varkappa$ -е слово — топологию терминального аксонового поля сети. Обратное движение интервала с крайне правой позиции к началу канона порождает топологии аксоновых полей сети.

Правило скользящего интервала можно представить в аналитической форме. Обозначим через t^λ слово, выделяемое подвижным интервалом на шаге λ . Нетрудно заметить, что

$$t^\lambda = \langle u_\lambda u_{\lambda+1} \dots u_{\varkappa-1} v_0 v_1 \dots v_{\lambda-1} \rangle. \quad (9)$$

Можно проверить, что данное правило соответствует схеме алгоритма БПФ в форме Гуда (см. рис. 9). Особенность порождающей схемы Гуда состоит в том, что все слои графа имеют одинаковую топологию, что упрощает алгоритм реализации. Однако выходные операнды на каждой базовой операции

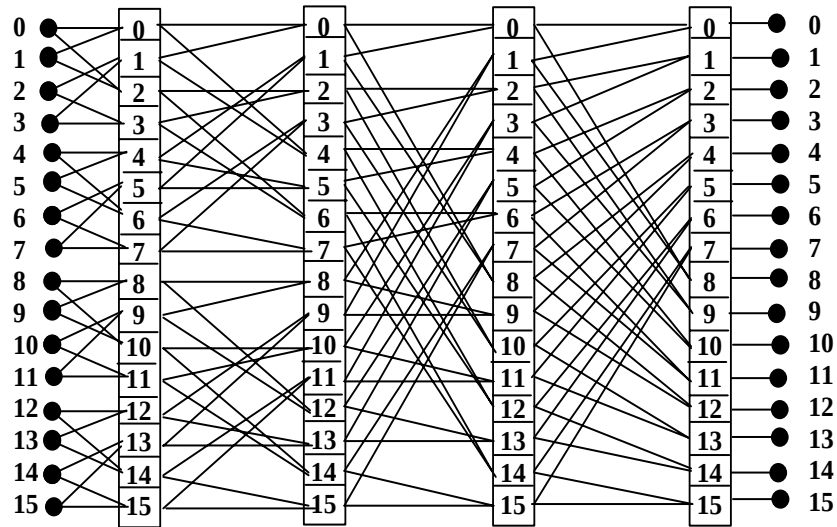


Рис. 8. Графическая интерпретация топологической траектории

не могут замещать входные отсчеты, поэтому требуется дополнительная память для хранения выходной последовательности. Далее по тексту будут представлены другие варианты порождающих схем (см. также [10]).

Граничные условия и топологические матрицы

В общем случае топологии терминальных полей сети могут быть записаны в следующем виде:

$$U = \langle u \rangle_S * q_b, \quad V = \langle v \rangle_C * q_a, \quad (10)$$

где q_a, q_b — регулярные перестановки. Векторы на входе и выходе сети имеют свою собственную нумерацию координат, для которой также будем использовать поразрядную форму, положив

$$U = \langle U \rangle_X \quad \text{и} \quad V = \langle V \rangle_Y,$$

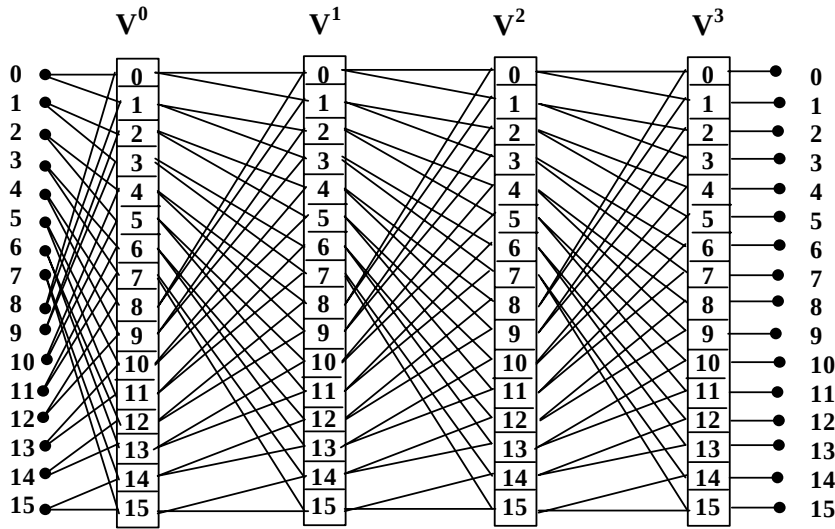


Рис. 9. Топологическая реализация схемы Гуда

где X, Y — линейно упорядоченные множества разрядных переменных. Например, можно выбрать

$$U = \langle U_{\varkappa-1} U_{\varkappa-2} \dots U_0 \rangle, \quad V = \langle V_{\varkappa-1} V_{\varkappa-2} \dots V_0 \rangle.$$

Разрядные числа U_i, V_i назовем *глобальными* разрядными числами в отличие от u_i, v_i для которых будем в дальнейшем использовать определение «*локальные*». В данном случае к определению «глобальные» следует добавить и определение «терминальные», поскольку они описывают граничные слои сети. Выражения (10) устанавливают соответствия между глобальными и локальными разрядными числами. Эти соответствия можно рассматривать как граничные условия для топологической траектории нейронной сети. Подобным образом можно ввести глобальные разрядные переменные для всех скрытых слоев, положив:

$$U^\lambda = \langle U_{\varkappa-1}^\lambda U_{\varkappa-2}^\lambda \dots U_0^\lambda \rangle, \quad V = \langle V_{\varkappa-1}^\lambda V_{\varkappa-2}^\lambda \dots V_0^\lambda \rangle.$$

Для каждого слоя справедливы следующие соответствия между локальными и глобальными переменными:

$$\langle U_{\kappa-1}^\lambda U_{\kappa-2}^\lambda \dots U_0^\lambda \rangle = t^\lambda * q_b^{\lambda-1}, \quad \langle V_{\kappa-1}^\lambda V_{\kappa-2}^\lambda \dots V_0^\lambda \rangle = t^\lambda * q_a^\lambda \quad (11)$$

(разложение $q^\lambda = (q_a^\lambda)^{-1} q_b^\lambda$ использовано из соображений симметрии). Последовательность слов вида $\langle \langle U \rangle_{X_\lambda} \oplus \langle V \rangle_{Y_\lambda} \rangle$ в алфавите терминальных переменных, соответствующая последовательности шагов топологической траектории, назовем *внешней траекторией топологий*. Пусть t^λ слово топологического предложения, определяемое некоторой регулярной порождающей схемой. Ограничимся компактными топологиями и представим соответствия (11) в табличном виде. Например, для топологии Гуда:

$$\begin{aligned} U^\lambda = t^\lambda &= \langle u_\lambda u_{\lambda+1} \dots u_{\kappa-2} u_{\kappa-1} v_0 v_1 \dots v_{\lambda-1} \rangle, \\ V^\lambda = t^{\lambda+1} &= \langle u_{\lambda+1} \dots u_{\kappa-2} u_{\kappa-1} v_0 v_1 \dots v_{\lambda-1} v_\lambda \rangle. \end{aligned} \quad (12)$$

Данное соответствие представляется таблицей 1. Построим по данной таблице матрицу, T_λ состоящую из нулей и единиц, используя при этом следующее правило: элемент матрицы $T_\lambda(U^\lambda, V^\lambda)$ равен единице, когда одноименные разрядные числа в конвенциях U^λ, V^λ совпадают по значениям. Следуя данному правилу и используя первую и третью строки таблицы, можно получить следующую аналитическую форму для определения элементов матрицы

ТАБЛИЦА 1. Соответствие локальных и глобальных переменных

$U^\lambda =$	u_λ	$u_{\lambda+1}$	$\dots$	$u_{\kappa-2}$	$u_{\kappa-1}$	v_0	$\dots$	$v_{\lambda-3}$	$v_{\lambda-2}$	$v_{\lambda-1}$
$U^\lambda =$	$U_{\kappa-1}^\lambda$	$U_{\kappa-2}^\lambda$	$\dots$	$U_{\kappa-\lambda+1}^\lambda$	$U_{\kappa-\lambda}^\lambda$	$U_{\kappa-\lambda-1}^\lambda$	$\dots$	U_2^λ	U_1^λ	U_0^λ
$V^\lambda =$	$u_{\lambda+1}$	$u_{\lambda+2}$	$\dots$	$u_{\kappa-1}$	v_0	v_1	$\dots$	$v_{\lambda-2}$	$v_{\lambda-1}$	v_λ
$V^\lambda =$	$V_{\kappa-1}^\lambda$	$V_{\kappa-2}^\lambda$	$\dots$	$V_{\kappa-\lambda+1}^\lambda$	$V_{\kappa-\lambda}^\lambda$	$V_{\kappa-\lambda-1}^\lambda$	$\dots$	V_2^λ	V_1^λ	V_0^λ

$$T_{\lambda}(U^{\lambda}, V^{\lambda}) = \delta(U_{\kappa-2}^{\lambda}, V_{\kappa-1}^{\lambda}) \delta(U_{\kappa-3}^{\lambda}, V_{\kappa-2}^{\lambda}) \dots \delta(U_0^{\lambda}, V_1^{\lambda}),$$

где $\delta(\cdot)$ — функция Кронекера. Например, для четырехслойной сети с основанием 2 (см. рис. 9) топологическая матрица слоя λ будет определяться выражением

$$T_{\lambda}(U^{\lambda}, V^{\lambda}) = \delta(U_2^{\lambda}, V_3^{\lambda}) \delta(U_1^{\lambda}, V_2^{\lambda}) \delta(U_0^{\lambda}, V_1^{\lambda}).$$

На рис. 10 показан вид топологической матрицы, отвечающей данному выражению; там же для удобства показана поразрядная нумерация строк и столбцов. На поле матрицы прямоугольникам выделены элементы, относящиеся к одному нейронному ядру, незаполненные позиции матрицы соответствуют нулевым значениям.

По существу данная матрица представляет собой шаблон синаптической карты нейронного слоя. При реализации нейронной сети ненулевые позиции топологической матрицы занимают синаптическим весами нейронных ядер. Обозначим матрицу синаптической карты через H_{λ} . Используя аналитическую форму для топологической матрицы и приведенную выше таблицу, можно получить следующее аналитическое представление для элементов матрицы H_{λ} :

$$h_{\lambda}(U^{\lambda}, V^{\lambda}) = w_{i_{\lambda}}^{\lambda}(U_{\kappa-1}^{\lambda}, V_0^{\lambda}) \delta(U_{\kappa-2}^{\lambda}, V_{\kappa-1}^{\lambda}) \delta(U_{\kappa-3}^{\lambda}, V_{\kappa-2}^{\lambda}) \dots \dots \delta(U_0^{\lambda}, V_1^{\lambda}),$$

здесь через $w_{i_{\lambda}}^{\lambda}(\cdot)$ обозначен элемент синаптической карты нейронного ядра i^{λ} . Номер ядра, выраженный в глобальных переменных, может быть представлен одним из выражений:

$$i^{\lambda} = \langle U_{\kappa-2}^{\lambda} U_{\kappa-3}^{\lambda} \dots U_0^{\lambda} \rangle = \langle V_{\kappa-1}^{\lambda} V_{\kappa-2}^{\lambda} \dots V_1^{\lambda} \rangle.$$

Алгоритм обработки данных в БНС

Обработка данных в нейронном ядре i^{λ} слоя λ определяется выражениями:

$$\begin{aligned} s_{i_{\lambda}}^{\lambda}(v_{\lambda}) &= \sum_{u_{\lambda}} x_{i_{\lambda}}^{\lambda}(u_{\lambda}) w_{i_{\lambda}}^{\lambda}(u_{\lambda}, v_{\lambda}), \\ y_{i_{\lambda}}^{\lambda}(v_{\lambda}) &= f_{v_{\lambda}}^{\lambda}(s_{i_{\lambda}}^{\lambda}(v_{\lambda})), \end{aligned} \quad (13)$$

V_3	0	0	0	0	0	0	0	1	1	1	1	1	1	1	1	1
V_2	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1
V_1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1
V_0	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1
$U_3U_2U_1U_0$																
0000	1	1														
0001			1	1												
0010					1	1										
0011							1	1								
0100									1	1						
0101											1	1				
0110													1	1		
0111															1	1
1000	1	1														
1001			1	1												
1010					1	1										
1011							1	1								
1100									1	1						
1101											1	1				
1110													1	1		
1111															1	1

Рис. 10. Топологическая матрица для порождающей схемы Гуда

где $f_{v_\lambda}^\lambda(\cdot)$ — функция активации, $x_{i_{\lambda+1}}^{\lambda+1}(v_\lambda)$ и $y_{i_\lambda}^\lambda(v_\lambda)$ — входные и выходные переменные. Для построения алгоритма обработки данных необходимо выполнить переход от локальных переменных нейронного ядра к глобальным переменным нейронного слоя. Рассмотрим в качестве примера нейронную сеть с топологией Гуда. Кортеж соответствующий номеру ядра можно получить, удалив из топологического слова t^λ (см. выражение (9)) букву u_λ . В результате получим:

$$i^\lambda = \langle u_{\lambda+1} \dots u_{\varkappa-1} v_0 v_1 \dots v_{\lambda-1} \rangle. \quad (14)$$

Следуя семантике слов, глобальный номер рецептора и глобальный номер аксона в слое λ можно вычислить, используя алгебраические выражения:

$$\begin{aligned} U^\lambda &= u_\lambda p_{\lambda+1} p_{\lambda+2} \dots p_{\varkappa-1} g_0 g_1 \dots g_{\lambda-1} + i^\lambda, \\ V^\lambda &= i^\lambda g_\lambda + v. \end{aligned} \quad (15)$$

Формулы (13), (14), (15) полностью определяют обработку данных в нейронном слое. В целом для нейронной сети алгоритм заключается в последовательной обработке данных по нейронным слоям.

Слабосвязанные сети

Условие слабой связанности [11–13] выражает собой отношение между окрестностями вершин и их проекциями на терминальные слои. В терминах нейронных сетей входной терминальный слой называется *афферентом*, а выходной — *эфферентом*. Под афферентной проекцией некоторой вершины B (далее обозначается $\text{Afr}(B)$) понимается подмножество вершин афферентного слоя, связанных дугами с этой вершиной. Аналогично под эфферентной проекцией вершины B (далее обозначается $\text{Efr}(B)$) понимается подмножество вершин эфферентного слоя, связанных дугами с этой вершиной. Определим также рецепторную окрестность $\Gamma^{-1}(B)$ как множество вершин непосредственно предшествующего слоя, связанных дугами с вершиной B и аксоновую окрестность $\Gamma(B)$, как множество вершин последующего слоя, связанных дугами с вершиной B . На рис. 11 показан пример структурной модели трехслойной ядерной нейронной сети. Для вершины B_1 данного примера терминальные проекции будут иметь следующий состав:

$$\text{Afr}(B_1) = \{A_3, A_4, A_5\}, \quad \text{Efr}(B_1) = \{C_0, C_1\}.$$

Принцип слабой связанности может быть выражен парой симметричных условий:

$$\begin{aligned} \text{Afr}(B) &= \sum_{A \in \Gamma^{-1}(B)} \text{Afr}(A), \\ \text{Efr}(B) &= \sum_{C \in \Gamma(B)} \text{Efr}(C), \end{aligned}$$

где символ $\sum$ обозначает прямую сумму соответствующих множеств. Первое выражение устанавливает, что в слабосвязанной сети для любой вершины B афферентные проекции вершин ее рецепторной окрестности не пересекаются. Второе — устанавливает аналогичное условие для эфферентных проекций. Фактически оба выражения двойственны друг другу и если выполняется одно из них, то обязательно будет выполнено и другое. Используя данное определение несложно показать, что в слабосвязанных сетях отсутствуют параллельные пути между вершинами, т. е. сети являются структурно безызбыточными.

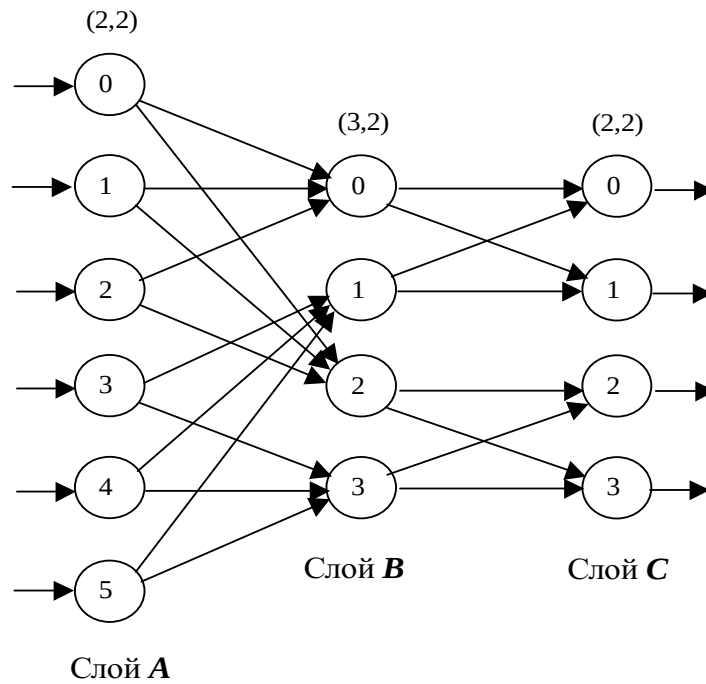


Рис. 11. Переход от восьмиточечного ДПФ к двум четырехточечным

Докажем теперь, что грамматика формального языка порождает слабо-связанные структуры. В интерпретирующем графе регулярной сети номер вершины в слое λ определяется выражением

$$i^\lambda = \langle \langle i \rangle_{I_\lambda} \oplus \langle j \rangle_{J_\lambda} \rangle, \quad (16)$$

а номера вершин начального и конечного слоя можно представить в виде:

$$i = \langle \langle i \rangle_{I_\lambda} \oplus \langle i \rangle_{I \setminus I_\lambda} \rangle, \quad j = \langle \langle j \rangle_{J_\lambda} \oplus \langle j \rangle_{J \setminus J_\lambda} \rangle, \quad (17)$$

где $I \setminus I_\lambda$ и $J \setminus J_\lambda$ обозначают разность множеств. По принципу построения интерпретирующего графа вершины двух смежных слоев связаны дугой, если одноименные переменные в смежных словах предложения имеют совпадающие значения. Выполняя индукцию по номеру слоя, нетрудно пока-

зат, что вершина i^λ будет связана с теми вершинами начального и конечного слоя, для которых одноименные разрядные переменные в выражениях (17) будут иметь те же значения, что и в выражении (16). Следовательно, терминальные проекции вершины i^λ можно выразить следующими параметрическими формами:

$$\text{Afr}(i^\lambda) = \langle \langle i \rangle_{I_\lambda} \oplus \langle i \rangle_{I \setminus I_\lambda} \rangle, \quad \text{Efr}(i^\lambda) = \langle \langle j \rangle_{J_\lambda} \oplus \langle j \rangle_{J \setminus J_\lambda} \rangle.$$

Здесь круглыми скобками выделены варьируемые разряды терминальных слов. Тогда для афферента вершины $i^{\lambda-1}$ слоя $\lambda - 1$ можно записать

$$\text{Afr}(i^{\lambda-1}) = \langle \langle i \rangle_{I_{\lambda-1}} \oplus \langle i \rangle_{I \setminus I_{\lambda-1}} \rangle. \quad (18)$$

Будем полагать, что вершины i^λ и $i^{\lambda-1}$ связаны между собой. По правилам порождающей грамматики $I_\lambda \subset I_{\lambda-1}$, причем разность множеств $I_{\lambda-1} \setminus I_\lambda$ состоит точно из одной буквы (разряда). Из свойств интерпретирующей функции (1) следует

$$\text{Afr}(i^\lambda) = \bigoplus_{I_{\lambda-1} \setminus I_\lambda} \text{Afr}(i^{\lambda-1}), \quad (19)$$

где прямая сумма подмножеств берется по всем значениям разрядного числа $I_{\lambda-1} \setminus I_\lambda$. Аналогично для эфферентов вершин смежных слоев λ и $\lambda + 1$ можно получить:

$$\text{Efr}(i^\lambda) = \bigoplus_{J_{\lambda+1} \setminus J_\lambda} \text{Efr}(i^{\lambda+1}). \quad (20)$$

Выражения, (19), (20) показывают, что интерпретирующий граф представляет собой слабосвязанную сеть.

Можно указать ряд причин, стимулирующих интерес к слабосвязанным нейронным сетям:

- К классу слабосвязанных сетей принадлежат как БНС, так и обычные многослойные полносвязанные сети.
- При реализации нейронных сетей на суперкомпьютерах структура нейронной сети подчиняется структуре вычислительной системы. Такие известные структуры суперкомпьютеров как Vapuan, «Омега», и др. [14] являются слабосвязанными сетями или строятся на основе слабосвязанных сетей.

- В слабосвязанных сетях реализуется принцип генетического подоби́я структуры [15], по заключениям нейрофизиологов [16] подобный принцип явился основным порождающим механизмом в эволюции новой коры головного мозга.
- Слабосвязанные сети обладают фрактальными свойствами [17].

Системные характеристики БНС

Вычислительная эффективность

Быстродействие вычислительных алгоритмов, как правило, оценивается числом операций умножения, необходимых для их реализации. Синаптическая карта нейронного ядра с размерностью рецепторного поля p и числом нейронов g представляет собой матрицу размерности $p \times g$. Обработка входного вектора в нейронном ядре соответствует умножению вектора на матрицу с последующим нелинейным преобразованием, определяемым функциями активации. Для выполнения матричного умножения требуется $p \cdot g$ скалярных операций умножения. Операции умножения, связанные с вычислением функций активации обычно не учитываются.

Рассмотрим $\varkappa$ -слойную БНС с характеристиками нейронных ядер:

$$(p_0, g_0), (p_1, g_1), \dots, (p_{\varkappa-1}, g_{\varkappa-1}).$$

Сеть имеет размерность по входу $N = p_0 p_1 \dots p_{\varkappa-1}$, а по выходу $M = g_0 g_1 \dots g_{\varkappa-1}$. Из (14) следует, что число ядер в слое m будет равно произведению

$$k_m = p_{m+1} p_{m+2} \dots p_{\varkappa-1} g_0 g_1 \dots g_{m-1}. \quad (21)$$

При этом в каждом нейронном ядре слоя m выполняется $p_m \cdot g_m$ операций умножения. Суммируя по всем нейронным слоям, получим:

$$Z = \sum_{m=0}^{\varkappa-1} p_m p_{m+1} \dots p_{\varkappa-1} g_0 g_1 \dots g_{m-1} g_m.$$

Пластичность

Высокое быстродействие БНС достигается за счет ограничения числа межнейронных связей. Расплатой за вычислительную эффективность является

снижение «уровня интеллекта» нейронной сети. Качественной характеристикой интеллекта является известное из биологии понятие пластичности, т. е. способности нейронной сети изменять свои параметры при обучении. К количественной оценке степени пластичности можно подойти следующим образом. Нейронная сеть, заданная на векторных пространствах E, D реализует отображение $E \rightarrow D$. Класс операторов нейронной сети, порождаемый изменением ее синаптических весов, образует многомерную поверхность в пространстве операторов, размерность которой определяется максимальной размерностью ее касательного пространства. Это числовая характеристика является наиболее подходящей оценкой степени параметрической пластичности нейронной сети. Нейронную сеть можно представить как материальную точку, дрейфующую в операторном пространстве при вариации синаптических весов. В этой интерпретации степень параметрической пластичности является прямым эквивалентом известного из механики понятия «число степеней свободы» материальной точки.

В касательном пространстве все операторы нейронной сети можно рассматривать как линейные, поэтому исходная нелинейная задача сводится к исследованию пластичности перестраиваемых линейных операторов. Для БНС найдено аналитическое выражение для расчета числа степеней свободы по структурным характеристикам сети [18]. Расчетная формула имеет вид:

$$S = \sum_{m=0}^{n-1} p_m g_m k_m - \sum_{m=0}^{n-2} D_m.$$

где (p_m, g_m) — размерность ядра слоя m ; k_m — число ядер в слое m ; $D_m = g_m k_m$ — число связей в переходе между слоями m и $m + 1$.

На рис. 12 приведены относительные зависимости пластичности и вычислительных затрат от размерности БНС. Графики построены для сети с одинаковой размерностью по входу и выходу. Базой для сравнения является однослойный персептрон для которого оценки равны, очевидно, $Z_0 = S_0 = N^2$.

Способность к обучению

Покажем, как степень параметрической пластичности связана с количеством линейно независимых образов, которое может распознавать нейронная сеть. Пусть H — линеаризованный оператор нейронной сети. Обозначим через X и Y , соответственно, входной и выходной векторы сети.

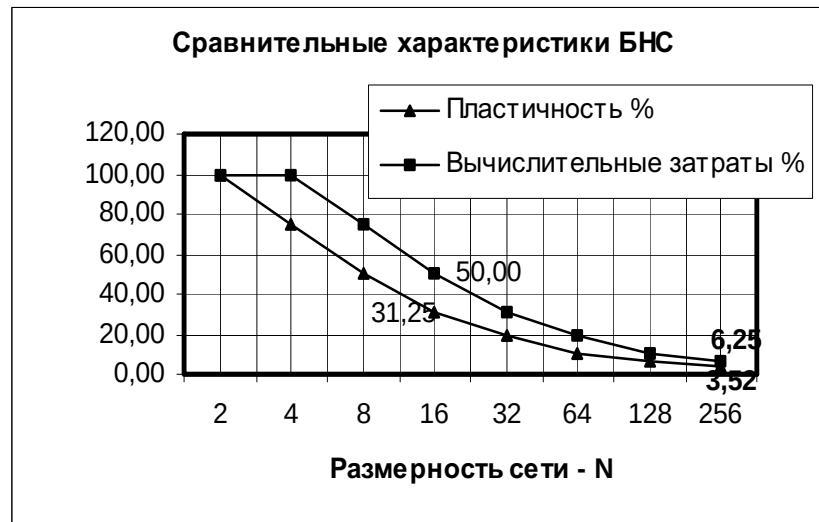


Рис. 12. Зависимость вычислительных затрат и быстродействия БНС от размерности сети

Множество обучающих примеров можно задать набором, состоящим из k пар векторов: $(X_1 Y_1), (X_2 Y_2), \dots, (X_k Y_k)$. Обозначим через $L(X)$ линейную оболочку входных векторов, а через $L(Y)$ — линейную оболочку выходных векторов обучающего множества. Линейные оболочки являются минимальными векторными пространствами, которые содержат все вектора обучающего множества, но, кроме того, они содержат несчетное множество векторов, которые могут быть получены как линейные комбинации образующих векторов.

Способность нейронной сети к обобщению позволяет расширить обучающее множество до непрерывного многообразия примеров W_{XY} , которое, очевидно, изоморфно тензорному произведению линейных оболочек: $W_{XY} \cong L(X) \otimes L(Y)$. Обозначим размерности линейных оболочек $L(X)$ и $L(Y)$, соответственно, через r_X и r_Y . Тогда размерность многообразия примеров будет равна $\dim W_{XY} = r_X r_Y$.

Необходимым условием непротиворечивости данных является выполнение неравенства: $r_Y \leq r_X$. Вариации синаптических весов нейронной сети порождают многообразие операторов W_H . Для того, чтобы нейронная



Рис. 13. Зависимость числа распознаваемых образов от размерности БНС

сеть была способна обучиться к многообразию примеров, необходимо выполнить условие $W_H \geq W_{XY}$, следствием которого является неравенство:

$$S \geq r_X r_Y, \quad (22)$$

где S — число степеней свободы нейронной сети. Если образы в обучающем множестве линейно независимы, то набор векторов Y_i образует базис линейной оболочки $L(Y)$, это означает, что $r_Y = k$. При отсутствии противоречий размерность многообразия $L(X)$ удовлетворяет условию: $r_X \geq k$. В этом случае из (22) непосредственно следует $k \leq \sqrt{S}$. На рис. 13 приведена расчетная зависимость числа распознаваемых образов от размерности сети. Например, для БНС размерности $N = 8$, показанной на рис. 4, расчетное количество независимо распознаваемых образов равно 5,66.

На рис. 14 приведены результаты эксперимента по распознаванию образов для данной нейронной сети. В эксперименте использовалась нейронная сеть с сигмоидными функциями активации в первом и во втором слое и линейными функциями активации в последнем, третьем слое. В качестве обучающего множества использовались ортогональные функции Уолша по

входу и унитарный единичный код по выходу. Сеть обучалась с помощью алгоритма Error BackPropagation.



Рис. 14. Зависимость ошибки обучения БНС (N=8) от размера обучающей выборки

Перестраиваемые спектральные преобразования

После разработки метода БПФ интенсивность исследований в области цифровой обработки резко возросла. Развитие исследований, во-первых, привело к расширению класса используемых спектральных преобразований, а во-вторых, к разработке эффективных методов их реализации. В практике цифрового спектрального анализа вошли преобразования Уолша–Адамара, Виленкина–Крестенсона, Хаара, косинусное, наклонное, вейвлет и другие. Каждый вид преобразования обладает теми или иными свойствами, определяющими область его применения. Довольно быстро выяснилось, что большинство используемых спектральных преобразований имеет один и тот же принцип формирования быстрого алгоритма. Этот принцип получил название «факторизация матриц» и заключался в разложении матрицы

дискретного спектрального преобразования H в произведение слабозаполненных матриц:

$$H = H_0 H_1, \dots, H_{\kappa-1}.$$

Структура слабозаполненных матриц устанавливалась подходящей теоремой факторизации. Было доказано множество теорем факторизации (это «творчество» продолжается и в наши дни), но постепенно стало понятным, что по одной и той же схеме факторизации можно реализовать быстрые алгоритмы для различных спектральных преобразований, путем изменения только коэффициентов в факторизованном представлении. В итоге, начиная с работ Г. Эндрюса (1970 г.), получило развитие направление обобщенного спектрального анализа, предметом исследования которого стали перестраиваемые спектральные преобразования. При реализации перестраиваемых преобразований выбирается одна из известных регулярных топологий (Гуда или Кули–Тьюки) и под данную топологию определяются значения коэффициентов слабозаполненных матриц. Ограничением развития этого направления вскоре стала сама топологическая схема, поскольку закрепленная топология, ограничивала возможности синтеза спектральных преобразований. На этом этапе перестраиваемые преобразования пересеклись с нейронными сетями, в тех и других присутствует массив настраиваемых коэффициентов и те и другие имеют послойную структурную организацию, но в отличие от нейронных сетей спектральные перестраиваемые преобразования являются линейными и ортогональными. Аналогом нейронного ядра в перестраиваемых преобразованиях является спектральное ядро, которое определяется квадратной ортогональной матрицей небольшой размерности. Функции активации нейронов ядра линейны, а смещения отсутствуют. В предыдущих параграфах были рассмотрены структурные модели и топологии БНС, полученные результаты в полной мере переносятся на перестраиваемые спектральные преобразования. Развитые методы топологического проектирования БНС позволяют существенно расширить класс перестраиваемых преобразований.

Настройка перестраиваемых преобразований

Перестраиваемое преобразование можно рассматривать как быструю нейронную сеть, где в каждом нейронном ядре i^λ выполнится линейное пре-

образование:

$$y_{i\lambda}^{\lambda}(v_{\lambda}) = \sum_{u_{\lambda}} x_{i\lambda}^{\lambda}(u_{\lambda}) w_{i\lambda}^{\lambda}(u_{\lambda}, v_{\lambda}). \quad (23)$$

Метод настройки основан на аналитическом представлении сигнальных передач рецептор–аксон между терминальными слоями сети. Функция передач определяется частной производной

$$h(U, V) = \frac{\partial y_{i_{\kappa-1}}^{\kappa-1}(v_{\kappa-1})}{\partial x_{i_0}^0(u_0)},$$

где U, V — глобальные номера рецепторов и аксонов терминальных слоев. В терминологии спектральных преобразований U — это индекс координаты вектора входного сигнала, а V — номер спектрального коэффициента. Элементы $h(U, V)$ образуют квадратную матрицу H спектрального преобразования $Y = XH$. Поскольку БНС — слабосвязанная сеть, то для каждой пары нейронных ядер, одно из которых принадлежит входному слою, а второе — выходному, существует единственный связывающий их путь. Последовательно дифференцируя вдоль выбранного пути цепочку уравнений (23), получим:

$$h(U, V) = w_{i_0}^0(u_0, v_0) w_{i_1}^1(u_1, v_1) \dots w_{i_{\kappa-1}}^{\kappa-1}(u_{\kappa-1}, v_{\kappa-1}). \quad (24)$$

Цель настройки перестраиваемого преобразования состоит в том, чтобы по известной матрице H найти элементы ядер и определить топологию слабозаполненных матриц H_{λ} . Топология данных матриц непосредственно связана с внешним представлением топологической траектории. Поэтому задача топологического синтеза интерпретируется как реализация заданной траектории. Процедура синтеза топологии, как правило, допускает множество возможных решений. Дополнительными требованиями могут быть компактность и регулярность топологической траектории, что упрощает реализацию алгоритма быстрого преобразования.

Настройка на базис Уолша

Функции базиса Уолша в упорядочении Пэли [9] задаются на интервале длиной $N = 2^{\kappa}$ следующим выражением :

$$\text{pal}(U, V) = \prod_{\lambda=0}^{\kappa-1} (-1)^{U_{\lambda} V_{\lambda}}, \quad (25)$$

где

$$U = \langle U_{\kappa-1} U_{\kappa-2} \dots U_0 \rangle \quad \text{и} \quad V = \langle V_{\kappa-1} V_{\kappa-2} \dots V_0 \rangle. \quad (26)$$

Все разрядные переменные в данных формулах принимают значения $\{0, 1\}$.

Построим матрицу спектрального преобразования так, чтобы функции базиса располагались вдоль столбцов. В процессе настройки необходимо определить, топологическую траекторию быстрого преобразования и параметры нейронных ядер. Дополнительно потребуем, чтобы топологическая траектория была компактной и регулярной.

Параметры нейронных ядер. Сравнивая (24) с определением функций Уолша, получим

$$w_{i\lambda}^\lambda(u_\lambda, v_\lambda) = (-1)^{u_\lambda v_\lambda}, \quad (27)$$

$$u_\lambda = U_\lambda, \quad v_\lambda = V_\lambda. \quad (28)$$

Выражению (27) соответствует матрица ядра:

$$W = \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}.$$

Очевидно, что нейронные ядра одинаковы во всех слоях сети.

Траектория топологий. Граничные условия для компактной траектории устанавливаются выражениями (26), (28) и имеют вид:

$$\begin{aligned} U &= \langle U_{\kappa-1} U_{\kappa-2} \dots U_0 \rangle = \langle u_{\kappa-1} u_{\kappa-2} \dots u_0 \rangle, \\ V &= \langle V_{\kappa-1} V_{\kappa-2} \dots V_0 \rangle = \langle v_{\kappa-1} v_{\kappa-2} \dots v_0 \rangle. \end{aligned}$$

На промежуточные шаги траектории, каких либо требований не накладывается, поэтому траекторию топологий можно выбрать с большим произволом. Воспользуемся регулярной порождающей схемой Кули-Тьюки «с прореживанием по времени»:

$$\begin{aligned} U^\lambda &= \langle u_{\kappa-1} u_{\kappa-2} \dots u_{\lambda+1} u_\lambda v_{\lambda-1} v_{\lambda-2} \dots v_0 \rangle, \\ V^\lambda &= \langle u_{\kappa-1} u_{\kappa-2} \dots u_{\lambda+1} v_\lambda v_{\lambda-1} v_{\lambda-2} \dots v_0 \rangle. \end{aligned}$$

Для данной схемы номер ядра определяется выражением

$$i^\lambda = \langle u_{\kappa-1} u_{\kappa-2} \dots u_{\lambda+1} v_{\lambda-1} v_{\lambda-2} \dots v_0 \rangle.$$

Структуру топологических матриц можно определить из табличного соответствия (см. таблицу 2). На основании данной таблицы получим, что аналитическая форма слабозаполненных матриц будет иметь вид

ТАБЛИЦА 2. Соответствие локальных и глобальных переменных в топологии Кули-Тьюки «с прореживанием по времени»

$U^\lambda =$	$u_{\kappa-1}$	$u_{\kappa-2}$	$\dots$	$u_{\lambda+1}$	u_λ	$v_{\lambda-1}$	$\dots$	v_2	v_1	v_0
$U^\lambda =$	$U_{\kappa-1}^\lambda$	$U_{\kappa-2}^\lambda$	$\dots$	$U_{\lambda+1}^\lambda$	U_λ^λ	$U_{\lambda-1}^\lambda$	$\dots$	U_2^λ	U_1^λ	U_0^λ
$V^\lambda =$	$u_{\kappa-1}$	$u_{\kappa-2}$	$\dots$	$u_{\lambda+1}$	v_λ	$v_{\lambda-1}$	$\dots$	v_2	v_1	v_0
$V^\lambda =$	$V_{\kappa-1}^\lambda$	$V_{\kappa-2}^\lambda$	$\dots$	$V_{\lambda+1}^\lambda$	V_λ^λ	$V_{\lambda-1}^\lambda$	$\dots$	V_2^λ	V_1^λ	V_0^λ

$$h_\lambda(U^\lambda, V^\lambda) = w_{i^\lambda}^\lambda(U_\lambda^\lambda, V_\lambda^\lambda) \delta(U_{\lambda-1}^\lambda, V_{\lambda-1}^\lambda) \dots \dots \delta(U_{\lambda+1}^\lambda, V_{\lambda+1}^\lambda) \delta(U_{\lambda-1}^\lambda, V_{\lambda-1}^\lambda) \dots \delta(U_0^\lambda, V_0^\lambda), \quad (29)$$

где $i^\lambda = \langle V_{\lambda-1}^\lambda V_{\lambda-2}^\lambda \dots V_{\lambda+1}^\lambda V_{\lambda-1}^\lambda \dots V_0^\lambda \rangle$. На рис. 15 показано матричное представление алгоритма быстрого преобразования Уолша для размерности $N = 2^3$. Матрица преобразования факторизуется в произведение трех матриц:

$$H = H_0 H_1 H_2.$$

Поразрядное представление строк и столбцов матрицы H можно записать в виде $U = \langle U_2 U_1 U_0 \rangle$, $V = \langle V_2 V_1 V_0 \rangle$. На основании (29) имеем, что слабозаполненные матрицы по слоям определяются выражениями:

- Слой 0:

$$h_0(U^0, V^0) = w_{i^0}^0(U_0^0, V_0^0) \delta(U_2^0, V_2^0) \delta(U_1^0, V_1^0), \quad i^0 = \langle V_2^0 V_1^0 \rangle.$$

- Слой 1:

$$h_1(U^1, V^1) = w_{i^1}^1(U_1^1, V_1^1) \delta(U_2^1, V_2^1) \delta(U_0^1, V_0^1), \quad i^1 = \langle V_2^1 V_0^1 \rangle.$$

- Слой 2:

$$h_2(U^2, V^2) = w_{i^2}^2(U_2^2, V_2^2) \delta(U_1^2, V_1^2) \delta(U_0^2, V_0^2), \quad i^2 = \langle V_1^2 V_0^2 \rangle$$

	V_2	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1	1
	V_1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1	1
	V_0	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	1
U_2	U_1	U_0																								
0	0	0	1	1						1		1						1			1					
0	0	1	1	-1						1		1						1			1					
0	1	0		1	1					1		-1							1				1			
0	1	1		1	-1					1		-1							1					1		
1	0	0				1	1						1	1				1			-1					
1	0	1				1	-1						1		1			1			-1					
1	1	0					1	1					1	-1				1				-1				
1	1	1					1	-1					1	-1					1				-1			

Рис. 15. Матричное представление быстрого алгоритма для преобразования Уолша–Пэли

Перемножив матрицы факторизованного представления, получим следующую матрицу базисных функций:

$$\begin{pmatrix} 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ 1 & -1 & 1 & -1 & 1 & -1 & 1 & -1 \\ 1 & 1 & -1 & -1 & 1 & 1 & -1 & -1 \\ 1 & -1 & -1 & 1 & 1 & -1 & -1 & 1 \\ 1 & 1 & 1 & 1 & -1 & -1 & -1 & -1 \\ 1 & -1 & 1 & -1 & -1 & 1 & -1 & 1 \\ 1 & 1 & -1 & -1 & -1 & -1 & 1 & 1 \\ 1 & -1 & -1 & 1 & -1 & 1 & 1 & -1 \end{pmatrix}$$

Эта же матрица может быть получена из определения (25).

Настройка на базис Фурье

По определению функции базиса Фурье задаются выражением:

$$F_N(U, V) = \frac{1}{\sqrt{N}} \exp \left(-j \frac{2\pi}{N} UV \right),$$

где U — временной отсчет, V — частота (или номер) базисной функции, $j = \sqrt{-1}$, $N = p_0 p_1 \dots p_{\kappa-1}$ — размерность преобразования (для быстрых преобразований размерность всегда является составным числом). Использование методов топологической настройки БНС позволило реализовать быстрое преобразование Фурье с естественным упорядочением по частотам следования (факторизованное представление для $N = 2^3$ представлено на рис.16, множители $1/\sqrt{2}$ с целью упрощения не показаны, поворачивающий множитель определяется выражением $\omega = \exp(-j\frac{2\pi}{8})$). Детали реализации БПФ приведены в [10]).

V_0	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1																			
V_1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1																			
V_2	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1																			
U_2	U_1	U_0																																									
0	0	0	1	1																			1	1																			
0	0	1	1		1																			1	1																		
0	1	0																			1	1																					
0	1	1																			1	1																					
1	0	0	1	-1																			1	-1																			
1	0	1	1		-1																			ω^2	$-\omega^2$																		
1	1	0																			1	-1																					
1	1	1																			ω	$-\omega$																					

Рис. 16. Факторизованное представление алгоритма БПФ с естественным упорядочением по частотам следования

Быстрое вейвлет-преобразование

Пакетный вейвлет-базис на интервале длиной $N = p_0 p_1 \dots p_{\kappa-1}$ может быть задан выражением:

$$h(U, V) = \varphi_m(U_{\kappa-m}, V_{\kappa-m}) \delta(\tau_m, \langle U_{\kappa-1} U_{\kappa-2} \dots U_{\kappa-m+1} \rangle), \quad (30)$$

где $U = \langle U_{\kappa-1} U_{\kappa-2} \dots U_0 \rangle$ — номер временного отсчета, $V = \langle V_0 V_1 \dots V_{\kappa-1} \rangle$ — номер вейвлет-функции (предполагается, что вейвлет-функции распо-

жены вдоль столбцов базисной матрицы), $\varphi_m(U_{\kappa-m}, V_{\kappa-m})$ — набор ортогональных образующих импульсов в частотной локализации номера m , $\delta(\cdot)$ — функция Кронекера, τ_m — порядковый номер временной локализации вейвлет-функции в частотной локализации m . Множество функций вейвлет-базиса, принадлежащих одной частотной локализации в дальнейшем будем называть *поликадой*. Функции в пределах поликады обладают одними и теми же частотными свойствами, но отличаются позициями образующих импульсов. Упорядочим функции в каждой поликаде по временным локализациям следующим правилом:

$$\tau_m = \langle V_{\kappa-1} V_{\kappa-2} \dots V_{\kappa-m+1} \rangle. \quad (31)$$

При таком упорядочении функцию Кронекера в выражении (30) можно разложить в произведение функций Кронекера, в результате получим:

$$h(U, V) = \varphi_m(U_{\kappa-m}, V_{\kappa-m}) \delta(U_{\kappa-m+1}, V_{\kappa-m+1}) \dots \dots \delta(U_{\kappa-2}, V_{\kappa-2}) \delta(U_{\kappa-1}, V_{\kappa-1}).$$

Сопоставив данное выражение с формулой вычисления элементов матрицы перестраиваемого преобразования (24), нетрудно найти правило определения элементов ядер (детали реализации алгоритма можно найти в работе [10]).

Одним из простейших вейвлет-преобразований является базис Хаара. Базис строится на интервале длиной $N = 2^\kappa$ и порождается двуполярным импульсом с временной базой, равной двум. В случае $N = 2^3$, базисные функции по частотным локализациям разбиваются на три октавы, Матрица образующих импульсов каждой октавы имеет вид:

$$\begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix},$$

а матрица преобразования факторизуется в произведение трех слабозаполненных матриц (см. рис. 17).

V_0	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1
V_1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1
V_2	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1
U_2	U_1	U_0																						
0	0	0	1	1					1	1							1	1						
0	0	1	1	-1						1	0							1	0					
0	1	0		1	1				1	-1										1	0			
0	1	1		1	-1					0	1										1	0		
1	0	0				1	1				1	1					1	-1						
1	0	1				1	-1					1	0					0	1					
1	1	0					1	1			1	-1							0	1				
1	1	1					1	-1				0	1							0	1			

Рис. 17. Факторизованное представление быстрого преобразования Хаара

Свернув факторизованное произведение, получим матрицу базиса в виде:

$$H = \begin{bmatrix} 1 & 1 & 1 & & 1 & & & \\ 1 & 1 & 1 & & -1 & & & \\ 1 & 1 & -1 & & & & 1 & \\ 1 & 1 & -1 & & & & -1 & \\ 1 & -1 & & 1 & & 1 & & \\ 1 & -1 & & 1 & & -1 & & \\ 1 & -1 & & -1 & & & & 1 \\ 1 & -1 & & -1 & & & & -1 \end{bmatrix}$$

Нейросетевая аппроксимация регулярных фракталов

Термин «фрактал» был введен *Бенуа Мандельбротом* в 1975 году для обозначения особого класса многомерных функций, обладающих свойством самоподобия и дробной размерности. Первоначально к фракталам относились как к математической экзотике, но за прошедшие десятилетия ситуация кардинальным образом изменилась. Оказалось, что окружающий нас

мир наполнен фракталами гораздо в большей степени, чем «привычными» нам объектами. Показательным примером являются, казалось бы, хорошо изученные динамические системы. Современные исследования показали, что в пространстве фазовых координат для нелинейных динамических систем характерно существование фрактальных притягивающих множеств (аттракторов) [19]. Фрактальные аттракторы связаны с так называемыми сценариями динамического хаоса [20], определяющего границы возможного прогнозирования поведения динамических систем.

В ряде работ [21, 22] было показано, что многослойные нейронные сети могут служить генераторами фрактальных структур и использоваться в этом качестве как инструмент моделирования и анализа нелинейных явлений. Вид фрактала накладывает определенные ограничения на структуру, топологию и параметры нейронной сети. В данном разделе будет показано использование БНС для генерации фрактальных структур.

Аналитическая форма регулярного фрактала

На рис. 18 показаны три итерации фрактала Кантора. На каждой итерации из непрерывной области удаляется средняя треть, что эквивалентно масштабной деформации исходной функции в отношении 1:3.

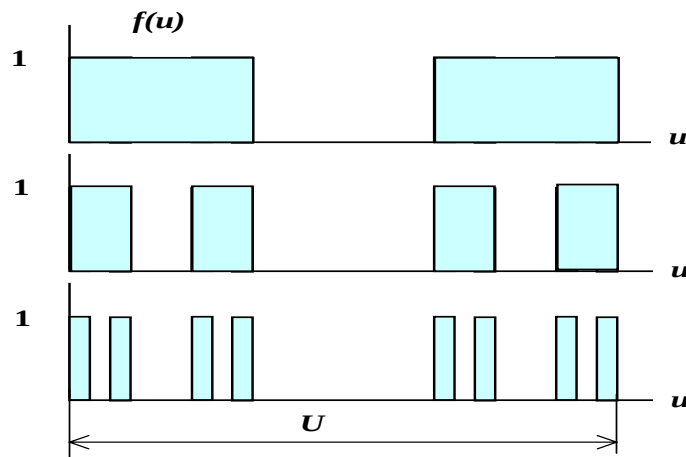


Рис. 18. Одномерный фрактал Кантора

Построим аналитическую форму вначале для фрактала Кантора, а затем обобщим ее на более широкий класс регулярных фракталов. Будем полагать, что фрактал Кантора задан на непрерывном интервале $U = [0, 1]$. Любую точку данного интервала $u \in U$ можно записать в троичной системе счисления в виде дроби:

$$u = 0, u_1 u_2 \dots u_n \dots, \quad (32)$$

где $u_i \in \{0, 1, 2\}$ — целые разрядные числа. Введем также непрерывные переменные $\tilde{u}_i \in [0, 3]$. Очевидно, что последовательность разрядных чисел (32) можно в любом месте оборвать, завершив ее непрерывной переменной, в результате точка интервала будет представлена в виде:

$$u = 0, u_1 u_2 \dots u_{n-1} \tilde{u}_n.$$

Разрядные числа u_i будем рассматривать как результат действия оператора выделяющего целую часть из непрерывной переменной $\tilde{u}_i$. Операцию масштабирования можно рассматривать как изменение правила представления точек интервала. Например, для базовой функции Кантора аргумент определяется правилом $u = \langle 0, \tilde{u}_1 \rangle = 3^{-1} \tilde{u}_1$. Для функции второй итерации: $u = \langle 0, u_1 \tilde{u}_2 \rangle = 3^{-1} u_1 + 3^{-2} \tilde{u}_2$. Для функции третьей итерации: $u = \langle 0, u_1 u_2 \tilde{u}_3 \rangle = 3^{-1} u_1 + 3^{-2} u_2 + 3^{-3} \tilde{u}_3$, и т. д. На непрерывном интервале $[0, 3]$ определим функцию:

$$\varphi(\tilde{u}) = \begin{cases} 1 & \text{для } \tilde{u} \in [0, 1), \\ 0 & \text{для } \tilde{u} \in [1, 2), \\ 1 & \text{для } \tilde{u} \in [2, 3) \end{cases}$$

Нетрудно видеть, что фрактал Кантора в аналитическом виде можно записать в виде бесконечного произведения:

$$f(u) = \varphi(\tilde{u}_1) \varphi(\tilde{u}_2) \varphi(\tilde{u}_3) \dots \varphi(\tilde{u}_n) \dots$$

Дискретная аппроксимация фракталов

В дискретном варианте фрактал строится на последовательности расширяющихся дискретных интервалах:

$$U_i = \{0, 1, 2, \dots, N_i - 1\}, \quad i = 0, 1, 2, \dots, \infty - 1 \dots$$

Точки интервалов определяются правилом $u = \langle u_{\kappa-1} u_{\kappa-2} \dots u_0 \rangle$. В общем случае разрядные переменные имеют различные локальные области определения: $u_i \in \{0, 1, \dots, p_i - 1\}$. На локальных интервалах определены кусочно-постоянные дискретные функции $\varphi_i(u_i)$. Аппроксимация фрактала записывается в виде конечного произведения:

$$f(u) \approx \varphi(u_0) \varphi(u_1) \dots \varphi(u_{\kappa-1}).$$

Например, для фрактала Кантора все дискретные образующие функции одинаковы и имеют вид:

$$\varphi(u) = \begin{cases} 1, & \text{для } u = 0, \\ 0, & \text{для } u = 1, \\ 1, & \text{для } u = 2. \end{cases}$$

Ковер Кантора. Двумерным вариантом фрактала Кантора являются плоская фигура, заданная на непрерывных интервалах $U = [0, 1)$, $V = [0, 1)$. На рис. 19 показана одна из итераций фрактала. Координаты точки в плоскости фигуры определяются бесконечными дробями:

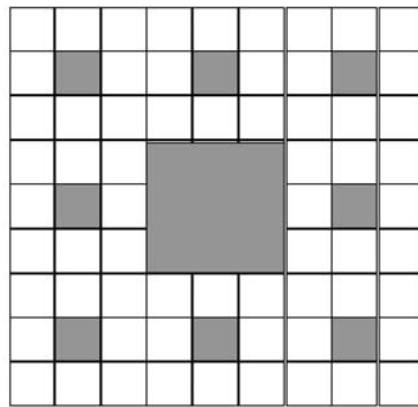


Рис. 19. Двумерный фрактал Кантора

$$u = 0, u_1 u_2 \dots u_n \dots; u_i \in \{0, 1, 2\}, \\ v = 0, v_1 v_2 \dots v_n \dots; v_i \in \{0, 1, 2\}.$$

Обозначим через $\tilde{u}_i \in [0, 3)$ и $\tilde{v}_i \in [0, 3)$ соответствующие непрерывные переменные, тогда ковер Кантора в аналитической форме можно записать в виде бесконечного произведения:

$$f(u, v) = \varphi(\tilde{u}_1, \tilde{v}_1) \varphi(\tilde{u}_2, \tilde{v}_2) \varphi(\tilde{u}_3, \tilde{v}_3) \dots \varphi(\tilde{u}_n, \tilde{v}_n) \dots$$

Дискретная аппроксимация для ковра Кантора представляет собой конечное произведение:

$$f(u, v) \approx \varphi(u_0, v_0) \varphi(u_1, v_1) \varphi(u_2, v_2) \dots \varphi(u_{\varkappa-1}, v_{\varkappa-1}), \quad (33)$$

где все дискретные образующие функции имеют вид:

$$\varphi(u_i, v_i) = \begin{pmatrix} 1 & 1 & 1 \\ 1 & 0 & 1 \\ 1 & 1 & 1 \end{pmatrix}.$$

Нетрудно заметить, что полученное выражение по форме совпадает с представлением (24) для элементов матрицы H нейронной сети. Сравнивая (24) и (33), получим следующие правила определения параметров нейронных ядер:

$$w_{i\lambda}^\lambda(u_\lambda, v_\lambda) = \varphi(u_\lambda, v_\lambda).$$

На рис. 20 показана нейросетевая аппроксимация для ковра Кантора на интервале длиной $N = 3^2$. Аппроксимирующая нейронная сеть имеет два слоя, каждой матрице сомножителю отвечает один нейронный слой. Все непоказанные элементы матриц равны нулю. Из рисунка видно, что результирующая матрица подобна ковра Кантора.

Фрактал Серпинского. На рис. 21 показана одна из итераций двумерного фрактала, называемого «салфеткой Серпинского». Фрактал задан на непрерывных интервалах $U = [0, 1)$, $V = [0, 1)$. Дискретная аппроксимация фрактала Серпинского определяется на интервале длиной $N = 2^\varkappa$ выражением (33), где образующая функция имеет вид

$$\varphi(u_i, v_i) = \begin{pmatrix} 1 & 1 \\ 1 & 0 \end{pmatrix}.$$

На рис. 22 показана нейросетевая аппроксимация для салфетки Серпинского на интервалах длиной $N = 2^3$. Аппроксимирующая нейронная сеть имеет три слоя, каждой матрице сомножителю отвечает один нейронный слой. Из рисунка видно, что результирующая матрица подобна непрерывной форме данного фрактала.

$$\begin{array}{|c|c|c|} \hline 1 & 1 & 1 \\ \hline 1 & 0 & 1 \\ \hline 1 & 1 & 1 \\ \hline \end{array} \cdot \begin{array}{|c|c|c|} \hline 1 & & 1 \\ \hline 1 & & 1 \\ \hline & 1 & 1 \\ \hline \end{array} = \begin{array}{|c|c|c|c|c|c|c|c|} \hline 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ \hline 1 & 0 & 1 & 1 & 0 & 1 & 1 & 0 \\ \hline 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ \hline 1 & 1 & 1 & 0 & 0 & 0 & 1 & 1 \\ \hline 1 & 0 & 1 & 0 & 0 & 0 & 1 & 0 \\ \hline 1 & 1 & 1 & 0 & 0 & 0 & 1 & 1 \\ \hline 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ \hline 1 & 0 & 1 & 1 & 0 & 1 & 1 & 0 \\ \hline 1 & 1 & 1 & 1 & 1 & 1 & 1 & 1 \\ \hline \end{array}$$

Рис. 20. Нейросетевая аппроксимация двумерного фрактала Кантора

Фрактальная фильтрация сигналов

Хорошо известны трудности связанные с выделением полезной информации из сигналов, обладающих частотно-локальными свойствами. Характерным примером являются сигналы, отражающие динамику турбулентных потоков, порожденных нелинейно взаимодействующими процессами в широких диапазонах пространственных частот и временных локализаций. Для анализа подобных сигналов в настоящее время широко используется вейвлет-преобразование, поскольку его элементы хорошо локализованы и обладают подвижным частотно-временным окном [23]. В основе вейвлет-анализа лежит представление сигналов в виде суперпозиции масштабных преобразований и сдвигов образующих импульсов. Самоподобие вейвлет-функций при масштабных преобразованиях является отличительным свойством фрактальных последовательностей, поэтому вейвлет-преобразование можно считать одним из представителей фрактальных методов обработки данных.

Принцип самоподобия можно также использовать при построении фрактальных фильтров. В обобщенном понимании фильтрация это процедура, реализующая проекцию сигнала, заданного в функциональном пространстве, в пространство меньшей размерности. В отличие от вейвлет-преобразования фрактальная фильтрация необратима и связана с потерей информации. Цель фрактальной фильтрации заключается в том, чтобы

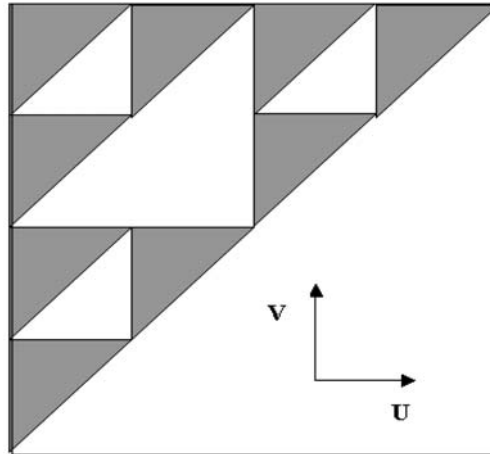


Рис. 21. Двумерный фрактал Серпинского

выявить пространственно распределенные свойства изучаемого объекта в кратных временных масштабах.

Фильтрация дискретных сигналов

Рассмотрим сигнал, определенный функцией $f(u)$, заданной на дискретном интервале длиной $N = p_0 p_1 \dots p_{\kappa-1}$. Представим аргумент этой функции в позиционной многоосновной системе счисления с основаниями $p_0, p_1, \dots, p_{\kappa-1}$:

$$u = \langle u_{\kappa-1} u_{\kappa-2} \dots u_0 \rangle = u_{\kappa-1} p_{\kappa-2} p_{\kappa-3} \dots p_0 + \\ + u_{\kappa-2} p_{\kappa-3} p_{\kappa-4} \dots p_0 + \dots + u_1 p_0 + u_0,$$

где $u_i \in [0, 1, \dots, p_i - 1]$ — разрядные переменные. В результате данного преобразования сигнал представляется как многомерная функция вида $f \langle u_{\kappa-1} u_{\kappa-2} \dots u_0 \rangle$, где каждый аргумент функции определяет некоторый масштабный срез сигнала. Зафиксируем все аргументы функции кроме u_m . Варьируя свободный аргумент u_m , получим выборку S_m (с числом элементов p_m). *Фрактальным фильтром* частотной локализации m будем называть произвольный функционал $F(S_m)$, определенный на выборке S_m .

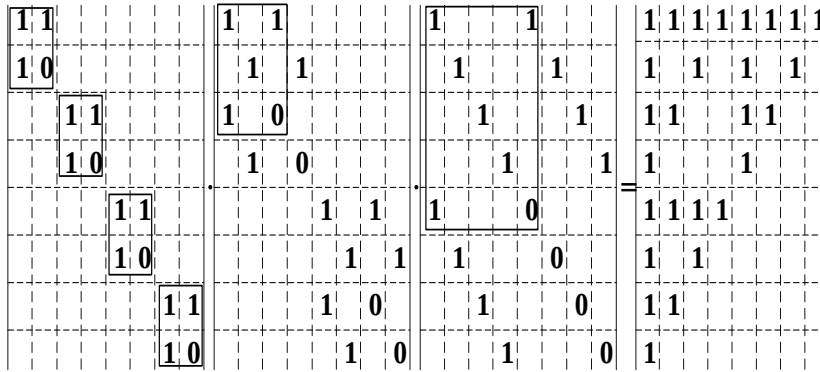


Рис. 22. Нейросетевая аппроксимация двумерного фрактала Серпинского

Операцию фрактальной фильтрации можно записать в виде:

$$f_{out} \langle u_{\kappa-1} u_{\kappa-2} \dots u_{m+1} u_{m-1} \dots u_0 \rangle = F_{u_m} (f_{inp} \langle u_{\kappa-1} u_{\kappa-2} \dots u_0 \rangle).$$

Если $m = 0$, то фильтр будем называть *фильтром нижних фрактальностей*, такой фильтр генерализует сигнал, сглаживая мелкие детали. Если $m = \kappa - 1$, то фильтр назовем *фильтром верхних фрактальностей*, такой фильтр сохраняет детали и нивелирует трендовые изменения. Фильтр с $0 < m < \kappa - 1$ можно назвать *препарирующим*. Фрактальный фильтр выполняет проекцию функционального пространства размерности N в подпространство размерности N/p_m , однако в ряде случаев результат фильтрации удобно представлять на исходном интервале $U = [0, 1, \dots, N - 1]$. Это можно сделать, положив равными значения функции в точках интервала, различающихся только по значению аргумента u_m . Полученная функция может быть вновь подвержена фрактальной фильтрации в некоторой частотной локальности $n \neq m$. Последовательная фрактальная фильтрация, реализуется принцип иерархической многомасштабной обработки данных.

Типы фрактальных фильтров

Линейные фильтры — могут быть реализованы произвольными линейными функционами. Примерами могут служить следующие типы:

- *Фильтр средних:*

$$\begin{aligned} f_{out} \langle u_{\mathcal{K}-1} u_{\mathcal{K}-2} \dots u_{m+1} u_{m-1} \dots u_0 \rangle &= \\ &= \frac{1}{p_m} \sum_{u_m} f_{inp} \langle u_{\mathcal{K}-1} u_{\mathcal{K}-2} \dots u_0 \rangle; \end{aligned}$$

- *Унарный линейный фильтр:*

$$\begin{aligned} f_{out} \langle u_{\mathcal{K}-1} u_{\mathcal{K}-2} \dots u_{m+1} u_{m-1} \dots u_0 \rangle &= \\ &= f_{inp} \langle u_{\mathcal{K}-1} u_{\mathcal{K}-2} \dots u_{m+1} t_m u_{m-1} \dots u_0 \rangle, \quad t_m = \text{const}; \end{aligned}$$

- *Аффинный фильтр:*

$$\begin{aligned} f_{out} \langle u_{\mathcal{K}-1} u_{\mathcal{K}-2} \dots u_{m+1} u_{m-1} \dots u_0 \rangle &= \\ &= \sum_{u_m} f_{inp} \langle u_{\mathcal{K}-1} u_{\mathcal{K}-2} \dots u_0 \rangle w(u_m) + w_0. \end{aligned}$$

Аффинный фильтр принято условно относить к классу линейных.

Нелинейные фильтры. Нелинейный функционал можно задать в алгоритмическом или аналитическом виде. Примерами алгоритмически заданных функционалов являются ранговые фильтры. При ранговой фильтрации выборка ранжируется. Поскольку длина выборки равна p_m то число возможных рангов не превышает p_m . Значением функционала является выборочное значение, соответствующее рангу $r \leq p_m$. Конкретными примерами могут служить следующие типы фильтров:

$$\begin{aligned} f_{out} \langle u_{\mathcal{K}-1} u_{\mathcal{K}-2} \dots u_{m+1} u_{m-1} \dots u_0 \rangle &= \min_{u_m} f_{inp} \langle u_{\mathcal{K}-1} u_{\mathcal{K}-2} \dots u_0 \rangle, \\ f_{out} \langle u_{\mathcal{K}-1} u_{\mathcal{K}-2} \dots u_{m+1} u_{m-1} \dots u_0 \rangle &= \max_{u_m} f_{inp} \langle u_{\mathcal{K}-1} u_{\mathcal{K}-2} \dots u_0 \rangle. \end{aligned}$$

Первый из них соответствует значению ранга $r = 0$, а второй — значению $r = p_m$. Примером нелинейных фильтров с аналитически заданными функционалами являются нормирующие фильтры, для которых функционал определяется той или иной векторной нормой, конкретными примерами могут служить следующие типы:

- *Фильтр Евклида:*

$$f_{out} \langle u_{\kappa-1} u_{\kappa-2} \dots u_{m+1} u_{m-1} \dots u_0 \rangle = \sqrt{\sum_{u_m} f_{inp} \langle u_{\kappa-1} u_{\kappa-2} \dots u_0 \rangle^2};$$

- *Фильтр Хэмминга:*

$$f_{out} \langle u_{\kappa-1} u_{\kappa-2} \dots u_{m+1} u_{m-1} \dots u_0 \rangle = \sum_{u_m} |f_{inp} \langle u_{\kappa-1} u_{\kappa-2} \dots u_0 \rangle|.$$

Многомерные фрактальные фильтры. Многомерные фильтры реализуются функционалами, заданными на многомерной выборке. Например, двумерный фрактальный фильтр в общем случае может быть определен в виде:

$$f_{out} \langle u_{\kappa-1} u_{\kappa-2} \dots u_0 \rangle = F_{u_m, u_n} (f_{inp} \langle u_{\kappa-1} u_{\kappa-2} \dots u_0 \rangle).$$

Такой фильтр выполняет проекцию сигнала из пространства размерности N в пространство размерности $N/(p_m p_n)$.

Фильтрация непрерывных сигналов

Пусть сигнал $f(u)$ определен на непрерывном интервале $U = [0, 1)$. Зададим позиционную систему счисления множеством оснований $p_1, p_2, p_3, \dots$, и представим аргумент функции в виде бесконечной дроби:

$$u = \langle 0, u_1 u_2 u_3 \dots \rangle = u_1 p_1^{-1} + u_2 p_2^{-1} p_1^{-1} + u_3 p_3^{-1} p_2^{-1} p_1^{-1} + \dots$$

Непрерывный фрактальный фильтр частотной локализации m определяется функционалом:

$$f_{out} \langle 0, u_1 u_2 \dots u_{m-1} u_{m+1} \dots \rangle = F_{u_m} (f_{inp} \langle 0, u_1 u_2 \dots u_m \dots \rangle).$$

В частности, фильтр верхних фрактальностей будет иметь вид:

$$f_{out} \langle 0, u_2 u_3 u_4 \dots \rangle = F_{u_1} (f_{inp} \langle 0, u_1 u_2 u_3 \dots \rangle),$$

а многомерный фильтр нижних фрактальностей можно задать функционалом:

$$f_{out} \langle 0, u_1 u_2 u_3 u_4 \dots u_m \rangle = \underset{u_{m+1}, u_{m+2}, \dots}{F} (f_{inp} \langle 0, u_1 u_2 u_3 \dots \rangle).$$

Для непрерывных фильтров выходной сигнал автоматически масштабируется в исходный интервал $U = [0, 1)$.

Приспособленные быстрые преобразования

В данном разделе будет рассмотрено использование фрактальной фильтрации для настройки множественных адаптивных фильтров реализованных на базе БНС. Наличие быстрого алгоритма позволяет использовать данные реализации при обработке сигналов высокой размерности в системах реального времени.

Методы цифровой адаптивной фильтрации (такие как согласованная фильтрация, инверсная фильтрация, фильтр Винера и т. п. [24]) используют априорную информацию о сигнале или помехе для настройки коэффициентов фильтра. Подобные фильтры позволяют получить максимальное отношение «сигнал/шум» при фильтрации опорных эталонных сигналов. Поэтому адаптивные фильтры находят широкое применение в технических приложениях.

Большой класс адаптивных фильтров представляют собой линейные преобразования среди которых можно выделить подкласс ортогональных преобразований. Перестраиваемые ортогональные преобразования, адаптируемые к одиночной функции получили название приспособленных преобразований [5]. Это название в дальнейшем будет использовано по отношению к любым линейным перестраиваемым преобразованиям, настраиваемых по эталонным функциям.

Алгоритм приспособления

Линейное преобразование будем называть приспособленным к функции, если один из столбцов матрицы преобразования совпадает с заданной функцией. Подобным образом можно определить приспособленность к нескольким функциям (в том случае, когда используется правое умножение матрицы на вектор, приспособление более удобно выполнять к строкам

матрицы). Идея метода настройки БНС к одной или нескольким функциям основана на представлении каждой функции заданного набора в виде фрактального произведения, отвечающем структуре БНС. Для одной функции задача фрактальной аппроксимации формулируется следующим образом.

Пусть на дискретном интервале длиной $N = p_0 p_1 \dots p_{\varkappa-1}$ задана вещественная функция $f(u)$. Ставится задача найти представление функции в виде произведения:

$$f(u) = \varphi_{i^0}(u_0) \varphi_{i^1}(u_1) \dots \varphi_{i^{\varkappa-1}}(u_{\varkappa-1}). \quad (34)$$

Правило выбора образующих функций для каждого шага фрактальной итерации подчинено условию $i^m = \langle u_{\varkappa-1} u_{\varkappa-2} \dots u_{m+1} \rangle$. Поскольку элементы матрицы БНС представляются через элементы ядер:

$$h(U, V) = w_{i^0}^0(u_0, v_0) w_{i^1}^1(u_1, v_1) \dots w_{i^{\varkappa-1}}^{\varkappa-1}(u_{\varkappa-1}, v_{\varkappa-1}), \quad (35)$$

где $i^m = \langle u_{\varkappa-1} u_{\varkappa-2} \dots u_{m+1} v_{m-1} v_{m-2} \dots v_0 \rangle$, то столбец с номером $V = \langle v_{\varkappa-1} v_{\varkappa-2} \dots v_0 \rangle$ матрицы H будет совпадать с заданной функцией, если элементы ядер совпадают с множителями фрактального разложения (34). Так как для нулевого слоя номер ядра $i^0 = \langle u_{\varkappa-1} u_{\varkappa-2} \dots u_1 \rangle$ не зависит от переменных v , то, возможно, одновременно приспособить БНС к такому количеству функций, сколько значений может принимать переменная v_0 . Для нахождения фрактального разложения (34) введем последовательность вспомогательных функций $f_0, f_1, \dots, f_{\varkappa-1}$, определив их рекуррентным правилом:

$$\begin{aligned} f(u) &= f_0 \langle u_{\varkappa-1} u_{\varkappa-2} \dots u_0 \rangle, \\ f_0 \langle u_{\varkappa-1} u_{\varkappa-2} \dots u_0 \rangle &= f_1 \langle u_{\varkappa-1} u_{\varkappa-2} \dots u_1 \rangle \varphi_{i^0}(u_0), \\ f_1 \langle u_{\varkappa-1} u_{\varkappa-2} \dots u_1 \rangle &= f_2 \langle u_{\varkappa-1} u_{\varkappa-2} \dots u_2 \rangle \varphi_{i^1}(u_1), \\ &\vdots \\ f_{\varkappa-2} \langle u_{\varkappa-1} u_{\varkappa-2} \rangle &= f_{\varkappa-1} \langle u_{\varkappa-1} \rangle \varphi_{i^{\varkappa-2}}(u_{\varkappa-2}), \\ f_{\varkappa-1} \langle u_{\varkappa-1} \rangle &= \varphi_{i^{\varkappa-1}}(u_{\varkappa-1}). \end{aligned} \quad (36)$$

Из этих соотношений непосредственно следует:

$$\varphi_{i^m}(u_m) = \frac{f_m \langle u_{\varkappa-1} u_{\varkappa-2} \dots u_m \rangle}{f_{m+1} \langle u_{\varkappa-1} u_{\varkappa-2} \dots u_{m+1} \rangle}, \quad m = 0, 1, \dots, \varkappa - 2. \quad (37)$$

Анализируя (36) нетрудно заметить, что последовательность вспомогательных функций представляют собой выходы цепочки фильтров нижних

фрактальностей (см. рис. 23). Для построения вспомогательных функций можно использовать любой метод фрактальной фильтрации. Это означает, что задача построения приспособленного преобразования имеет множество возможных решений.

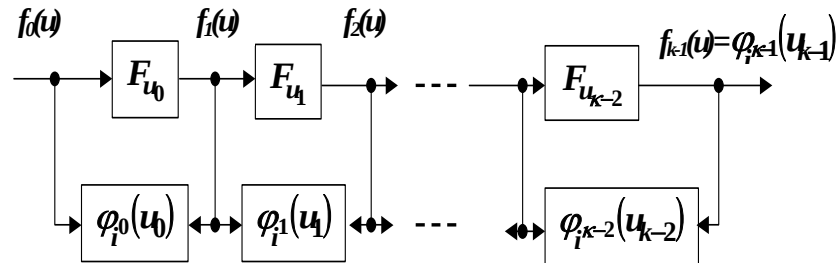


Рис. 23. Схема расчета аппроксимирующих компонент

Положительным аспектом множественности является простое разрешение проблемы цензурирования нулей для знаменателя выражения (37). Нулевое значение знаменателя всегда можно заменить на любое удобное ненулевое (например, единичное). Это ведет к локальному изменению фрактального фильтра, но не влияет на результат аппроксимации.

Множественное приспособление. Если область возможных значений разрядной переменной v_0 соответствует интервалу $[0, g_0 - 1]$, то с помощью выше изложенного алгоритма можно настроить преобразование точно на g_0 функций. При этом остаются элементы ядер, которые могут быть выбраны произвольно.

На рис. 24 показано факторизованное представление быстрого преобразования приспособленного к двум функциям, размещенным в двух первых столбцах матрицы. Символом «*» выделены элементы ядер, которые выбираются произвольно. Не задействованные степени свободы можно использовать для расширения опорной базы приспособления. Это можно реализовать за счет использования фрактально фильтрованных образов добавочных функций. Рассмотрим алгоритм подобного приспособления.

Пусть F — это система из $M = g_0 g_1 \dots g_{k-1}$ функций заданных на интервале длиной $N = p_0 p_1 \dots p_{k-1}$. Для аппроксимации функций системы F будем использовать БНС, реализующую перестраиваемое линейное преобразование размерности $N \times M$.

v_2	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1
v_1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1
v_0	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1
u_2	u_1	u_0																						
0	0	0	1	1					1		*						1			*				
0	0	1	1	1					1		*						1			*				
0	1	0		1	1				1		*							*			*			
0	1	1		1	1				1		*							*				*		
1	0	0				1	1					1		*			1		*					
1	0	1				1	1						1		*		1		*					
1	1	0					1	1				1		*				*			*			
1	1	1					1	1					1		*			*				*		

Рис. 24. Факторизованное преобразование приспособленное к одной функции

Пусть $x = \langle x_{\mathcal{K}-1} x_{\mathcal{K}-2} \dots x_0 \rangle$ — некоторый базовый столбец матрицы преобразования, который должен быть точно приспособлен к соответствующей функции системы. Тогда все столбцы $x = \langle x_{\mathcal{K}-1} x_{\mathcal{K}-2} \dots x_1 v_0 \rangle$, где v_0 варьируется в интервале $[0, g_0 - 1]$, также будут точно приспособлены к соответствующим функциям опорной системы. В алгоритме множественного приспособления настройка БНС выполняется путем последовательной фрактальной фильтрации всего массива функций одновременно, при этом на каждом шаге строится опорное множество функций, фрактальные множители которых используются для заполнения нейронных ядер БНС. Опорные множества функций вложены друг в друга и последовательно окутывают базовый столбец. Правило формирования окутывающей последовательности в зависимости от номера итерации (номера слоя) показано в табл. 3.

Сравнивая правило формирования опорных множеств с правилом нумерации нейронных ядер (последний столбец таблицы 3), можно сделать вывод, что в результате последовательного выполнения алгоритма все ядра преобразования будут заполнены полностью. Перед выполнением алго-

Таблица 3. Правило формирования опорных множеств аппроксимируемых функций

Номер слоя m	Опорное множество функций	Размер опорного множества	Нумерация нейронных ядер i^m
0	$\langle x_{\varkappa-1} x_{\varkappa-2} \dots x_1 v_0 \rangle$	g_0	$\langle u_{\varkappa-1} u_{\varkappa-2} \dots u_1 \rangle$
1	$\langle x_{\varkappa-1} x_{\varkappa-2} \dots x_2 v_1 v_0 \rangle$	$g_1 g_0$	$\langle u_{\varkappa-1} u_{\varkappa-2} \dots u_2 v_0 \rangle$
2	$\langle x_{\varkappa-1} x_{\varkappa-2} \dots x_3 v_2 v_1 v_0 \rangle$	$g_2 g_1 g_0$	$\langle u_{\varkappa-1} u_{\varkappa-2} \dots u_3 v_1 v_0 \rangle$
...	...	...	...
$\varkappa - 2$	$\langle x_{\varkappa-1} v_{\varkappa-2} \dots v_1 v_0 \rangle$	$g_{\varkappa-2} \dots g_1 g_0$	$\langle u_{\varkappa-1} v_{\varkappa-3} \dots v_1 v_0 \rangle$
$\varkappa - 1$	$\langle v_{\varkappa-1} v_{\varkappa-2} \dots v_1 v_0 \rangle$	$g_{\varkappa-1} \dots g_1 g_0$	$\langle v_{\varkappa-2} \dots v_1 v_0 \rangle$

ритма, функции опорной системы упорядочиваются по требуемой глубине приспособления. Функции с номером базового столбца и ближайшие к ней, составляющие опорное множество для $m = 0$ будут точно приспособлены, функции опорного множества для $m = 1$ будут приспособлены по одному фильтрованным образам, функции опорного множества для $m = 2$ — будут приспособлены по дважды фильтрованным образам и т. д.

Приспособленные преобразования в нечетком пространстве

Реализация приспособленных преобразований в нечетком пространстве имеет по крайней мере два преимущества: во-первых, обеспечивается высокая вычислительная эффективность — поскольку логические операции выполняются существенно быстрее арифметических и, во-вторых, отсутствует необходимость цензурирования нулей на выходе фрактальных фильтров — поскольку при вычислении фрактальных компонент функции не используется операция деления.

Геометрия нечеткого пространства определяется операциями $\min$ и $\max$, действующих на ограниченном отрезке $[0, 1]$ множества вещественных чисел. Пусть $f(u)$ — произвольная вещественная функция, заданная на интервале длиной $N = p_0 p_1 \dots p_{\varkappa-1}$. В нечетком пространстве область возможных значений функции принадлежит диапазону $[0, 1]$. С целью упрощения записи для операций $\min$ и $\max$ далее будем использовать обозначения « $\circ$ » и « $\oplus$ », соответственно, рассматривая их как логическое умножение и логическое сложение.

Ставится задача построить быстрое преобразование в нечетком пространстве, приспособленное к заданной функции. Будем использовать для этой цели БНС, в которой арифметические операции сложения и умножения вещественных чисел заменены соответствующими логическими операциями. Задача сводится к нахождению для функции $f(u)$ фрактального логического разложения:

$$f(u) = f\langle u_{\varkappa-1}u_{\varkappa-2}\dots u_0\rangle = \varphi_{i^0}(u_0) \circ \varphi_{i^1}(u_1) \circ \dots \circ \varphi_{i^{\varkappa-1}}(u_{\varkappa-1}),$$

где $i^m = \langle u_{\varkappa-1}u_{\varkappa-2}\dots u_{m+1}\rangle$ — позиционный номер функции-компоненты масштабного уровня m . Функции сомножители можно определить, рекуррентно раскладывая исходную функцию в последовательность логических произведений:

$$\begin{aligned} f(u) &= f_0\langle u_{\varkappa-1}u_{\varkappa-2}\dots u_0\rangle, \\ f_0\langle u_{\varkappa-1}u_{\varkappa-2}\dots u_0\rangle &= f_1\langle u_{\varkappa-1}u_{\varkappa-2}\dots u_1\rangle \circ \varphi_{i^0}(u_0), \\ f_1\langle u_{\varkappa-1}u_{\varkappa-2}\dots u_1\rangle &= f_2\langle u_{\varkappa-1}u_{\varkappa-2}\dots u_2\rangle \circ \varphi_{i^1}(u_1), \\ &\vdots \\ f_{\varkappa-2}\langle u_{\varkappa-1}u_{\varkappa-2}\rangle &= f_{\varkappa-1}\langle u_{\varkappa-1}\rangle \circ \varphi_{i^{\varkappa-2}}(u_{\varkappa-2}), \\ f_{\varkappa-1}\langle u_{\varkappa-1}\rangle &= \varphi_{i^{\varkappa-1}}(u_{\varkappa-1}). \end{aligned} \quad (38)$$

Существует определенная свобода в выборе вспомогательных функций $f_1, \dots, f_{\varkappa-1}$. Оставаясь в рамках нечеткого пространства, будем использовать для построения вспомогательных функций ранговый фрактальный фильтр следующего вида:

$$f_{m+1}\langle u_{\varkappa-1}u_{\varkappa-2}\dots u_{m+1}\rangle = \max_{u_m}(f_m\langle u_{\varkappa-1}u_{\varkappa-2}\dots u_m\rangle).$$

В этом случае из (38) следует, что функции-компоненты будут определяться выражением:

$$\varphi_{i^m}(u_m) = f_{m+1}\langle u_{\varkappa-1}u_{\varkappa-2}\dots u_{m+1}\rangle \circ f_m\langle u_{\varkappa-1}u_{\varkappa-2}\dots u_m\rangle.$$

Дальнейшая процедура настройки БНС не отличается от аналогичной процедуры арифметического пространства. Преобразование, приспособленное к одной функции, имеет свободные элементы нейронных ядер, которые могут быть заполнены при добавлении опорных функций, подобно тому, как это было сделано для преобразований в арифметическом пространстве.

При обработке данных с помощью нейронной сети выполняются операции « $\circ$ » и « $\oplus$ ». Если входной образ совпадает с опорной функцией приспособленного преобразования, то на соответствующем выходе сети будет получено максимальное значение. На этом может быть основано классификационное правило распознавания сигналов.

В нечетком пространстве операции « $\circ$ » и « $\oplus$ » полностью симметричны, поэтому существует двойственная форма фрактального представления функций, которая может быть задана выражениями:

$$\begin{aligned} f(u) &= f\langle u_{\mathcal{X}-1} u_{\mathcal{X}-2} \dots u_0 \rangle = \varphi_{i^0}(u_0) \oplus \varphi_{i^1}(u_1) \oplus \dots \oplus \varphi_{i^{\mathcal{X}-1}}(u_{\mathcal{X}-1}) \\ f_{m+1}\langle u_{\mathcal{X}-1} u_{\mathcal{X}-2} \dots u_{m+1} \rangle &= \min_{u_m} (f_m\langle u_{\mathcal{X}-1} u_{\mathcal{X}-2} \dots u_m \rangle), \\ \varphi_{i^m}(u_m) &= f_{m+1}\langle u_{\mathcal{X}-1} u_{\mathcal{X}-2} \dots u_{m+1} \rangle \oplus f_m\langle u_{\mathcal{X}-1} u_{\mathcal{X}-2} \dots u_m \rangle. \end{aligned}$$

Спектральные приспособленные преобразования

Для задач классификации и распознавания сигналов существенное значение имеют процедуры предварительной обработки, ориентированные на устранение избыточности и выделении информативных признаков. Использование спектральных преобразований для этих целей позволяет представить информацию, содержащуюся в исходном сигнале в виде взаимно-независимых ортогональных составляющих. Поскольку энергия сигнала определяется суммой квадратов спектральных коэффициентов, то по величине спектрального коэффициента можно непосредственно судить о значимости информативного признака. Известно, что максимальное сокращение избыточности данных обеспечивается ортогональным преобразованием Карунена-Лоэва, которое образуется собственными векторами ковариационной матрицы сигнальной группы. Однако использование данного преобразования сопряжено со значительными вычислительными затратами. По этой причине метод Карунена-Лоэва не применяется для обработки данных высокой размерности. Если при построении преобразования Карунена-Лоэва ограничиться только одним собственным вектором, имеющим максимальный ранг (главной компонентой), то объем вычислений резко сокращается. Достаточно часто главную компоненту можно выделить из сигнала непосредственно, как опорный образ, относительно которого наблюдаются все сигнальные вариации. Ортогональное преобразование, которое настроено на одну главную компоненту относится к классу приспособленных. В силу ортогональности результатом воздействия приспособленного преобразования на опорный образ является выходной вектор,

у которого одна из координат равна единице, в то время как все остальные координаты равны нулю. Таким образом, приспособленное ортогональное преобразование идеально подходит для селективного распознавания эталонного сигнала. Построение быстрых приспособленных ортогональных преобразований для сигналов определенных на интервале с длиной кратной степени двойки рассматривалось в работе [5]. Ниже будет показано использование БНС для синтеза приспособленного ортогонального преобразования любой составной размерности.

Достаточные условия приспособленности. Для спектральных преобразований удобно использовать конструктивное определение приспособленности, которое можно сформулировать следующим образом. Ортогональное преобразование, заданное матрицей $h(u, v)$, приспособлено к функции $f(u)$ по столбцу, x если выполнены условия:

$$\begin{cases} \sum_u h(u, v) f(u) = 1, & \text{при } v = x, \\ \sum_u h(u, v) f(u) = 0, & \text{при } v \neq x. \end{cases} \quad (39)$$

Здесь предполагается, что опорная функция нормирована условием

$$\sum_u f^2(u) = 1.$$

Подставив в левую часть (39) выражения (34) и (35), после преобразований получим:

$$\begin{aligned} \sum_u h(u, v) f(u) &= \sum_{u_0} w_{i^0}(u_0, v_0) \varphi_{i^0}(u_0) \sum_{u_1} w_{i^1}(u_1, v_1) \varphi_{i^1}(u_1) \times \dots \\ &\dots \times \sum_{u_{\varkappa-1}} w_{i^{\varkappa-1}}(u_{\varkappa-1}, v_{\varkappa-1}) \varphi_{i^{\varkappa-1}}(u_{\varkappa-1}). \end{aligned}$$

Из данной формулы следует, что приспособленность спектрального преобразования обеспечивается, когда нейронные ядра любого слоя m приспособлены к соответствующим компонентам фрактального разложения опорной функции. Иначе это достаточное условие можно записать в виде системы уравнений:

$$\begin{cases} \sum_{u_m} w_{i^m}(u_m, v_m) \varphi_{i^m}(u_m) = 1, & \text{при } v_m = x_m, \\ \sum_{u_m} w_{i^m}(u_m, v_m) \varphi_{i^m}(u_m) = 0, & \text{при } v_m \neq x_m. \end{cases}$$

Таким образом, задача настройки БНС сводится к построению приспособленных ортогональных ядер. Рассмотрим возможный алгоритм такого построения.

Алгоритм построения приспособленных ортогональных ядер. Пусть компонента фрактального разложения представлена вектором вида $\varphi = (\varphi_0 \varphi_1 \dots \varphi_{p-1})$, размерности p . Разместим этот вектор в первом столбце ортогональной матрицы. Понятно, что существует множество способов достроить остальные столбцы матрицы. Алгоритм построения во многом определяется выбранной топологией матрицы. Поскольку никаких других требований на ортогональную матрицу не накладывается, то при выборе ее топологии целесообразно руководствоваться принципом минимальной длины описания (иначе именуемым «брита Оккама»).

Вектор-образ может содержать нулевые значения, не теряя общности, можно полагать, что все нулевые значения размещены в конце вектора начиная с позиции $k + 1$. Предлагаемая топология матрицы показана на рис. 25. В данной матрице $y_1, y_2, \dots, y_k$ неизвестные параметры, определяемые из условий ортогональности столбцов:

$$\begin{pmatrix} \varphi_0 & \varphi_0 & \varphi_0 & \dots & \varphi_0 & \varphi_0 & 0 & 0 & 0 & 0 \\ \varphi_1 & \varphi_1 & \varphi_1 & \dots & \varphi_1 & y_1 & 0 & 0 & 0 & 0 \\ \dots & \dots & \dots & \dots & y_2 & 0 & 0 & 0 & 0 & 0 \\ \dots & \dots & \varphi_{k-2} & \dots & 0 & 0 & 0 & 0 & 0 & 0 \\ \dots & \varphi_{k-1} & y_{k-1} & \dots & 0 & 0 & 0 & 0 & 0 & 0 \\ \varphi_k & y_k & 0 & \dots & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & \dots & 0 & 0 & 0 & 0 & 0 & 1 \end{pmatrix}$$

Рис. 25. Топология матрицы приспособленного ядра

$$y_{\alpha} \varphi_{\alpha} + \sum_{u=0}^{\alpha-1} \varphi_u^2 = 0.$$

Для построения ортогонального ядра W достаточно в построенной матрице произвести нормировку столбцов, так чтобы выполнялись условия:

$$\sum_{u=0}^{p-1} w^2(u, v) = 1.$$

Множественное приспособление. Ортогональное приспособленное преобразование не полностью использует степени свободы БНС. Если точка приспособления совпадает с первым столбцом результирующей матрицы, то будут настроены только ядра

$$i^m = \langle u_{\kappa-1} u_{\kappa-2} \dots u_{m+1} 0_{m-1} 0_{m-2} \dots 0_0 \rangle.$$

На рис. 26 приведено факторизованное представление ортогонального преобразования приспособленного к одной функции размещенной в первом столбце. Символом «*» отмечены элементы незаполненных ядер.

Для того чтобы преобразование было ортогональным, достаточно заполнить свободные ядра элементами любых ортогональных матриц соответствующей размерности. В практических алгоритмах, можно либо заполнить свободные ядра единичными матрицами вида:

$$W = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix},$$

либо повторить ядра, уже найденные для текущего слоя. В первом случае будут получены вейвлет-подобные ортогональные преобразования с явно выраженными свойствами пространственной локализации, во втором случае базисные функции будут фрактально близки друг к другу (т. е. их фильтрованные образы могут совпадать на некоторых масштабных уровнях). Использовать более глубокое приспособление, например, к фильтрованным образам опорных функций (как это было сделано в предыдущем параграфе) возможно только тогда, когда опорные функции фрактально ортогональны между собой (т. е. ортогональны на всех масштабных уровнях). В общем случае условие фрактальной ортогональности не выполняется.

v_2	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1	0	0	0	0	1	1	1	1	1
v_1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1	0	0	1	1	1
v_0	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	0	1	1
u_2	u_1	u_0																							
0	0	0	1	1													1								
0	0	1	1	1													*		*						
0	1	0			1	1											1		1						
0	1	1			1	1											*		*						
1	0	0					1	1											1		1				
1	0	1					1	1						*	*		*				*				
1	1	0							1		1							*				*			
1	1	1									*	*						*				*			

Рис. 26. Топология ортогонального преобразования, приспособленного к одной функции

Быстрые нейронные сети в квантовых вычислениях

Квантовые вычисления (вычисления, основанные на методах квантовой механики) были инициированы в начале 80-х годов работами *Ричарда Фейнмана* [25, 26]. В 1994 году *Питер Шор* показал, что NP-сложная задача — разложение целых чисел на множители — может быть решена на квантовом компьютере за полиномиальное время [27]. Это результат дал толчок к бурному развитию квантовых вычислений. В августе 2000 года агентством Reuters было объявлено о первом практическом успехе: “IBM says it develops most advanced quantum computer («фирма IBM разработала первый квантовый компьютер»)). В данной разработке впервые был реализован квантовый регистр на 5 атомах. Практически была доказана принципиальная возможность реализации квантовых компьютеров.

В классическом компьютере регистр длиной n имеет конечное множество состояний равно 2^n , поскольку каждый бит имеет только два состояния. В квантовом компьютере функциональным аналогом бита является q -бит, который ассоциируется с двумерным комплексным пространством

и поэтому имеет континуум возможных состояний. В пространстве q -бита выделены два базисных вектора $|0\rangle$ и $|1\rangle$. Кроме базисных состояний для q -бита возможны также состояния, представляющие собой любые линейные комбинации базисных векторов: $|\phi\rangle = \alpha|0\rangle + \beta|1\rangle$, где $\alpha^2 + \beta^2 = 1$. Квантовый регистр представляет собой конечный набор q -битов. Базисными состояниями квантового регистра являются тензорные произведения базисных состояний входящих в него q -бит. Для обозначения тензорных произведений в квантовой механике используются сокращения:

$$|x_{n-1}\rangle \otimes |x_{n-2}\rangle \otimes \dots \otimes |x_0\rangle = |x_{n-1}x_{n-2} \dots x_0\rangle.$$

Согласно общим принципам квантовой механики, возможным состоянием квантового регистра длиной n может быть любой вектор единичной длины в комплексном пространстве размерности 2^n . Квантовое состояние регистра записывается в виде:

$$|\varphi\rangle = \sum_{|x\rangle \in |x_{n-1}x_{n-2} \dots x_0\rangle} a_x |x\rangle,$$

где суммирование производится по всем возможным базисным состояниям. Для вектора единичной длины сумма квадратов коэффициентов a_x всегда равна единице. Квантовые состояния в комплексном пространстве задаются с точностью до фазового множителя $e^{j\phi}$. Квадрат значения каждого коэффициента a_x интерпретируется как вероятность нахождения квантового регистра в данном базовом состоянии. Состояние квантового регистра называется запутанным, если его не возможно представить в виде произведения квантовых состояний входящих в него q -бит.

Эволюция состояний квантового регистра во времени описывается уравнением Шредингера, решением которого являются унитарные преобразования размерности 2^n . При квантовых вычислениях в q -регистре готовится начальное состояние в виде базовых состояний q -бит. Затем используется некоторое унитарное преобразование над отдельными q -битами, зависящее от решаемой задачи, в итоге возникает новое состояние квантового регистра. Результатом вычислительной операции считается состояние квантового регистра, которое имеет вероятность достаточно близкую к единице. Эффективность квантовых вычислений достигается за счет того, что одновременно с унитарным преобразованием над состояниями q -бита «бесплатно» выполняется тензорное произведение, экспоненциально расширяющее область действия унитарного преобразования. Наибольший

эффект достигается для факторизуемых квантовых состояниях, однако и для широкого класса запутанных состояний, возможна реализация полиномиально эффективных алгоритмов.

Элементарным квантовой операция на регистре определяется двумерным унитарным преобразованием, действующим в пространстве одного q -бита. Данное преобразование (унитарный вентиль) задается комплекснозначной матрицей размером 2×2 :

$$W = \begin{array}{c|cc} & |0\rangle & |1\rangle \\ \hline |0\rangle & w_{00} & w_{01} \\ |1\rangle & w_{10} & w_{11} \end{array}$$

В канонической форме унитарная матрица второго порядка может быть записана в виде:

$$W = \begin{pmatrix} e^{j\gamma_{00}} \cos \theta & e^{j\gamma_{01}} \sin \theta \\ e^{j\gamma_{10}} \sin \theta & -e^{j\gamma_{11}} \cos \theta \end{pmatrix}. \quad (40)$$

Условие унитарности в этом случае сводится к соотношению между углами комплексных экспонент:

$$\gamma_{00} - \gamma_{01} = \gamma_{10} - \gamma_{11}.$$

Бинарные вентили задаются унитарными матрицами размера 4×4 . Примером может служить элемент «управляемый фазовый сдвиг» [28], задаваемый матрицей:

$$W = \begin{array}{c|cccc} & |00\rangle & |01\rangle & |10\rangle & |11\rangle \\ \hline |00\rangle & 1 & 0 & 0 & 0 \\ |01\rangle & 0 & 1 & 0 & 0 \\ |10\rangle & 0 & 0 & 1 & 0 \\ |11\rangle & 0 & 0 & 0 & e^{j\varphi} \end{array}$$

Можно показать, что результирующее состояние регистра после бинарного вентиля в общем случае является запутанным. Последовательная комбинация элементарных преобразований в принципе позволяет выполнить любое квантовое вычисление. Существенным вопросом является эффективность вычислений. В работе [27] для построения эффективного квантового вычисления Шор использовал модификацию быстрого алгоритма Фурье с основанием 2. Как было показано в предыдущих разделах,

быстрое преобразование Фурье может быть реализовано в классе быстрых нейронных сетей, а в общем случае БНС могут быть использованы для формирования широкого набора унитарных факторизуемых преобразований.

Применительно к квантовым вычислениям размерность нейронной сети должна быть равна степени двойки. Это ведет к тому, что все нейронные ядра сети структурно-подобны и представляют собой унитарные матрицы размера 2×2 . Если все ядра в пределах слоя параметрически одинаковы, то нейронный слой однозначно сопоставляется с одним q -битом квантового регистра. В этом случае эквивалентный унитарный оператор нейронной сети, описывающий эволюцию состояния q -битного регистра размерности n , выражается кронекеровским произведением матриц нейронных ядер: $H = W_0 \times W_1 \times \dots \times W_{n-1}$. На рис. 27 показана квантовая схема подобного вычисления. Для данной схемы число квантовых операций равно размерности квантового регистра и, следовательно, данное преобразование является максимально эффективным. Результирующее состояние для данной квантовой схемы всегда факторизуемо.

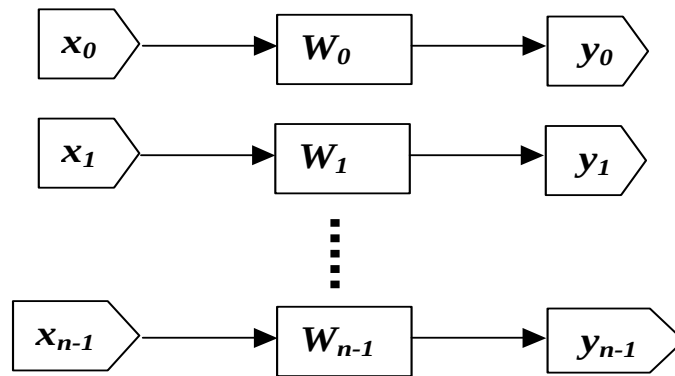


Рис. 27. Квантовая схема БНС с однородными ядрами в слое

Следует отметить, что моделирование квантового регистра на обычном компьютере приводит к NP -сложной задаче. На рис. 28 показана квантовая схема быстрого преобразования Фурье для размерности 2^3 . В отличие от предыдущего случая квантовая схема содержит бинарные вентили управляемого фазового сдвига. Значение фазового сдвига и число бинар-

ных вентилях зависит от номера нейронного слоя. Все унарные вентили определяются матрицей преобразования Адамара:

$$W = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}.$$

Детали построения квантовой схемы изложены в работе [29]. Результирующее состояние квантового регистра является запутанным, тем не менее, квантовая схема является полиномиально эффективной.

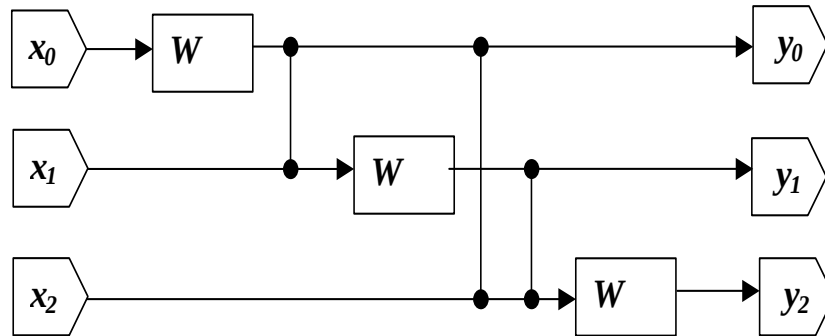


Рис. 28. Квантовая схема быстрого преобразования Фурье

Нейронная технология обычно имеет дело с образами. В квантовом варианте образом является квантовое состояние регистра. Задача распознавания образа сводится к построению такого унитарного преобразования, которое формирует на выходе одно из базисных состояний регистра с вероятностью близкой или равной единице. По существу подобная задача была рассмотрена при построении приспособленных спектральных преобразований. Было показано, что настройка сводиться к фрактальной декомпозиции функции и выбору параметров нейронных ядер. Методы декомпозиции основаны на фрактальной фильтрации, и отличаются только комплексным характером фрактальных фильтров.

Для нейронных ядер размерности два можно предложить более общий способ выбора параметров ядер, чем было изложено ранее. Обозначим через $\dot{w}_i(u, v)$ коэффициенты нейронного ядра в слое i , а через $\dot{\varphi}_i(u)$ (здесь и далее точка над символом означает комплексную величину, а «черточка» — комплексно сопряженную) — соответствующую компоненту фрактального

разложения функции-образа $\dot{f}(u)$. Значениями функции $\dot{f}(u)$, заданной на интервале длиной 2^n , являются коэффициенты базисных векторов состояния квантового регистра. Как было показано ранее, спектральный оператор будет приспособлен к функции, если нейронные ядра будут приспособлены к соответствующим компонентам фрактального разложения функции. Для унитарных ядер условие приспособленности имеет вид:

$$\begin{cases} \sum_{u_m} \dot{\varphi}_{im}(u_m, v_m) \bar{\varphi}_{im}(u_m) = 1, & \text{при } v_m = x_m, \\ \sum_{u_m} \dot{\varphi}_{im}(u_m, v_m) \bar{\varphi}_{im}(u_m) = 0, & \text{при } v_m \neq x_m. \end{cases}$$

Например, для компоненты $\dot{\varphi}_{im}(u_m) = \varphi_{im}(u_m) e^{j\alpha_{u_m}}$ и ядра вида (40) приспособленного по первому столбцу данные условия приводятся к виду:

$$\begin{aligned} \varphi_{im}(0) e^{j(\gamma_{00}-\alpha_0)} \cos \theta + \varphi_{im}(1) e^{j(\gamma_{10}-\alpha_1)} \sin \theta &= 1, \\ \varphi_{im}(0) e^{j(\gamma_{01}-\alpha_0)} \sin \theta - \varphi_{im}(1) e^{j(\gamma_{11}-\alpha_1)} \cos \theta &= 0. \end{aligned}$$

Откуда следуют расчетные формулы для определения параметров ядер:

$$\operatorname{tg} \theta = \varphi_{im}(1)/\varphi_{im}(0), \quad \gamma_{00} = \alpha_0, \quad \gamma_{10} = \alpha_1, \quad \gamma_{01} - \alpha_0 = \gamma_{11} - \alpha_1.$$

Данными формулами углы γ_{01}, γ_{11} определяются неоднозначно, для определенности можно положить $\gamma_{01} = \alpha_0, \gamma_{11} = \alpha_1$. Если исходное состояние квантового регистра совпадает с функцией образа, то квантовое состояние после воздействия оператора нейронной сети будет в точности совпадать с одним из базовых состояний. В общем случае БНС выполняет элементарные преобразования над состояниями квантового регистра, а не отдельных q -бит. Поэтому в случае сильно запутанных состояний вычислительная эффективность квантовых БНС падает до NP -класса.

Многомерные БНС

Парадигма многослойных БНС достаточно просто распространяется на многомерное пространство. Грамматика формального языка полностью сохраняется, основные отличия заключаются в многомерной интерпретации букв алфавита. Для d -мерной сети любая буква алфавита заменяется символьным вектором, в котором каждая координата соответствует одной из размерностей сети.

Рассмотрим, например БНС размерности $d = 2$. Любая буква алфавита в этом случае представляется двумерным вектором. Будем кодировать координаты каждого вектора двухсимвольным набором, где первый символ выбирается из множества i, j, u, v и означает родовое имя буквы, а второй символ (x либо y) означает пространственное направление. В такой интерпретации номер ядра в слое λ обозначается символьным вектором $i^\lambda = (ix^\lambda \ iy^\lambda)$. Ядра слоя располагаются в плоскости, координаты ядер задаются значениями компонент вектора i^λ . Синаптическая карта каждого нейронного ядра представляет собой четырехмерную матрицу с элементами $w_{i^\lambda}^\lambda(u_\lambda, v_\lambda)$, где значения $u_\lambda = (ux_\lambda \ uy_\lambda)$ определяют локальные пространственные координаты рецепторов, а $v_\lambda = (vx_\lambda \ vy_\lambda)$ — локальные координаты аксонов ядра (см. рис. 29).

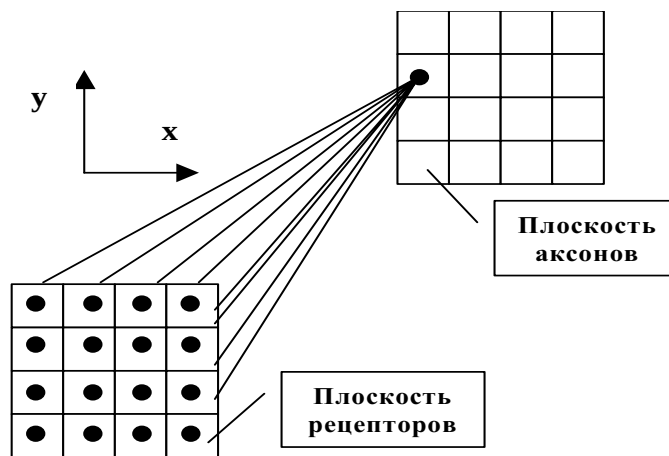


Рис. 29. Нейронное ядро двумерной БНС

Структурная модель $\mathcal{K}$ -слойной БНС как и прежде задается каноническим предложением состоящим из $\mathcal{K}$ слов, упорядоченных по числу вхождений двумерных букв с родовым именем i . При построении графической интерпретации, дугами соединяются нейронные ядра смежных слоев, которые имеют совпадающие значения для одноименных двумерных букв. На рис. 30 показано правило построения структурной модели для канонического предложения:

$$[\langle i_1 i_2 \rangle \langle i_2 j_1 \rangle \langle j_1 j_0 \rangle] .$$

Все разрядные переменные имеют основания равные двум. Пространственные координаты ядер определяются словами $\langle ix_1ix_2 \rangle, \langle iy_1iy_2 \rangle$ — для слоя 0, $\langle ix_2jx_1 \rangle, \langle iy_2jy_1 \rangle$ — для слоя 1, и $\langle jx_1jx_0 \rangle, \langle jy_1jy_0 \rangle$ — для слоя 2.

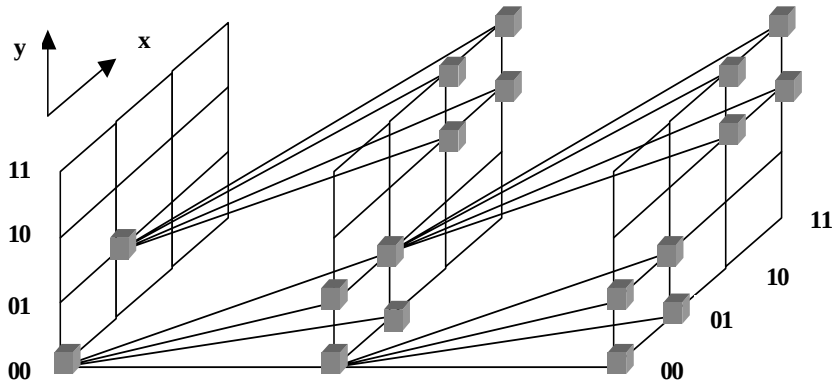


Рис. 30. Правило построения двумерной БНС

Литература

1. Дорогов А. Ю., Алексеев А. А., Буторин Д. А. Нейронные сети со структурой быстрого алгоритма // *Нейроинформатика и ее приложения: 6 Всерос. семинар, 20–25 октября 1998г., Красноярск. Тезисы докл.* – Красноярск, КГТУ, 1998. – С. 53.
2. Good I. J. The Interaction Algorithm and Practical Fourier Analysis // *Journal of Royal Statistical Society.* – Ser. B. – 1958. – Vol. 20. – No. 2. – P. 361–372.
3. Andrews H. C., Caspari K. L. A General Techniques for Spectral Analysis // *IEEE Trans. Computers.* – 1970. – Vol. C-19. – Jan., No. 1. – P. 16–25.
4. Эндрюс Г. Применение вычислительных машин для обработки изображений: Перевод с англ. под ред. Б. Ф. Курьянова. – М., 1977. – 160 с.
5. Солодовников А. И., Спиваковский А. М. Основы теории и методы спектральной обработки информации. – Л., 1986. – 272 с.
6. Лабунец В. Г. Единый подход к алгоритмам быстрых преобразований // *Применение ортогональных методов при обработке сигналов и анализа систем: Межвуз. Сб.* – Свердловск: Уральск. Политехн. Ин-т. – 1980. – С. 4–14.

7. *Cooley J. W., Tukey J. W.* An Algorithm for the Machine Computation of Complex Fourier Series // *Math. Comp.* – 1965. – No. 19. – P. 297–301.
8. *Рабинер Л., Гоулд.* Теория и применение цифровой обработки сигналов: Пер. с англ. – М.: Мир, 1978. – 848 с.
9. *Трахтман А. М., Трахтман В. А.* Основы теории дискретных сигналов на конечных интервалах. – М.: Сов. Радио, 1975. – 208 с.
10. *Дорогов А. Ю.* Быстрые нейронные сети. – СПб.: Изд-во С. Петерб. ун-та, 2002. – 80 с.
11. *Дорогов А. Ю.* Структурный анализ слабосвязанных нейронных сетей // *Управление в социальных, экономических и технических системах*. Кн. 3. Управление в технических системах: Труды Межреспубликанской научной конференции, г. Кисловодск 28 июня–2 июля 1998 г. – С. 75–80.
12. *Дорогов А. Ю.* Структурный синтез модульных слабосвязанных нейронных сетей. Часть 1. Методология структурного синтеза модульных нейронных сетей // *Кибернетика и системный анализ*. – 2001. – № 2. – С. 34–42.
13. *Дорогов А. Ю.* Структурный синтез модульных слабосвязанных нейронных сетей. Часть 2. Ядерные нейронные сети // *Кибернетика и системный анализ*. – 2001. – № 4. – С. 13–20.
14. *Uhr L.* Algorithm-Structured Computer Arrays and Networks: Parallel Architectures for Perception and Modelling. – Academic Press, New York, 1984.
15. *Дорогов А. Ю.* Генезис слабосвязанных нейронных сетей // Всероссийская научно-техническая конференция «Нейроинформатика-99», Москва, 20–22 января 1999 г. – Сб. науч. тр. Часть 1. – М.: 1999. – С. 64–70.
16. *Эдельман Дж., Маункасл В.* Разумный мозг: Кортикальная организация и селекция групп в теории высших функций головного мозга / Пер. с англ. *Н. Ю. Алексеенко*, под ред. *Е. К. Соколова*. – М.: Мир. – 1981. – 133 с.
17. *Дорогов А. Ю.* Фракталы и нейронные сети // *Проблемы нейрокибернетики* (материалы Юбилейной международной конференции по нейрокибернетике посвященной 90-летию со дня рождения проф. А. Б. Когана, 23–29 октября 2002 г., Ростов-на-Дону). Том 2. – Ростов-на-Дону. – 2002. – С. 9–14.
18. *Дорогов А. Ю., Алексеев А. А.* Пластичность многослойных слабосвязанных нейронных сетей // *Нейрокомпьютеры: разработка и применение* № 11, 2001, с. 22–40.
19. *Малинецкий Г. Г., Потапов А. Б.* Современные проблемы нелинейной динамики. – М.: Эдиториал УРСС, 2000. – 336 с.
20. *Кроновер Р. М.* Фракталы и хаос в динамических системах. – М.: Постмаркет, 2000. – 350 с.

21. Макаренко Н. Г. Фракталы, аттракторы, нейронные сети и все такое // «Нейроинформатика 2002»: 4-й Всеросс. науч. техн. конф. 23–25 января 2002 г., Москва – Лекции по нейроинформатике. – Часть 2. – с. 136–169.
22. Stark J. Iterated function systems as neural networks // *Neural Networks*. – 1991. – V. 4. – pp. 679–690.
23. Астафьева Н. М. Вейвлет-анализ: основы теории и примеры применения // *Успехи физических наук*. – Т. 166. – 1996. – № 11. – С. 1145–1170.
24. Каппелини В. и др. Цифровые фильтры и их применение: Пер. с англ. – М.: Энергоатомиздат, 1983. – 360 с.
25. Feynman R. Simulating physics with computers // *International Journal of Theoretical Physics*, **21**, 6&7, 467–488, 1982.
26. Feynman R. Quantum mechanical computers // *Optics News* **11**, 1985. Also in *Foundations of Physics*, **16**: 507–531, 1986.
27. Shor P. W. Algorithms for quantum computation: Discrete log and factoring // In: *Proceedings of the 35th Annual Symposium on Foundations of Computer Science* (Nov. 1994). pp. 124–134, 1994. Institute of Electrical and Electronic Engineers Computer Society Press.
URL: <ftp://netlib.att.com/netlib/att/math/shor/quantum.algorithms.ps.Z>.
28. Риффель Э., Полак В. Основы квантовых вычислений // *Квантовый компьютер и квантовые вычисления*. – Т. 1, № 1, 2000. – с. 4–57.
29. Ekert A., Hayden P., Inamori H. Basic concepts in quantum computation. – Centre for Quantum Computation University of Oxford OX1 3PU, United Kingdom.

Александр Юрьевич ДОРОГОВ, кандидат технических наук, доцент кафедры автоматизации и процессов управления Санкт-Петербургского государственного электротехнического университета («ЛЭТИ»). Области научных интересов: нейронные сети, быстрые перестраиваемые преобразования, фракталы, системный анализ. Автор более 100 печатных работ и одной монографии.

В. Г. ЯХНО

Институт прикладной физики РАН, Нижний Новгород

E-mail: yakhno@appl.sci-nnov.ru

НЕЙРОНОПОДОБНЫЕ МОДЕЛИ ОПИСАНИЯ ДИНАМИЧЕСКИХ ПРОЦЕССОВ ПРЕОБРАЗОВАНИЯ ИНФОРМАЦИИ

Аннотация

Приведено конспективное рассмотрение вариантов базовых нейроноподобных моделей, описывающих динамические процессы преобразования информационных потоков.

V. YAKHNO

Institute for Applied Physics, RAS,

Nizhny Novgorod

E-mail: yakhno@appl.sci-nnov.ru

NEURON-LIKE MODELS DESCRIBING DYNAMICAL TRANSFORMATION PROCESSES FOR INFORMATION FLOWS

Abstract

Some variants of basic neuron-like models for dynamical processes of information flow transformation are considered.

Введение

Распределенные системы называют *нейроноподобными*, если они состоят из активных элементов с *несколькими устойчивыми (или «квазиустойчивыми») состояниями* и взаимодействие между такими неравновесными элементами осуществляется за счет *нелокальных пространственных связей*. Динамические режимы преобразования информационных сигналов в нейроноподобных системах и соответствие их экспериментальным данным зависит от вида используемых базовых моделей.

Различные пространственно-временные решения базовых моделей определяются:

- алгоритмами выделения признаков на различных уровнях описания;
- вариантами алгоритмов управления;
- модельными описаниями распознаваемых объектов и моделями работы самой распознающей системы.

Функциональные режимы сложных информационных систем обычно легче понимать, если они рассматриваются с помощью блочных моделей, соответствующих разным функциональным уровням преобразования и управления информационными сигналами. Например, *однородные нейроноподобные системы* позволяют осуществлять режимы параллельного выделения признаков из изображений.

Системы, позволяющие автоматически оценивать качество работы алгоритмов кодирования–декодирования, необходимы для повышения точности распознающих систем и отслеживания особенностей обрабатываемых изображений.

Иерархически организованные взаимосвязанные распознающие системы необходимы для понимания эволюционной динамики мотиваций в работе таких сформировавшихся симбиозов. В них происходит обмен алгоритмическими и «энергетическими» ресурсами, который приводит к процессам усовершенствования, модернизации алгоритмов, используемых в таких взаимосвязанных распознающих системах.

Модели для таких уровней описания могут быть представлены как в виде уравнений, так и в виде функциональных схем, определяющих основные пути и типы операций преобразования изображений. Обычно в распознающих системах информационный сигнал можно представить в виде изображений.

Для понимания особенностей динамики сложных систем очень важен выбор «эффективных» переменных, т. е. тех переменных, которые лучше всего соответствуют функциональному режиму системы. Известно, что более понятной выглядит модель, в которой число необходимых «эффективных» переменных сведено к минимуму. С другой стороны, число используемых переменных в модели должно обеспечивать адекватное описание процессов.

Рассмотрим три типа «базовых» моделей для описания разных функциональных уровней, которые удовлетворяют этим требованиям. Последовательное использование таких моделей позволяет описать широкий круг

динамических процессов во взаимодействующих информационных системах.

Модели с адаптацией параметров

Первый тип моделей представлен однородными нейроноподобными системами, в которых особенности преобразования входного изображения $u = u_0(t_0, \vec{r})$ определяются изменением дополнительной переменной $g_i(t, \vec{r})$. Вариант такой системы имеет вид [1-4, 6-8, 13]:

$$\frac{du}{dt} = -\frac{du}{\tau_1} + \beta_{F_1}(g) * F_1 \left[-\theta_1(g) + \int_{-\infty}^{\infty} \Phi_1[(\xi - r), g] u(\xi, t) d\xi \right] + D_1 \frac{\partial^2 u(t, \vec{r})}{\partial r^2}, \quad (1)$$

$$\frac{dg}{dt} = -\frac{dg}{\tau_2} + \beta_{F_2}(g) * F_2 \left[-\theta_2(g) + \int_{-\infty}^{\infty} \Phi_2[(\xi - r), g] u(\xi, t) d\xi \right] + D_2 \frac{\partial^2 u(t, \vec{r})}{\partial r^2}. \quad (2)$$

Двухкомпонентная модель (1)–(2) получена из уравнений, описывающих взаимодействие нейронов с возбуждающими и тормозными связями в участке коры головного мозга животных, содержащем сотни тысяч нервных клеток в приближении однородности рассматриваемого участка [1–3]. При интегральном описании активности многих ансамблей из нейроноподобных элементов используются две наиболее важные (или можно назвать их — «эффективные») переменные $u(t, \vec{r})$ и $g(t, \vec{r})$. Первая переменная обычно описывает уровень активности популяций возбуждающих нейронов, а вторая переменная — уровень активности популяций тормозных нейронов. Функции

$$\Phi_{mp}(t - \tau, \vec{\xi} - \vec{r})$$

имеют вид так называемого «латерального» торможения,

$$\Phi(R) = (1 - akR^2)e^{-aR^2},$$

а $F[\dots]$ имеет ступенчатый вид. Система (1)–(2) применяется для описания процессов в однородных нейронных сетях сетчатки глаза живых объектов,

коры некоторых отделов головного мозга и т. п. [1–4]. Показано, что в таких системах базовые пространственно-временные автоволновые процессы представлены:

- фронтами переключений между различными стационарными состояниями;
- разнообразными импульсными решениями;
- автономными источниками волн;
- разнообразными структурами синхронизации и фазировки для систем из автоколебательных нейроноподобных элементов.

Примеры результатов исследований процессов подобного рода приведены в работах [1–7]. В моделях таких процессов выходной информационный сигнал в момент времени t_0 представляет собой $u(t_0, \vec{r})$ и $g(t_0, \vec{r})$ — два изображения структур пространственной активности нейроноподобных элементов системы.

Модели адаптивных распознающих систем

Ко **второму типу** моделей относятся «элементарные» *системы принятия решений с фиксированными алгоритмами*, в которых предполагается использование управляющих алгоритмов с целью оптимизации режимов распознавания, как под особенности вида сигнала, так и под задачи, решаемые такой адаптивной распознающей системой. Функциональная модель для элементарных распознающих систем с адаптивными свойствами (см. рис. 1) состоит из:

- *блоков преобразования* (кодирования и восстановления) информационных сигналов;
- *блоков сохранения* кодового описания этих сигналов и используемых для их обработки алгоритмов (моделей);
- *блоков оценки* близости восстановленного сигнала к виду исходного сигнала, и принятия решений [7–13].

Если требуется использование режимов параллельного преобразования сигналов, то для выполнения операций во всех этих блоках могут быть задействованы модели однородных нейроноподобных систем [4, 6–8]. Это означает, что модели второго типа можно полностью построить из слоев, преобразование сигналов в которых описывается моделями первого типа.

Основная особенность представленной базовой модельной системы для процессов адаптивного принятия решений заключается в формальном представлении необходимых и достаточных управляющих воздействий между информационными потоками и используемыми алгоритмами обработки этих потоков. В частности, вычисленные параметры изображений информационных сигналов (входные и промежуточные изображения, разнообразные кодовые описания, оценки точности соответствия, или «невязки» — мотивационные оценки) влияют на величины управляющих параметров для каждой из операций $A_n(E n_{n-1})$ — алгоритмов кодирования–декодирования, алгоритмов вычисления мотивационных оценок (невязок), принятия решений), которые в свою очередь изменяют параметры информационного потока. По виду динамики этого процесса вычисляются оценки точности для используемых алгоритмов кодирования–декодирования и формируются мотивационные сигналы для реакций системы. По ним, например, определяются условия, при которых необходимо корректировать параметры действующих алгоритмов или проводить замену «старых» алгоритмов на «новые» алгоритмы кодирования. Важно обратить внимание, что термины «старые» и «новые» алгоритмы фактически обозначают различные модельные представления распознающей системы о распознаваемом объекте и условиях работы этой системы. В результате система в ответ на входной информационный сигнал I_n формирует выходной информационный сигнал в виде

$$I_{n+1} \left[D_n, R_n, A_n(E n_{n-1}), Int_n \right].$$

Здесь используются следующие обозначения:

D_n — решения, принятые распознающей системой;

R_n — оценки уверенности, статистической достоверности принятого решения;

$A_n(E n_{n-1})$ — формализованное представление знаний в системе, включающих используемые идеи, методы, модели и алгоритмы;

$E n_{n-1}$ — описывает параметры среды, обеспечивающие активное состояние распознающего устройства при работе того или иного алгоритма. В самом простом случае, этот параметр может описывать только одну величину: потребляемые энергетические ресурсы, требуемые для работы конкретного алгоритма;

Int_n — изображение, интерпретирующее входной информационный сигнал, которое генерируется распознающей системой из кодового описания входного изображения;

I_{n+1} сигнал на выходе распознающей системы, состоящий из набора величин $D_n, R_n, A_n(E_{n-1}), Int_n$.

Уровень выходной величины E_n определяется особенностями работы исполнительных механизмов, которые запускаются на основании решения распознающей системы D_n, R_n .

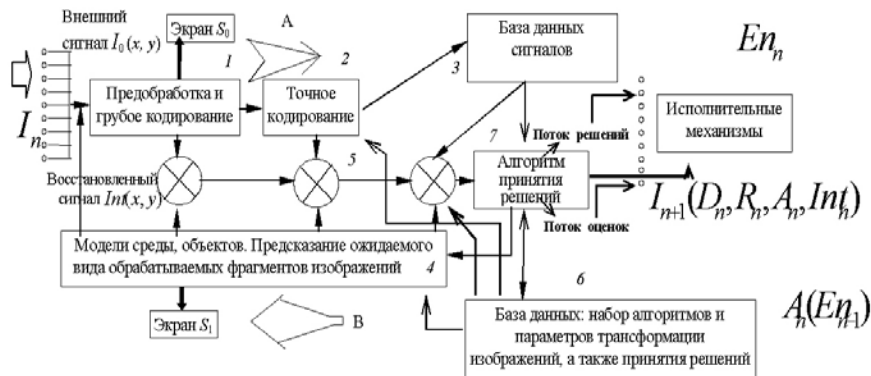


Рис. 1. Этапы трансформации потоков информационных данных и взаимодействие между различными обрабатывающими блоками в адаптивной системе принятия решений с A_n — фиксированным набором алгоритмов и уровнем E_n «энергетического» обеспечения системы (нейроподобные системы второй группы)

Используя схему адаптивной распознающей системы (рис. 1), можно дать формализованные определения для понятий, которые обычно вызывают споры при обсуждении информационных систем.

Данными об изображениях называются наборы признаков или кодовых описаний, которые хранятся в блоке 3 (см. рис. 1).

Знаниями в данной системе естественно назвать наборы алгоритмов $A_n(E_{n-1})$ в блоке 6.

Ценность входного информационного сигнала определяется по величинам невязок в блоке 5, вычисляемых из сравнения наборов кодовых опи-

саний для исходно ожидаемых системой изображений (блок 4) и реально вычисленными кодами от входного изображения (блоки 1, 2).

Процесс поиска оптимальных алгоритмов системы, на основе контроля величины на поле невязок, по-видимому, аналогичен одному из режимов, который в живых системах называют *сознанием*.

Рассмотрим, например, с помощью формализованной схемы на рис. 1 некоторые, весьма понятные с житейской точки зрения, *режимы непонимания* при общении между адаптивными распознающими системами [13].

1. Во взаимодействующих системах используются для распознавания разные модели ситуаций и связанные с ними алгоритмы кодирования–декодирования (это может быть связано с разными целями, мотивациями разных систем или просто с отсутствием подходящих алгоритмов), в результате наблюдается ситуация полного непонимания (нарушения в блоках 1, 2). В частном случае, алгоритм одной из систем может быть ориентирован на «образное» кодирование, ассоциативные связи в семантической сети «образов», в то время как другая система строит восприятие на логической модели из цепочки выводов и опирается на оценки обоснованности выводов. Это пример взаимодействия людей с образным и логическим мышлением. Известно, что такие люди очень часто не понимают друг друга [14].
2. Входные сигналы, передаваемые между адаптивными системами, представлены в некотором заранее закодированном «жаргонном» виде (несоответствия алгоритмов в блоках 1, 2, 4). Взаимодействующие распознающие системы либо имеют возможность (мотивации) произвести перекодировку к уже имеющимся у них кодовым представлениям (адаптация в блоке 7), и тогда проявляются элементы понимания, либо не имеют такой возможности (например, потому что перестройка используемых моделей требует дополнительных затрат энергии), и тогда — непонимание (отсутствие настроек в блоке 7).
3. Во входных сигналах распознающая система в первую очередь выделяет такие признаки ситуации, которые препятствуют включению тех адекватных алгоритмов обработки, которые необходимы для распознавания принимаемого потока информации (нарушение работы в блоке 7 ведет к нарушениям в блоках 1, 2, 4). Например, распознающий автомат для понимания потока информации должен затрачивать определенный уровень энергии, в то же самое время он получает сиг-

налы от соседей о малой важности этой информации. В этом случае, если приоритеты сигналов от соседей (или экономии своей энергии) оказываются выше, чем коэффициент мотивации к восприятию входного информационного потока, адаптации алгоритмов распознавания не происходит, и адаптивная система отключается от процесса анализа входных информационных сигналов и т. д.

Приведенные примеры объединяются одним известным афоризмом: *«Вы не воспринимаете нас не потому, что наши представления неверны или сложны, а потому, что наши представления не входят в ваши понятия».*

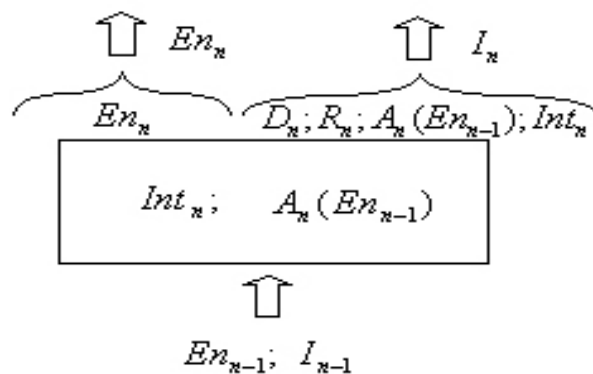


Рис. 2. Схематическое представление трансформации входных потоков $En_{n-1}; I_{n-1}$ в выходной поток $En_n; I_n$ информационных данных в адаптивной системе принятия решений с фиксированным набором алгоритмов

Один из известных выходов в ситуациях непонимания состоит в представлении внутренних кодовых описаний в такой форме, когда они понятны и другой распознающей системе, участвующей в общении. Например, вместо текстового описания архитектуры и функций базовых моделей, автор мог бы представить варианты программного продукта, в которых нажатием кнопок вызываются демонстрации соответствующих преобразований с изображениями, выбираются адекватные алгоритмы, согласуются кодовые описания, и т. д. Человеку такие действия воспринимать, конечно же, легче. Известно только, что подготовка такого «демонстрационного» решения

обычно требует дополнительных ресурсов En_{n-1} и знаний $A_n(En_{n-1})$. Часто — очень немалых.

Сопоставления с данными нейрофизиологических исследований [14–17] показывают также, что разработанная базовая модель для адаптивных распознающих систем позволяет на количественном уровне описывать различные варианты психологических режимов реагирования животных в процессе осознания ими действующих на них информационных сигналов.

Если опустить внутренние переменные в функциональной схеме адаптивного распознавателя, показанной на рис. 1, то ее упрощенное изображение можно представить в виде рис. 2. На этой схеме показаны только основные переменные, участвующие в преобразовании входного информационного потока I_n , вместе с предоставляемыми для работы системы ресурсами En_{n-1} , в выходной набор

$$\{D_n; R_n; A_n(En_n); Int_n\}$$

информационного потока и новый уровень En_n — ресурсов для работы распознающих систем на следующем слое иерархии.

Известно, что в реальной жизни адаптивные распознающие системы заинтересованы в получении точных решений D_n, R_n, Int_n для того, чтобы обеспечить доступ к необходимым ресурсам En_n при выполнении поставленных целей, а также получить возможность усовершенствования или модернизации доступного им набора знаний $A_n(En_{n-1})$.

Модели взаимодействующих распознающих систем

К **третьему типу** естественно отнести модели, описывающие различные варианты возможных взаимодействий между «элементарными» адаптивными распознающими системами.

Из всех возможных комбинаций взаимодействия распознающих систем рассмотрим иерархическую архитектуру, известную для многих *популяционных систем* (см., например, [19–20]) и показанную на рис. 3.

Процесс трансформации входных величин $(En_{n-1}), (I_{n-1})$, отвечающих за «энергетические» ресурсы и информационный сигнал, в выходные значения $En_n; I_n$ на выделенном слое взаимодействующих адаптивных систем принятия решений зависит от используемого набора A_n алгоритмов. Особенности этого процесса зависят от распределения дополнительных воздействий ΔEn_{n-1} и ΔA_{n-1} «энергетического» и «информационного»

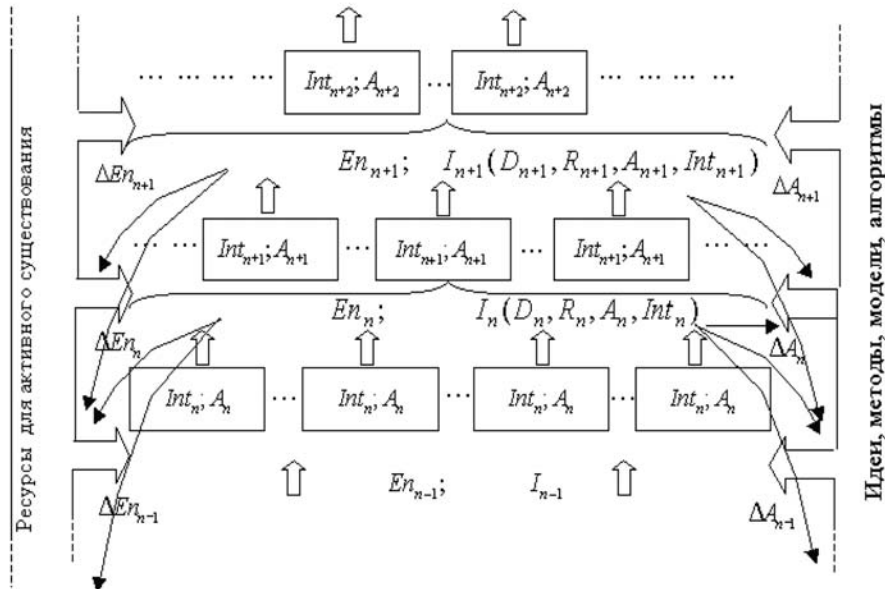


Рис. 3. Схематическое представление базовых взаимодействий между адаптивными распознающими системами при их работе в иерархической организации

Трансформация входных величин (En_{n-1}) и (I_{n-1}) , отвечающих за «энергетические» ресурсы и информационный сигнал, соответственно, в выходные значения $En_n; I_n$ на выделенном слое взаимодействующих адаптивных систем принятия решений зависит от используемого набора A_n алгоритмов. Особенности этого процесса зависят от распределения дополнительных воздействий ΔEn_{n-1} и ΔA_{n-1} «энергетического» и «информационного» характера (в общем случае, такое распределение имеет пространственно неоднородное распределение), которые обычно контролируются принятыми решениями D_{n-1}, R_{n-1} , а также решениями с вышележащих иерархических слоев.

характера (в общем случае, такое распределение будет пространственно неоднородным), которые обычно контролируются принятыми решениями D_{n-1}, R_{n-1} , а также решениями с вышележащих иерархических слоев.

На схеме показан процесс передачи с нижних уровней на верхние фоновый поток ресурсов En_n и фоновый поток информационных данных I_n . На каждом слое составляющие его системы могут взаимодействовать друг с другом. Динамику их коллективного поведения можно описывать с помощью моделей первого типа. Каждый вышележащий слой систем имеет возможность управлять притоком дополнительных ресурсов из двух дополнительных, специализированных резервуаров как для энергетических ресурсов ΔEn_n , обеспечивающих активное состояние распознающих устройств, так и для знаний ΔA_n , обеспечивающих выполнение «правильных» действий. При этом управление на каждый нижележащий слой подается в виде фильтров, которые в текущий момент разрешают нижележащему слою систем:

- дополнительную энергетическую поддержку с инструкциями ее использования и
- дополнительные знания (алгоритмы), необходимые для исполнения инструкций.

Представленная архитектура управления энергетическими ресурсами считается вполне привычной.

Гораздо больше вопросов вызывает происхождение и механизмы работы специализированного резервуара для знаний.

Достаточно хорошо известны две точки зрения на функционирование механизмов создания новых знаний (идей, методов, моделей, алгоритмов).

В соответствии с одной из них, новые знания создаются в результате комбинирования из ранее известных алгоритмических блоков. Этим могут заниматься, например, специализированные «распознающие системы». Это аналогично конструированию более совершенных алгоритмов из имеющихся в распоряжении этих систем упрощенных базовых блоков — «алгоритмических функций».

Другая точка зрения основана на представлении, что все необходимые вариации знаний или алгоритмических блоков уже существуют в некотором хранилище знаний (место этого хранилища разные авторы указывают по-разному). По мере необходимости адаптивные распознающие системы получают требуемые алгоритмы в уже готовом виде. Жизненный опыт показывает, что обычно такая доставка новых знаний происходит, если в них

возникает очень сильная потребность.

В рамках схемы на рис. 3 можно рассматривать функционирование обоих этих механизмов. Было бы очень интересно рассмотреть другие возможные варианты механизмов генерации новых знаний.

Разнообразные динамические режимы коллективной активности систем в иерархической модели (рис. 3) получаются в зависимости от приоритетов для каждой распознающей системы на такие параметры, как особенности ее взаимосвязей с соседями, $\Delta E n_n$ или ΔA_n . Описание этих режимов, определение их устойчивости и классификация наиболее характерного поведения систем в такой иерархии представляет весьма интересную область исследований.

Модели «элементарных адаптивных распознающих систем», а также модели для взаимодействующих ансамблей распознающих систем, позволяют формализовать описание гипотез, связанных с режимами реагирования в социальных системах. Например, режимы оптимизации при выполнении поставленной задачи позволяют формализовать определение особенностей психологических реакций распознающих автоматов, включая разделение на режимы, аналогичные реакциям сангвиника, флегматика, холерика и меланхолика. Можно построить интерпретацию режимов отстаивания своих интересов информационными системами за счет «латерального подавления» соседей; режимов проявления «двойной морали» в иерархической организации, и ряда других режимов, присущих социальным системам.

Выводы

Приведенные здесь разные типы базовых моделей показывают перспективность их практического применения. Известные автору примеры процессов принятия решений активными информационными системами удастся объяснить с помощью этих моделей. Было бы интересно найти ограничения предлагаемого модельного описания для объяснения экспериментальных данных о динамике нейроноподобных систем. Важно, что за каждой функциональной связью или операцией в предложенных здесь моделях стоят строго детерминированные правила преобразования сигналов. Варианты таких систем уже сейчас могут программироваться или реализовываться в аппаратуре. Конечно, следует признать, что в действительности разнообразие реакций в таких нейроноподобных системах определяется использованием достаточно обширного набора алгоритмов (возможных мо-

дельных преобразований). Однако разнообразие таких наборов модельных представлений, как известно, наблюдаются и в реальной жизни.

Для описания динамики сложных систем сейчас, как известно, используются разные виды базовых моделей. При этом результаты получаемых описаний часто вполне адекватно соответствуют опытным данным. Основная трудность состоит в переводе одних результатов на язык других модельных представлений. Формирование набора модельных представлений для нейроноподобных систем, ориентированного на использование его в практике общения между заинтересованными исследователями, составляет основную цель данной лекции.

Работа выполнялась при частичной поддержке грантов INTAS-01-0690 и АФГР РМО-10214-BNL №36943.

Литература

1. Wilson H. R., Cowan J. D. A mathematical theory of the functional dynamics of cortical and thalamic neuron tissue // *Kybernetic*. – 1973. – v. 13. – pp. 55–80.
2. Сбитнев В. И. Преобразования потока спайков в статистических нейронных ансамблях // *Биофизика*. – V1975. – Т. 20. – С. 699–702; 1976. – Т. 21. – С. 1072–1076; 1977. – Т. 22. – С. 523–528.
3. Кудряшов А. В., Яхно В. Г. Распространение областей повышенной импульсной активности в нейронной сети // *Динамика биологических систем*. – 1978. – Вып. 2. – С. 45–59.
4. Yakhno V. G., Belliustin N. S., Krasilnikova I. G., Kuznetsov S. O., Nuidel I. V., Panfilov A. I., Perminov A. O., Shadrin A. V., Shevyrev A. A. Research decision-making system operating with composite image fragments using neuron-like algorithms // *Radiophysics*. – Vol. 37, No. 8, pp. 961–986, 1994.
5. Vasiliev V. A., Romanovskii Y. M., Chernavskii D. C., Yakhno V. G. Autowave processes in kinetic systems. Spatial and temporal self-organization in physics, chemistry, biology, and medicine. – D. Reidel Publishing Company, 1987.
6. Belliustin N. S., Kuznetsov S. O., Nuidel I. V., Yakhno V. G. Neural networks with close nonlocal coupling for analyzing composite images // *Neurocomputing*. – v. 3. – 1991. – pp. 231–246.
7. Yakhno V. G. Basic models of hierarchy neuron-like systems and ways to analyse some of their complex reactions // *Optical Memory & Neural Network*. – 1995. – v. 4, No. 2. – pp. 141–155.

8. Яхно В. Г. Процессы самоорганизации в распределенных нейроподобных системах. Примеры возможных применений // *«Нейроинформатика 2001». Лекции по нейроинформатике*. – М.: МИФИ, 2001. – С. 103–141.
9. Анохин П. К. Избранные труды. – М.: Наука, 1978. – 400 с.
10. Fukushima K. Neural network model of selective attention in visual pattern recognition and associative recall // *Applied Optics*. – 1983. – v. 26, No. 23. – pp. 4985–4992; Neural network for visual pattern recognition // *Computer*. – 1988. – pp. 65–67.
11. Zverev V. A. Physical Foundation of image formation by wave fields. – IAP RAS. – 1998. – 252 pp.
12. Яхно В. Г., Нуйдель И. В., Тельных А. А., Бондаренко Б. Н., Сборщиков И. Ф., Хилько А. И. Метод адаптивного распознавания информационных образов и система для его осуществления. – Российский патент № 2160467, 1999.
13. Яхно В. Г. Модели нейроподобных систем. Динамические режимы преобразования информации. Нелинейные волны 2002 / Отв. ред. А. В. Гапонов-Грехов, В. И. Некоркин. – Нижний Новгород: ИПФ РАН, 2003, С. 90–114.
14. Грэй У. Живой мозг. – М.: Мир, 1966. – 295 с.
15. Иваницкий А. М. Физиологические основы психики // *Природа*. – 1999. – № 8. – С. 156–162.
16. Иваницкий А. М. Главная загадка природы: как на основе работы мозга возникают субъективные переживания // *Психологический журнал*. – 1999. – том. 20, № 3. – С. 93–104.
17. Архипов В. И. Воспроизведение следов долговременной памяти, зависимой от внимания // *Журнал высшей нервной деятельности*. – 1998. – том. 48, вып. 5. – С. 836–845.
18. Сергин В. Я. Психофизиологические механизмы осознания: гипотеза самоотожествления // *Журнал высшей нервной деятельности*. – 1998. – том. 48, вып. 3. – С. 558–570.
19. Манифест эволюции, 2003. URL: <http://www.sarasvati.comtv.ru>
20. Пригожин И. Р. (ред.) Человек перед лицом неопределенности. – Москва-Ижевск: Институт компьютерных исследований, 2003. – 304 с.

Владимир Григорьевич ЯХНО, ведущий научный сотрудник, доктор физико–математических наук, заведует лабораторией в Институте прикладной физики РАН (Нижний Новгород). Научные интересы связаны с исследованием процессов самоорганизации в распределенных неравновесных системах и применением автоволновых представлений для моделирования процессов обработки сенсорных сигналов, развития компьютерных алгоритмов кодирования сложных изображений, рассмотрения характерных адаптационных процессов при работе систем распознавания сложных изображений. Имеет более 120 научных публикаций (в том числе 2 монографии и 4 патента).

Н. Г. ЯРУШКИНА

Ульяновский государственный технический университет (УлГТУ),

г. Ульяновск

E-mail: jng@ulstu.ru

НЕЧЕТКИЕ НЕЙРОННЫЕ СЕТИ С ГЕНЕТИЧЕСКОЙ НАСТРОЙКОЙ

Аннотация

Нечеткие нейронные сети представляют собой реализацию систем нечеткого логического вывода методами нейронных сетей. Такие сети включают в себя слои, состоящие из специальных *И* и *ИЛИ* нейронов. Для настройки нечетких нейронных сетей используются методы генетической оптимизации. В лекции изложено краткое введение в построение и обучение нечетких нейронных сетей. Важным преимуществом таких сетей является сочетание методов машинного обучения и вербализации правил вывода.

N. G. YARUSHKINA

Ulyanovsk State Technical University (UISTU),

Ulyanovsk

E-mail: jng@ulstu.ru

FUZZY NEURON NETWORKS WITH GENETIC OPTIMIZATION

Abstract

Fuzzy neuron networks with genetic optimization are fuzzy logical inference systems. These nets include special *AND*, *OR* neurons. Genetic algorithms optimize a structure and parameters of fuzzy neuron networks. This survey lecture represents brief introduction to creation and learning of fuzzy neuron networks. Important practical advantage of fuzzy neuron networks is hybridization of machine learning methods and fuzzy logical rules representation in linguistic terms.

Введение

Исследования восприятия на искусственных моделях нейронных сетей успешно развивались, начиная с 40-х годов прошлого века (МакКаллок и Питтс, 1947; Хебб, 1949). Со знаменитой ныне статьи Л. Заде [20] началось постепенное развитие нечетких систем. В 1979 году Заде ввел теорию приближенных рассуждений [5]. Однако увлечение символьной обработкой в течение 70–80-х годов приводило к тому, что рождение и становление вычислительного интеллекта оставалось почти не замеченным. Тем не менее в 1982 году Хопфилд обеспечил математические основы понимания динамики нейронных сетей с обратными связями. В 1984 году Кохонен предложил алгоритм обучения без учителя для самоорганизующихся сетей. В 1986 году Румельхарт и МакКлелленд ввели алгоритм обратного распространения ошибки для нелинейных многослойных сетей. Книга Холланда «Адаптация в природе и искусственных системах» (1975) [18] заложила основы эволюционного моделирования и генетические алгоритмы стали методом структурного синтеза в интеллектуальных системах. Успехи нейронных сетей, нечетких систем и генетических алгоритмов привели к формированию направления, которое можно обозначить, как «*вычислительный интеллект*». В результате в 1994 году в Беркли в Калифорнии Л. Заде ввел «зонтичный» термин «*мягкие вычисления*», который можно пояснить следующей формулой [21]:

$$\begin{aligned} \text{Мягкие вычисления} = & \text{нечеткие системы} + \\ & + \text{нейронные сети} + \text{генетические алгоритмы}. \end{aligned}$$

Появление термина «мягкие вычисления» обычно объясняют тем, что мягкие или гибридные системы, такие как нечеткие нейронные сети с генетической настройкой параметров, демонстрируют взаимное усиление достоинств и погашение недостатков отдельных методов. Очевидно, что представление знаний в нейронных сетях в виде матриц весов не позволяет строить объяснение проделанного распознавания или прогнозирования, в то время как системы вывода на базе нечетких правил позволяют строить объяснения как обратные протоколы вывода (ответы на вопрос ПОЧЕМУ?). Нейронные сети обучаются с помощью универсального алгоритма, то есть трудоемкое извлечение знаний заменяется сбором достаточной по объему обучающей выборки. Для нечетких систем вывода построение включает в себя трудоемкие процессы формализации понятий, построения функций

принадлежности, формирования правил вывода. А *нечеткие нейронные сети* обучаются как нейронные сети, но строят объяснения как системы нечеткого вывода. В гибридных системах генетические алгоритмы (ГА) способны выполнить настройку функций принадлежности. Функция принадлежности задается параметризованной функцией формы, параметры которой оптимизирует ГА. ГА может оптимизировать и состав больших баз нечетких продукций или структуру нейронной сети, в результате возникают *генетические нечеткие нейронные сети*. Однако названные причины формирования гибридных систем, которые составляют основное содержание вычислительного интеллекта, являются технологическими прикладными причинами.

Существует, на наш взгляд, более важная и общая фундаментальная причина развития вычислительного интеллекта. Эту причину можно обозначить как необходимость объединить в единую систему восприятие и логическую обработку. Исторически исследование восприятия (нейронные сети, распознавание образов) развивается отдельно от классического искусственного интеллекта, моделирующего логическое мышление.

Существуют разрыв между бионическим интеллектом искусственных нейронных сетей, моделирующих восприятие, и логическим интеллектом систем логического вывода. Между тем естественный интеллект, видимо, не имеет такой резкой границы. Грубую схему «слоев естественного интеллекта» можно описать следующей формулой [1]:

$$\text{сенсорика} + \text{моторика} + \text{безусловные и условные рефлексy} + \\ + \text{врожденные программы (инстинкты)} + \text{мышление} = \text{интеллект.}$$

Возможно, эффективность человеческого интеллекта состоит в том, что все такие слои работают согласованно, помогая ему решать многие недоступные искусственным системам задачи. Вычислительный интеллект объединяет в гибридные системы искусственные модели таких слоев. Сказанное позволяет рассматривать вычислительный интеллект не как технологическое достижение, а как парадигму развития искусственного интеллекта в XXI веке. Поэтому вычислительный интеллект, сочетающий восприятие (*perception based theory*) и мышление (*computing with words*), во всяком случае, уместен в решении интеллектуальных задач.

Вычислительный интеллект в широком смысле ищет и предлагает методологическую схему, содержащую неполную, нечеткую и неточную информацию, но решающую интеллектуальные задачи. Вычислительный интел-

лект в узком смысле ищет схемы гибридизации интеллектуальных технологий, обладающие преимуществами перед их автономным использованием. Общая задача объединения моделей восприятия и логической обработки на уровне структуры должна проявляться в том, что появляются схемы, работающие с вещественными числами (перцепция) и дискретными сигналами истины (логика). Как мы увидим ниже, схемы гибридных систем включают в себя не только классические нейроны, но и *И*, *ИЛИ*–нейроны.

Наряду с термином «мягкие вычисления» используется и другой интегрирующий термин — вычислительный интеллект. В рамках вычислительного интеллекта очень многие исследователи разрабатывают системы, объединяющие составные части мягких вычислений. В настоящее время существуют в основном нейро-нечеткие системы. Однако растет число нечетко-генетических, нейро-генетических и нейро-нечетко-генетических систем. В данной лекции анализируется эффективное использование и полезные комбинации составных частей мягких вычислений (МВ): нечеткой логики (НЛ) для контроля параметров генетических алгоритмов (ГА) и нейронных сетей (НС), ГА для формирования НС (топологии и весов), НС для настройки нечетких контроллеров.

Нечеткие системы

Нечеткие множества

Нечеткие системы опираются на лингвистическое описание отношений между переменными процессов. В отличие от базовых значений нечеткие значения представляют собой слова повседневного языка, то есть они могут описывать человеческий опыт и знания о процессах.

Пример 1. Определение множеств нечетких значений.

Объектная переменная — дата.

Множество нечетких значений:

$$F_1 = \{\text{начало недели, середина недели, конец недели}\}.$$

Функции принадлежности: Определение и смысл

В любой ситуации признак объекта проблемной области имеет одно и только одно четкое значение из согласованного множества базовых и одно

ТАБЛИЦА 1. Функции принадлежности для понятия «дата»

b	f_1	f_2	f_3	$\mu(b, f)$
Пн.	1.0	0.0	0.0	1.0
Вт.	0.6	0.6	0.0	1.2
Ср.	0.0	1.0	0.0	1.0
Чт.	0.0	0.6	0.3	0.9
Пт.	0.0	0.0	0.7	0.7
Сб.	0.0	0.0	0.9	0.9
Вс.	0.0	0.0	1.0	1.0

или более чем одно нечеткое значение из соответствующего множества нечетких значений. Выберем для объектной переменной «дата» множество базовых значений:

$$B = \{\text{Пн., Вт., Ср., Чт., Пт., Сб., Вс.}\}$$

и множество нечетких значений:

$$F = \{\text{начало недели, середина недели, конец недели}\}.$$

Сущности «дата» и четкому значению «среда», соответствует нечеткое — середина недели. «Среду» с трудом можно отнести на «начало недели» или «конец недели».

Интуитивное отношение между базовым и нечетким значением объектной переменной выражают количественно с помощью *функции принадлежности*, обозначаемой греческой буквой μ .

Функция $\mu(b, f)$ отображает базовое значение b и нечеткое значение f в интервал $[0; 1]$. По определению, $0 \leq \mu(b, f) \leq 1$ для b и f . Если $\mu(b_0, f_0) = 0$ для пары значений (b_0, f_0) , то четкое значение b_0 не принадлежит нечеткому f_0 и не является членом f_0 . Если $\mu(b_0, f_0) = 1$, то $b_0 \in f_0$. Частичная или неопределенная принадлежность выражается значением функции, лежащими в интервале между 0 и 1.

Рассмотрим таблицу значений $\mu(b, f)$ для объектной переменной «дата» (табл. 1). Сумма чисел по горизонтали часто, но не обязательно близка к 1.0. Реально это означает, что нечеткие значения не покрывают формально спектр базовых значений.

Нечеткие числа

Часто признаки объектов проблемной области с одной стороны описываются четкими числами как базовыми значениями, а с другой — нечеткими, такими как ~ 1 , ~ 2.5 , «грубо 6.0». Соответствующие функции принадлежности строятся как кривые Гаусса или треугольные функции.

Нечеткие интервалы

Рассмотрим объектную переменную с базовыми упорядоченными действительными числами и нечеткими значениями, заданными в виде нечетких интервалов, например: грубо между 1 и 5, приблизительно в интервале от 100 до 120, вероятно, между 5 и 6. Функция принадлежности обычно строится как трапецевидная по форме.

Операции с нечеткими множествами

Дополнительное множество НЕ. Рассмотрим определение операции дополнения в классической теории множеств. Пусть M — нечеткое множество элементов из множества базовых значений B , тогда дополнительное множество $\neg M$, определяется как множество из тех базовых элементов B , которые не принадлежат множеству

$$M : \neg M = \{b; b \in B, b \notin M\}.$$

Определим дополнение в теории нечетких множеств. Рассмотрим нечеткое множество $M = \{(b, \mu); b \in B, 0 \leq \mu \leq 1\}$, дополнительное нечеткому множеству $\neg M$, определяемому как

$$\neg M = \{(b, (1 - \mu)), b \in B, 0 \leq \mu \leq 1\}.$$

Пересечение И. Рассмотрим определение операции пересечения в традиционной теории множеств. Пусть заданы два множества M_1 и M_2 элементов из множества базовых значений B , тогда пересечение множеств $M = M_1 \cap M_2$ определяет множество, элементы которого из B принадлежат к обоим множествам M_1 и M_2 одновременно

$$M = M_1 \cap M_2 = \{b; b \in M_1 \text{ и } b \in M_2\}.$$

Определим операцию пересечения в теории нечетких множеств. Пусть заданы два нечетких множества

$$M_1 = \{(b, \mu_1) \mid b \in B, 0 \leq \mu_1 \leq 1\},$$

$$M_2 = \{(b, \mu_2) \mid b \in B, 0 \leq \mu_2 \leq 1\}.$$

Пересечение нечетких множеств определяется как

$$M = M_1 \cap M_2 = \{(b, \min(\mu_1, \mu_2))\}.$$

Значение $b_0 \in M_1$ и $b_0 \in M_2$ полностью содержится в $M_1 \cap M_2$:

$$\mu_1(b_0) = 1; \quad \mu_2(b_0) = 1 \Rightarrow \min(\mu_1, \mu_2) = 1.$$

Значение b_0 , частично содержащееся в M_1 или в M_2 , только частично содержится в $M_1 \cap M_2$:

$$0 \leq \mu_1(b_0) \leq 1; \quad 0 \leq \mu_2(b_0) < 1 \Rightarrow 0 \leq \min(\mu_1, \mu_2) \leq 1.$$

Если значение $b_0 \notin M_1$ и $b_0 \notin M_2$, то оно не содержится в

$$M_1 \cap M_2 : \mu_1(b_0) = 0, \quad \mu_2(b_0) = 0 \Rightarrow \min(\mu_1, \mu_2) = 0.$$

Объединение ИЛИ. Рассмотрим определение операции объединения в традиционной теории множеств. Пусть даны два множества M_1 и M_2 элементов базового множества B , тогда объединением множеств $M = M_1 \cup M_2$ называется множество, содержащее такие элементы B , которые принадлежат множествам M_1 или M_2 или обоим одновременно:

$$M = M_1 \cup M_2 \Rightarrow \{b; \quad b \in M_1 \text{ или } b \in M_2\}.$$

Определим объединение в теории нечетких множеств. Пусть даны два нечетких множества

$$M_1 = \{(b, \mu_1) \mid b \in B, \quad 0 \leq \mu_1 \leq 1\},$$

$$M_2 = \{(b, \mu_2) \mid b \in B, \quad 0 \leq \mu_2 \leq 1\},$$

тогда объединением множеств $M = M_1 \cup M_2$ называется

$$M = \{(b, \max(\mu_1, \mu_2)) \mid b \in B, \quad 0 \leq \mu_1 \leq 1, \quad 0 \leq \mu_2 \leq 1\}.$$

Для заданного множества $F = \{f\}$ нечетких значений объектной переменной A , мы можем добавить новые термины, используя связку *ИЛИ*. Функция принадлежности нового значения объектной переменной A определяется как:

$$\begin{aligned}\mu(b, f_1 \text{ ИЛИ } f_2) &= \max[\mu(b, f_1), \mu(b, f_2)], \quad \forall f_1 \text{ ИЛИ } f_2 \in A. \\ M(f_1 \text{ ИЛИ } f_2) &= \{b \mid \mu(b, f_1 \text{ ИЛИ } f_2)\} = \\ &= \{b \mid \max(\mu(b, f_1), \mu(b, f_2))\} = M(f_1) \cup M(f_2).\end{aligned}$$

Следовательно, нечеткое множество, принадлежащее к нечетким значениям f_1 *ИЛИ* f_2 , это нечеткая дизъюнкция множеств, принадлежащих f_1 и f_2 .

Обобщенные определения пересечения и объединения нечетких множеств. Выше мы определили операции пересечения и объединения двух нечетких множеств следующим образом:

- пересечение:

$$M_1 \cap M_2 = \{b \mid \min(\mu_1, \mu_2)\},$$

- объединение:

$$M_1 \cup M_2 = \{b \mid \max(\mu_1, \mu_2)\}.$$

Фактически эти определения — специальные случаи общего, использующего понятия T -норм и S -конорм [19]. Эти выражения обозначают функции со специальными свойствами, которые будут детально обсуждаться далее. T -нормы и S -конормы обладают двумя важными свойствами. $T(x, y)$ и $S(x, y)$ — это функции двух действительных переменных x и y , определенных на интервалах $[0, 1] \times [0, 1]$. Значения функций $T(x, y)$ и $S(x, y)$ также находятся в интервале $0 \leq T(x, y) \leq 1$ и $0 \leq S(x, y) \leq 1$. Используя любые T -нормы $T(x, y)$ и S -конормы $S(x, y)$, обобщим определения операций над нечеткими множествами:

- пересечение:

$$M_1 \cap M_2 = \{b \mid T(\mu_1, \mu_2)\},$$

Таблица 2. T -нормы

$T(x, y)$	Имя функции
$y = \begin{cases} 0 & , \text{ если } \max(x, y) > 1 \\ \min(x, y) & , \text{ если } \max(x, y) \leq 1 \end{cases}$	ограниченное произведение
$\max(0, x + y - 1)$	усиленная разность
$\frac{x * y}{1 + (1-x) * (1-y)}$	произведение Эйнштейна
$x * y$	алгебраическое произведение
$\frac{x * y}{1 - (1-x) * (1-y)}$	произведение Гамахера
$\min(x, y)$	минимум

- объединение:

$$M_1 \cup M_2 = \{b \mid S(\mu_1, \mu_2)\}.$$

- операция И:

$$\mu(b, f_1 \text{ И } f_2) = T(\mu(b, f_1), \mu(b, f_2)).$$

- операция ИЛИ:

$$\mu(b, f_1 \text{ ИЛИ } f_2) = S(\mu(b, f_1), \mu(b, f_2)).$$

Далее будет показано что для $0 \leq x \leq 1$ и $0 \leq y \leq 1$ функция $\min(x, y)$ соответствует T -норме и функция $\max(x, y)$ — S -конорме. В табл. 2 и 3 перечислены некоторые широко известные T -нормы и S -конормы [3].

Общие свойства T -норм и S -конорм

T -нормы и S -конормы имеют следующие общие свойства:

- интервалы определения:

$$0 \leq x \leq 1, \quad 0 \leq y \leq 1;$$

- интервал значений:

$$0 \leq T(x, y) \leq 1, \quad 0 \leq S(x, y) \leq 1;$$

ТАБЛИЦА 3. S -конормы

$S(x, y)$	Имя функции
$\max(x, y)$	максимум
$1 - (1 - x) * \frac{(1-y)}{(1-x*y)}$	сумма Гамахера
$1 - (1 - x) * (1 - y)$	алгебраическая сумма
$1 - (1 - x) * \frac{(1-y)}{(1+x*y)}$	сумма Эйнштейна
$\min(1, x + y)$	ограниченная сумма
$y = \begin{cases} 1 & , \text{ если } \min(x, y) > 0 \\ \max(x, y) & , \text{ если } \min(x, y) = 0 \end{cases}$	усиленная сумма

- коммутативность:

$$T(x, y) = T(y, x), \quad S(x, y) = S(y, x);$$

- ассоциативность:

$$T[T(x, y), z] = T[x, T(y, z)],$$

$$S[S(x, y), z] = S[x, S(y, z)];$$

- монотонность:

$$\text{Если } x = u \text{ и } y = v, \text{ то } T(x, y) = T(u, v), \quad S(x, y) = S(u, v);$$

- специальные значения:

$$T(0, y) = 0,$$

$$S(0, y) = y,$$

$$T(1, y) = y,$$

$$S(1, y) = 1.$$

Каждая функция со свойствами T называется T -нормой или триангулярной нормой, а со свойствами S — S -конормой или триангулярной ко-нормой. Каждая пара триангулярных норм и конорм, которые подчиняются обобщенным отношениям Де-Моргана, называется парой относительных триангулярных норм. Все функции принадлежности (как для I , так и для $ИЛИ$ -связок нечетких значений) выбираются такими же далекими друг

от друга, как пары относительных триангулярных T -норм и S -конорм. T -нормы выбираются для *И*, S -конормы — для *ИЛИ*. Функции называют триангулярными потому, что соотношения

$$T(x, y) = T(y, x) \text{ и } S(x, y) = S(y, x)$$

выполняются только в треугольной области $0 \leq x \leq 1, 0 \leq y \leq x$.

Нечеткие отношения

Общие определения. В практических приложениях бывает полезно соединить два стандартных нечетких множества в одно комбинированное нечеткое множество, называемое *декартовым произведением*. Подобная концепция широко используется в традиционной теории множеств.

Предположим имеется два объекта проблемной области A и B с двумя множествами базовых значений $\{a\}$ и $\{b\}$, нечетких значений $\{f\}$ и $\{g\}$, и функциями принадлежности $p(a, f)$ и $q(b, g)$, соответственно. Сущности A и B могут быть как различными, так и фактически идентичными.

Рассмотрим два нечетких множества $P(f) = \{a | p(a, f)\}$ объектной переменной A и $Q(g) = \{b | q(b, g)\}$ переменной B . Операция $n(x, y)$ будет нормой или конормой, соответственно. Определим произведение множеств $P(f)$ и $Q(g)$ следующим образом:

$$Z(f, g) = P(f) \otimes Q(g) = \{(a, b) | n[p(a, f), q(b, g)]\}, \\ \{a, p(a, f)\} \in P(f), \{b, q(b, g)\} \in Q(g).$$

Другими словами, Z — это множество всех триплетов a, b, n , где a и b базовые значения A и B , f и g нечеткие значения A и B , а $n[p(a, f), q(b, g)] = \mu[(a, b), (f, g)]$ определяет комбинированную принадлежность базовых значений (a, b) нечетким значениям (f, g) .

Если пара (f, g) нечетких значений связывается связкой *И*, т. е. $(f, g) = f$ И g , то требуется выбрать T -норму $n(x, y)$. Если $(f, g) = f$ ИЛИ g , то необходимо связывать f и g функцией $n(x, y)$, которая должна быть S -конормой.

Двухместные нечеткие множества. Нечеткое бинарное отношение. Пусть X и Y — непустые множества. Нечеткое отношение R — это нечеткое подмножество $X \otimes Y$. Другими словами, $R \in F(X \otimes Y)$. Если $X = Y$,

то говорят, что R — это бинарное нечеткое отношение на X . Если R — это бинарное нечеткое отношение на $X \otimes X$, тогда $R(x, y)$ можно интерпретировать как степени принадлежности упорядоченных пар (x, y) .

Пример 2. Бинарное нечеткое отношение.

Бинарное нечеткое отношение на множестве $U = \{1, 2, 3\}$ назовем «приблизительно равно»

$$\begin{aligned} R(1, 1) &= R(2, 2) = R(3, 3) = 1, \\ R(1, 2) &= R(2, 1) = R(2, 3) = R(3, 2) = 0.6, \\ R(1, 3) &= R(3, 1) = 0.2. \end{aligned}$$

Значимость нечетких отношений обусловлена тем, что они описывают взаимодействие между переменными.

Операции над нечеткими отношениями. Пусть R и S — два бинарных нечетких отношения на $X \times Y$. Пересечение R и S определяется следующим образом:

$$(R \cap S)(x, y) = \min(R(x, y), S(x, y)),$$

где $R, S : X \times Y \rightarrow [0, 1]$.

Объединение R и S определяется как

$$(R \cup S)(x, y) = \max(R(x, y), S(x, y)).$$

Проекция двухместных функций принадлежности. Если в двухместных функциях принадлежности $[(a, b), (f, g)] = n[p(a, f), q(b, g)]$, переменные b и g зафиксировать на уровне значений $b = b_0$ и $g = g_0$ объектной переменной B , то функция $\nu(a, f) = \mu[(a, b_0), (f, g_0)]$ называется *проекцией* $\mu[(a, b), (f, g)]$ на значения (b_0, g_0) переменной B . Функция $\nu(a, f)$ зависит только от двух переменных a, f и представляет собой принадлежность базовых значений объектной переменной A к нечетким, причем имеется столько проекций $\nu(a, f)$ двухместных функций принадлежности $\mu[(a, b), (f, g)]$, сколько возможно комбинаций (b_0, g_0) . Проекцией нечеткого отношения является тоже нечеткое множество. Для того чтобы задать проекции, необходимо указать их функции принадлежности. Нечеткие проекции [15] нечеткого отношения определяются следующим образом:

$$\begin{aligned} \Pi_X R(x, y) &= \sup R(x, y) \mid y \in Y, \\ \Pi_Y R(x, y) &= \sup R(x, y) \mid x \in X. \end{aligned}$$

Определение проекций позволит в дальнейшем определить композицию нечетких отношений и композицию нечеткого множества с нечетким отношением.

Графически проекция нечеткого отношения приведена на рис. 1.

Нечеткое множество A нормально тогда и только тогда, когда $\sup_x \mu_A(x) =$

1. Если A и B — нормальные множества, то

$$\begin{aligned}\Pi_y(A \otimes B) &= B, \\ \Pi_x(A \otimes B) &= A, \\ \Pi_x(x) &= \sup\{(A \otimes B)(x, y) \mid y \in Y\} = \\ &= \sup\{A(x) \cap B(y) \mid y \in Y\} = \\ &= \min\{A(x), \sup\{B(y) \mid y \in Y\}\} = \\ &= \min\{A(x), 1\} = A(x).\end{aligned}$$

Функции нечетких переменных

Функции с одной независимой переменной. Для описания проблемной области используются либо четкие числа x из некоторого множества четких $\{x\}$, либо нечеткие числа f из некоторого множества нечетких значений $\{f\}$. Представления связаны друг с другом функцией принадлежности $\mu(x, f)$ соответствующего нечеткого множества $X(f) = \{(x, \mu(x, f))\}$.

Функция $w = \text{fun}(x)$ может быть задана алгоритмом, который каждому четкому числу x из $\{x\}$ ставит в соответствие только одно четкое число w из определенного множества $\{w\}$ действительных чисел. Функция $w = \text{fun}(x)$ — это «образ» x после выполнения преобразования $\text{fun}(x) : x \rightarrow w$. Существует ли подобное определение функции для нечетких чисел? Пусть f_0 — некоторое нечеткое число (например, $f_0 = \text{«около 3.4»}$), а $X(f_0) = \{(x, \mu(x, f_0))\}$ — соответствующее нечеткое множество. Тогда найдем нечеткое множество W .

- Пусть (x_0, μ_0) — любой элемент нечеткого множества $X(f_0)$, тогда четкое значение w_0 определяется по формуле $w_0 = \text{fun}(x_0)$.
- Принадлежность $r(w_0)$ вычисляется следующим образом:
 - соберем все $(x, \mu) \in X(f_0)$ с $\text{fun}(x) = w_0$, пусть этими элементами будут $(x_0, \mu_0), (x_1, \mu_1), \dots, (x_N, \mu_N)$,
 - вычислим $r_0 = S(\mu_0, \mu_1, \dots, \mu_N)$ по некоторой ранее согласованной S -конорме,

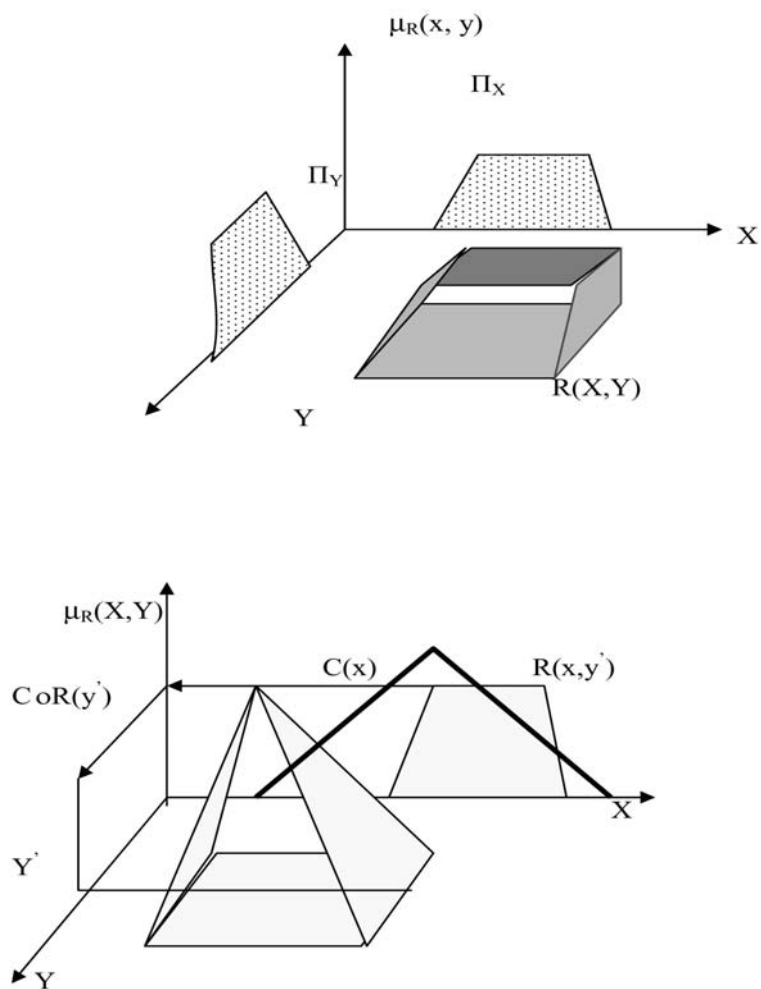


Рис. 1. Проекция и композиция нечеткого отношения

- В результате $W = \{(w, r)\}$ будет множеством всех пар $(w, r) = (w_0, r_0)$, вычисленных для всех возможных элементов $(x_0, m_0) \in X(f_0)$.

Рассмотрим W как нечеткое множество нечетких чисел p_0 , определенных на множестве $\{w\}$ базовых значений w . Нечеткое число p_0 рассматривается как «образ» $p_0 = \text{fun}(f_0)$. Так W получает значения $W(\text{fun}(f_0))$. В общем случае для любого нечеткого числа из $\{f\}$ справедливо:

$$x \rightarrow w = \text{fun}(x), f \rightarrow p = \text{fun}(f), X(f) \rightarrow W(p).$$

Если нечеткое множество W получено из нечеткого множества X , то говорят, что существует отношение между X и W или $X \rightarrow W$ [13].

Композиция нечеткого множества и отношения. Выразим процедуру вычисления функции от нечеткого значения с помощью формул операций $\sup - \min$ [15]. Композиция $\sup - \min$ нечеткого множества $C \in F(X)$ и нечеткого отношения $R \in F(X \otimes Y)$ определяется как

$$(C \circ R)(y) = \sup_{x \in X} \min\{C(x), R(x, y)\}$$

для всех $y \in Y$.

Композицию нечеткого множества C и нечеткого отношения R можно представить как тень отношения R на нечеткое множество C (рис. 1).

Пусть A и B — нечеткие числа, а $R \subset A \otimes B$ будет нечетким отношением. Очевидны следующие свойства композиции

$$\begin{aligned} A \circ R &= A \circ (A \otimes B) = A, \\ B \circ R &= B \circ (A \otimes B) = B, \end{aligned}$$

Определим $\sup - \min$ композицию нечетких отношений. Пусть $R \in (X \otimes Y)$ и $S \in F(Y \otimes Z)$. Тогда $\sup - \min$ композиция R и S , обозначаемая $R \circ S$, определяется как

$$(R \circ S)(u, w) = \sup_{x \in X} \{R(u, v), S(v, w)\}.$$

Очевидно, что $R \circ S$ — бинарное нечеткое отношение на X и Z (рис. 1).

Нечеткие системы

Определение лингвистической переменной. Системы нечетких продукций строятся на основе понятия лингвистической переменной, которой называют пятерку объектов:

$$(x, T(x), U, G, M),$$

где x — собственное имя переменной; $T(x)$ — терминальное множество, т. е. набор значений переменной (нечетких меток); U — множество объектов (или универсум); G — синтаксические правила употребления; M — семантические правила употребления [8]. Примерами лингвистических переменных служат: количество денежных средств в финансовых решениях; температура, давление воздуха в технических процессах; средняя отметка в процессах образования; плотность населения в социальных процессах. Количество денежных средств может быть выражено как: большое, малое, 1000 рублей, значительное, незначительное; динамика давления воздуха — выросло, осталось прежним, максимально низкое, 1023 Па; отметка — отлично, хорошо, удовлетворительно, неудовлетворительно; численность населения — значительная, средняя, малая.

Схема приближенного логического вывода. Задача интерполяции. Рассмотрим логический вывод по нечетким продукциям на основе теории приближенных вычислений. Пусть

$$\begin{aligned} y &= f(x) && \text{— предпосылка,} \\ x &= x' && \text{— факт,} \\ y' &= f(x') && \text{— заключение.} \end{aligned}$$

Предпосылку рассуждения $f(x)$ зададим с помощью N значений функции:

$$\begin{aligned} R_1 &: \text{если } x_1, \text{ то } y_1; \\ R_2 &: \text{если } x_2, \text{ то } y_2; \\ &\dots \\ R_n &: \text{если } x_n, \text{ то } y_n. \end{aligned}$$

Текущее значение $x = x'$.

Необходимо найти значение y' , зависящее от x' . Нечеткая задача интерполяции формулируется с использованием зависимостей нечетких множеств:

$$\begin{aligned} R_1 : & \text{ если } A_1(x_1), \text{ то } B_1(y_1); \\ R_2 : & \text{ если } A_2(x_2), \text{ то } B_2(y_2); \\ & \dots \\ R_n : & \text{ если } A_n(x_n), \text{ то } B_n(y_n); \end{aligned}$$

и нечеткого значения входной переменной — $C(x')$.

Необходимо найти значение $D(y')$, входящего в заключение.

Обобщив нечеткую зависимость переменных x и y , получим продукцию: если $A(x)$, то $B(y)$. Цель нечеткого логического вывода — определить значения выходной переменной $B'(y)$ от входной переменной $C(x')$. Результирующее нечеткое множество $B' = C \circ (A \rightarrow B)$ — композиция нечеткого множества C и импликации нечетких множеств $A \rightarrow B$. Нечеткое множество B' заключения вычисляется как композиция двух нечетких множеств — факта и импликации. Композиция факта и импликации представляет собой проекцию импликации на факт.

$$A' \circ (A \rightarrow B) = \sup \min(A', A \rightarrow B); \quad y \in Y.$$

На основании задачи нечеткой интерполяции строят оболочку экспертной системы, осуществляющей вывод на словах [15].

Основные правила умозаключений. Основными правилами умозаключений являются два силлогизма: *modus ponens* (МР) и *modus tollens* (МТ). Рассмотрим изменение основных правил вывода в случае, если предпосылка и заключение выражаются в терминах нечетких множеств.

Нечеткий modus ponens (Fuzzy МР):

$$\text{если } A(x) \rightarrow B(x), A'(x), \text{ то } B'(x).$$

Нечеткий modus tollens (FMT):

$$\text{если } A(x) \rightarrow B(x), \neg B'(x), \text{ то } \neg A'(x).$$

При использовании систем нечеткого вывода особое значение имеет *дефазификация*, т.е. определение четкого представляющего элемента по

нечеткому множеству. Самым распространенным методом дефазификации является определение центра тяжести фигуры, образуемой функцией принадлежности и осью абсцисс (центроидный метод). С его помощью вычисляют центр тяжести фигуры:

$$\frac{\int_w z * C(z) dz}{\int_w C(z) dz} \Rightarrow \frac{\sum_{i=1}^n z_i * C(z_i)}{\sum_{i=1}^n C(z_i)}.$$

На практике используют метод выбора максимума:

$$Z = \min\{Z | C(Z) = \max(C(Z))\}.$$

В некоторых системах применяют также методы минимального и среднего максимума. Правильно выбранная дефазификация существенно влияет на эффективность работы системы нечеткого вывода.

Универсальная аппроксимация с помощью систем нечеткого вывода. Свойство универсальности применения систем нечеткого вывода доказано следующей фундаментальной теоремой. В 1992 году *У. Ванг* показал, что справедливо следующее утверждение: если нечеткая импликация основана на использовании операции минимум, функция принадлежности задается гауссовым распределением:

$$\mu(t) = \exp(-1/2 * ((a - t^2)/b))$$

и используется центроидный метод дефазификации, то система нечеткого вывода является универсальным аппроксиматором.

В 1995 году *К. Кастро* доказал справедливость следующей теоремы: если импликация основана на использовании операции произведения по Ларсену, а функции принадлежности — треугольные, то при использовании центроидной дефазификации нечеткий контроллер является универсальным аппроксиматором [11].

Схемы нечеткого вывода Рассмотрим различные схемы вывода на базе нечетких правил. Пусть входная переменная x является каким-либо нечетким множеством A , и y — нечетким множеством B , тогда выходная переменная z является нечетким множеством C . Если задан факт $x = x_0, y = y_0$, то необходимо найти $z = z_0$.

Схема вывода на базе нечетких правил сводится к решению следующей задачи:

База правил

R_1 : если $(X \text{ is } A_1)$ И $(Y \text{ is } B_1)$ тогда $(Z \text{ is } C_1)$,
в противном случае

R_2 : если $(X \text{ is } A_2)$ И $(Y \text{ is } B_2)$ тогда $(Z \text{ is } C_2)$,
в противном случае

...

R_n : если $(X \text{ is } A_n)$ И $(Y \text{ is } B_n)$ тогда $(Z \text{ is } C_n)$.

Факт: $x = x_0, y = y_0$.

Следствие $z = ?$

Для реализации интеллектуальной системы логического вывода по базе нечетких правил, необходимо определить функции:

- представления в системе нечетких понятий (функций принадлежности);
- вычисления логических выражений условных частей правил с логическими связками *И*, *ИЛИ*;
- вычисления импликации;
- усреднения результата, получаемого по разным правилам путем композиции.

Каждая из четырех перечисленных выше функций задает вариативность схемы нечеткого вывода.

Интеллектуальная система осуществляет логический вывод по базе нечетких правил по этапам.

- Фазификация фактических данных, т. е. точное значение x_0 интерпретируется как нечеткая точка.
- Композиция входной переменной и условной части правила: $x_0 \circ A_i, y_0 \circ B_i$, т. е. вычисляется уровень пригодности правила к ситуации. Если факт задан нечеткой точкой, то композиция сводится к выявлению соответствующей степени принадлежности.
- Вычисление нечеткой импликации.

$$((x_0 \circ A_i) \cap (y_0 \circ B_i)) \rightarrow C_i$$

для $\forall R$. Результатом выполнения для всех правил являются N нечетких значений для выхода Z .

- Агрегация среднего значения, т. е. построение нечеткого значения выхода по результатам предыдущих этапов.

$$C = \cup_{i=1}^n C_i.$$

- Дефазификация, т. е. выбор представляющего элемента по агрегированному нечеткому понятию.

Схемы нечеткого вывода у разных авторов уточняются тем не менее только до оператора нечеткой импликации. Распространены пять схем нечеткого вывода [15]. В *схеме 1 нечеткого вывода Мамдани* импликация моделируется минимумом, а агрегация — максимумом (рис. 2).

Схема 2 разработана Цукамото для монотонных функций принадлежности.

На рис. 2 использованы следующие обозначения: L'_1, L'_2 — уровни достоверности применения правил, z_1, z_2 — значения выходной переменной по первому и второму правилу,

$$\begin{aligned} z'_1 &= C_1^{-1}(L_1), \\ z'_2 &= C_2^{-1}(L_2), \\ z &= \frac{L'_1 z_1 + L'_2 z_2}{L'_1 + L'_2}. \end{aligned}$$

Ввиду монотонности функций вычисления выходной переменной сводят к усреднению значений, полученных по разным правилам.

Схема 3 по Суджено ограничивает правые части правил вывода линейным случаем:

$$\begin{aligned} \text{Если } (x \text{ is } A_1) \text{ И } (y \text{ is } B_1) \text{ тогда } (z &= a_1 * x + a_2 * y), \\ \text{Если } (x \text{ is } A_2) \text{ И } (y \text{ is } B_2) \text{ тогда } (z &= b_1 * x + b_2 * y). \end{aligned}$$

Схема 4 по Ларсену выполняет импликацию с помощью произведения.

Схема 5 — это упрощенная схема нечеткого вывода:

$$\text{Если } (x \text{ is } A_i) \text{ И } (y \text{ is } B_i) \text{ тогда } (z = Z_i),$$

где Z_i — четкое значение.

Так как правые части правил в упрощенной схеме логического вывода задаются четко, то в результате вывода получается дискретное множество решений, для каждого элемента которого задана определенная степень уверенности. В качестве значений выходной переменной выбирается значение с максимальной уверенностью.

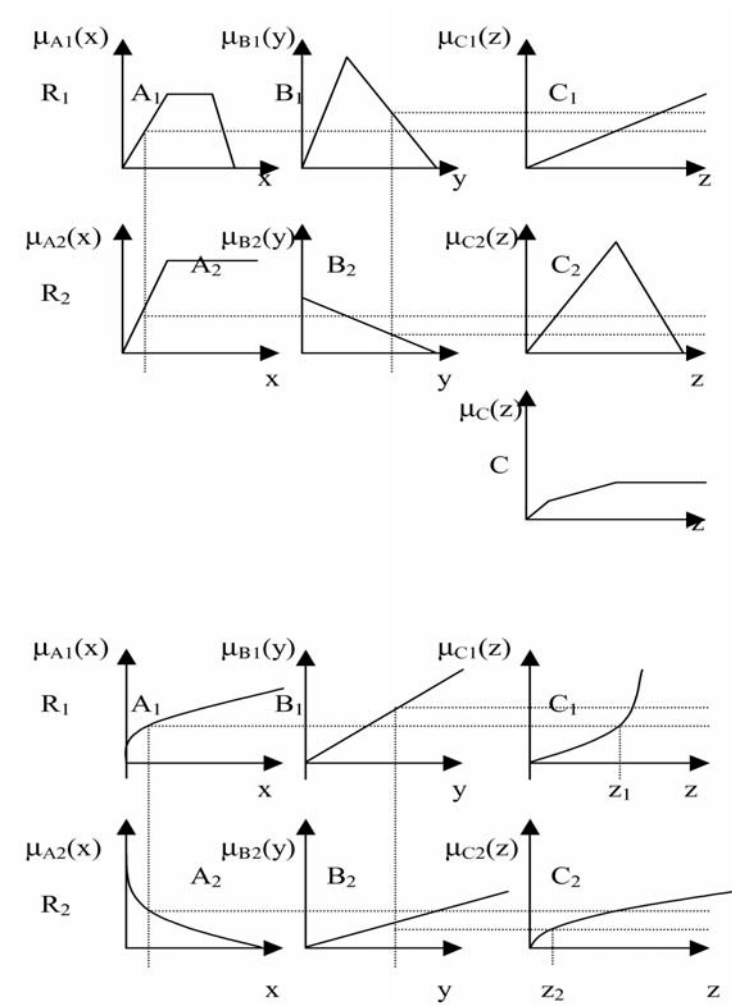


Рис. 2. Схемы нечеткого вывода по Мамдани и Цукamoto

Генетические вычисления

Генетические алгоритмы

Генетические алгоритмы — это адаптивные методы поиска, которые в последнее время часто используются для решения задач функциональной и структурной оптимизации. Генетическими они называются потому, что строятся на принципах эволюции биологических организмов Ч. Дарвина. Популяции развиваются в течение нескольких поколений, подчиняясь законам естественного отбора, т.е. принципу «*выживает наиболее приспособленный*» (survival of the fittest). Генетические алгоритмы (ГА) способны «развивать» решения реальных задач, если те закодированы соответствующим образом. ГА могут использоваться для проектирования структуры механизмов, поиска оптимальной формы детали, раскроя ткани. Например, израильская компания Schema на основе ГА разработала программный продукт Channeling для оптимизации работы сотовой связи путем выбора оптимальной частоты, на которой будет вестись разговор.

Основные принципы ГА были сформулированы Дж. Голландом. ГА приблизительно моделирует процессы, происходящие в популяциях. В природе особи в популяции конкурируют друг с другом за различные ресурсы, за возможность породить потомство. Наиболее приспособленные к окружающим условиям особи будут иметь больше шансов воспроизвести потомков. Следовательно, гены высоко адаптированных особей распространяются в популяции и количество их потомков возрастает в каждом последующем поколении, т.е. вид развивается за счет приспособления к среде обитания. ГА по аналогии с эволюционным механизмом работают с популяцией, каждая из хромосом которой представляет возможное решение данной задачи. Каждая хромосома оценивается мерой ее «*приспособленности*» (fitness function), которую называют также *функцией оптимальности*. В частности, мерой приспособленности хромосом, кодирующих монтажную схему радиотехнического изделия, может служить длина проводников в задаче размещения элементов. Наиболее приспособленные особи получают возможность «воспроизвести» потомство с помощью «*перекрестного скрещивания*» (recombination) с другими особями популяции. В результате появляются новые особи, сочетающие в себе характеристики, наследуемые ими от родителей. Наименее приспособленные особи постепенно исчезают из популяции в процессе эволюции. Новое поколение обладает лучшими характеристиками по сравнению с предыдущим. Скрещивание наиболее

приспособленных особей приводит к тому, что эволюция отыскивает перспективные решения в широком пространстве поиска. В конечном итоге, популяция сходится к оптимальному решению задачи. Существует много способов реализации идеи биологической эволюции в рамках ГА. Традиционным считается ГА, представленный ниже с использованием псевдокода. Большими буквами записаны управляющие операторы алгоритма, а курсивом — генетические операторы [4, 10]:

НАЧАЛО — генетический алгоритм

Создать начальную популяцию

Оценить приспособленность каждой особи

останов := ЛОЖЬ

ПОКА НЕ останов ВЫПОЛНЯТЬ

НАЧАЛО — создать популяцию нового поколения

ПОВТОРИТЬ (размер популяции / 2) **РАЗ**

НАЧАЛО — цикл воспроизводства

Выбрать две особи с высокой приспособленностью

из предыдущего поколения для скрещивания

Скрестить выбранные особи и получить двух потомков

Оценить приспособленности потомков

Поместить потомков в новое поколение

КОНЕЦ

ЕСЛИ приспособленность лучшего потомка > порога **ТО**

останов := ИСТИНА

КОНЕЦ

КОНЕЦ

Применение генетических алгоритмов

ГА применяются для решения многих конкретных научных и технических проблем, но самое популярное приложение генетических алгоритмов — оптимизация многопараметрических функций. Эффективность ГА заключена в его способности манипулировать одновременно многими параметрами, которая использовалось во многих прикладных программах, включая проектирование интегральных схем, настройку параметров алгоритмов, поиск решений нелинейных дифференциальных уравнений и др.

Эффективность применения ГА [9] зависит от характеристик задачи. Если пространство поиска в задаче небольшое, то решение можно найти методом полного перебора и ГА не нужен для решения задачи. Если

пространство гладкое и унимодальное, то любой градиентный алгоритм будет эффективнее ГА. Если существует дополнительная информация о пространстве поиска, то методы поиска, использующие эвристики, часто превосходят ГА.

Если существует дополнительная информация о пространстве поиска, то методы поиска, использующие эвристики, часто превосходят ГА. Область применения ГА — задачи с большим пространством поиска решений и отсутствием эвристической информации. Эффективность ГА зависит не только от параметров задачи, но и от характеристик метода кодировки решений, выбора генетических операторов, настройки параметров.

Стандартный ГА

Для реализации ГА сначала выбирают подходящую структуру для представления оптимизируемых решений. С точки зрения задачи поиска, одно решение представляет точку в пространстве поиска всех возможных решений. Популяция решений состоит из одной или большего количества хромосом. Обычно хромосома — это битовая строка, хотя ГА используют не только бинарное представление. Хромосомы строят с использованием векторов вещественных чисел, строк фиксированной или переменной длины. Каждая хромосома представляет собой конкатенацию ряда подстрок, называемых генами. Гены располагаются в различных позициях или локусах хромосомы, и принимают значения, называемые аллелями, т. е. ген — бит, локус — его позиция в строке, и аллель — его значение (0 или 1). Биологический термин «генотип» относится к полной генетической модели и соответствует структуре хромосомы в ГА. Термин «фенотип» относится к внешним наблюдаемым признакам и соответствует вектору в пространстве параметров. Обычно, методика кодирования реальной переменной x состоит в ее преобразовании в двоичные целочисленные строки заданной длины.

С помощью пропорционального отбора назначают каждой i -ой хромосоме вероятность $P_s(i)$, равную отношению ее приспособленности к суммарной приспособленности популяции:

$$P_s(i) = \frac{f(i)}{\sum_{i=1}^n f(i)}$$

Отбор и замещение всех n особей происходит в соответствии с величиной $P_s(i)$. Простейший оператор пропорционального отбора — «рулетка»

(roulette-wheel selection), которая отбирает особей с помощью n «запусков» рулетки. Колесо рулетки содержит по одному сектору для каждого члена популяции. Размер i -ого сектора пропорционален соответствующей величине $P_s(i)$. При таком отборе члены популяции с более высокой приспособленностью выбираются чаще, чем особи с низкой. ГА выполняет оператор рекомбинации (кроссовер) для n выбранных особей с заданной вероятностью P_c , для чего n строк разбиваются на $(n/2)$ пар случайным образом.

Для каждой пары применяется кроссовер с вероятностью P_c . Выражение $(1 - P_c)$ выражает вероятность сохранения особей, которые переходят на стадию мутации. Потомки, полученные с помощью кроссовера, заменяют собой родителей и также переходят к мутации. Одноточечный кроссовер работает следующим образом. Сначала, случайным образом выбирается одна из $(l - 1)$ точек разрыва. Точка разрыва — граница между соседними битами в строке. Родительские структуры разрываются в этой точке на два сегмента. Затем, соответствующие сегменты различных родителей склеиваются и получаются два генотипа потомков.

После стадии кроссовера выполняются операторы мутации, которые с вероятностью P_m изменяют случайный бит в каждой строке на противоположный. Популяция, полученная после мутации, замещает старую и обработка одного поколения завершается. Последующие поколения обрабатываются в вышеописанной последовательности: отбор, кроссовер и мутация. В настоящее время используют разные операторы отбора, кроссовера и мутации. Например, турнирный отбор реализует n турниров, чтобы выбрать n особей. Каждый турнир построен на выборке k элементов из популяции, и выбора лучшей особи среди них. Наиболее распространен турнирный отбор с $k = 2$. Элитные методы отбора гарантируют, что при отборе обязательно будут выживать лучшие члены популяции [9]. Двухточечный и равномерный кроссовер могут использоваться вместо одноточечного оператора. В двухточечном кроссовере выбираются две точки разрыва, и родительские хромосомы обмениваются сегментом, который находится между двумя этими точками. В равномерном кроссовере, каждый бит первого родителя наследуется первым потомком с заданной вероятностью; в противном случае этот бит передается второму потомку.

Гибридные системы

Нечеткие нейронные сети

Преимущества аппарата нечетких нейронных сетей Эффективность аппарата нейросетей определяется их аппроксимирующей способностью, причем НС являются универсальными функциональными аппроксиматорами. С помощью НС можно выразить любую непрерывную функциональную зависимость на основе обучения НС, без предварительной аналитической работы по выявлению правил зависимости выхода от входа. Недостатком нейросетей является невозможность объяснить выходной результат, так как значения распределены по нейронам в виде значений коэффициентов весов. Основной трудностью в применении нечетких экспертных систем служит необходимость явно сформулировать правила проблемной области в форме продукции. В нечетких экспертных системах легко построить объяснение результата в форме протокола рассуждений. Поэтому в настоящее время создаются гибридные технологии, сочетающие преимущества нечетких систем и нейронных сетей. Примером гибридной технологии служит реализация системы нечетких правил на основе нейросети. База нечетких правил для двух входных и одной выходной переменных имеет следующую структуру:

$$R_i : \text{ЕСЛИ } x_{1i} \text{ is } A_{1i} \text{ И } x_{2i} \text{ is } A_{2i} \text{ ТО } z_i \text{ is } C_i.$$

Для реализации базы нечетких правил будем интерпретировать ее как таблицу определения некоторой функции, т. е. базу правил можно представить обучающей выборкой: $\{((A_{1i}, A_{2i}), C_i)\}$. Например, $\{((\text{малое, большое}), \text{около нуля})\}$. Для интеграции двух технологий — нечетких систем и нейрокомпьютинга, необходимо предложить способ четкого дискретного представления непрерывных функций принадлежности, для чего выберем максимально большой интервал $[x_1, x_2]$, в котором представлены все нечеткие множества условных частей правил. Если разбить интервал с равным шагом, то любое нечеткое значение представляется четким вектором. Другой способ представления нечеткого понятия в виде четких данных состоит в представлении нечеткого множества в виде совокупности α -срезов (рис.3).

При использовании α -срезов каждое α_j -подмножество представляется двумя числами — левой и правой границами: $\alpha^{L_{ij}}, \alpha^{R_{ij}}$, где j — номер

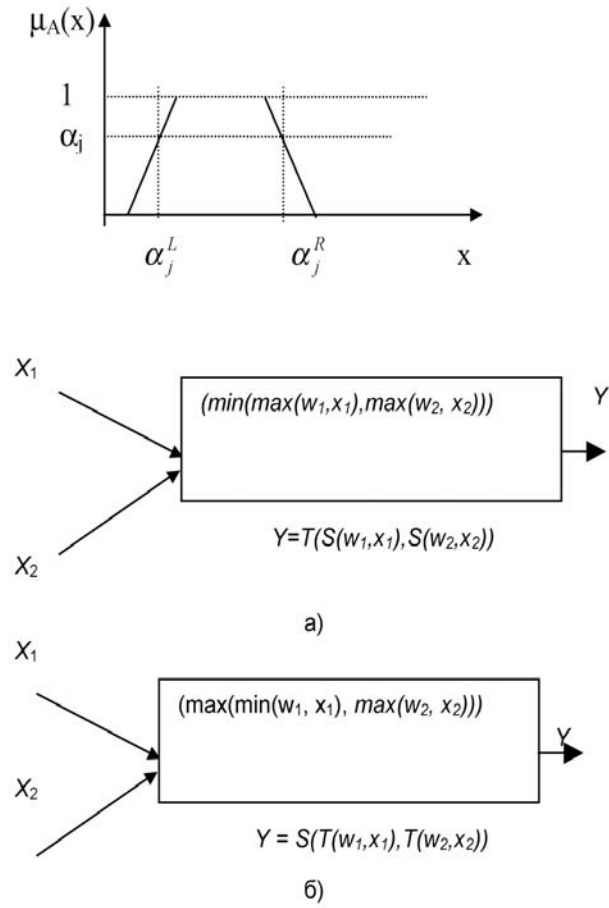


Рис. 3. α -срез, И-ИЛИ нейроны

α -среза, а i — номера точек на его левой и правой границах, т. е. α -срезы четко представляют непрерывную функцию принадлежности.

Понятие нечеткой нейросети. Глубинная интеграция нечетких систем и нейросетей связана с разработкой моделей нейронов, функции которых отличаются от функций традиционного нейрона. Модификация модели нейрона для адаптации к нечетким системам касается выбора функции активации, реализации операций сложения и умножения, так как в нечеткой логике сложение моделируется любой треугольной конормой (например, $\max, a + b - a * b, \dots$), а операция умножения — треугольной нормой ($\min, a * b, \dots$).

И-нейроном называется нейрон, в котором умножение веса w на вход x моделируется конормой $S(w, x)$, а сложение — нормой $T(w, x)$.

Для двухвходового И-нейрона справедлива формула:

$$Y = T(S(w_1, x_1), S(w_2, x_2)).$$

ИЛИ-нейроном называется нейрон, в котором умножение веса w и входа x моделируется нормой $T(w, x)$, а сложение взвешенных весов — конормой $S(w, y)$. Для двухвходового ИЛИ-нейрона справедлива формула:

$$Y = S(T(w_1, x_1), T(w_2, x_2)).$$

Если выбрать в качестве T -нормы $\min$, а S -нормы — $\max$, то формула преобразования ИЛИ-нейрона уточняется следующим образом:

$$\max(\min(w_1, x_1), \min(w_2, x_2)).$$

В качестве функции активации обычно используют радиальную базисную функцию $F(x) = \exp(-b * (x^2 - a))$.

Нечеткой нейронной сетью (ННС) называют четкую нейронную сеть прямого распространения сигнала, которая построена на основе многослойной архитектуры с использованием И-ИЛИ-нейронов [11, 13, 15]. Нечеткая нейросеть функционирует стандартным образом на основе четких действительных чисел. Нечеткой является только интерпретация результатов. При создании гибридной технологии можно использовать нейрокомпьютинг для решения частной задачи нечетких экспертных систем, а именно настройки параметров функции принадлежности. Традиционно функции принадлежности формируют двумя способами: методом экспертной оценки

или на основе статистики. Гибридные технологии предлагают третий способ: в качестве функции принадлежности выбирается параметризованная функция формы (например, параметризованная гауссова кривая), параметры которой настраиваются с помощью нейросетей. Настройка параметров может быть получена с помощью алгоритма обратного распространения ошибки.

Рассмотрим применение алгоритма обратного распространения ошибки для обучения ННС. Пусть задана следующая система нечетких правил:

ЕСЛИ x_1 is A_{1j} И $\dots$ x_n is A_{nj} ТО z_j is C_j .

где A_{ij} — нечеткие числа; C_j — действительные числа. Значение $\alpha_j = \Pi_i a_{ji}$ — сила или достоверность правила, α_{ji} — степень принадлежности нечеткого множества условной части j -го правила

x_1 is A_{1i} И $\dots$ x_n is A_{ni} ,

$i = 1, \dots, n$ — номер входной переменной, $j = 1, \dots, m$ — количество правил. Значение выхода вычисляется по формуле средневзвешенного выхода:

$$z = \frac{\sum_{j=1}^m \alpha_j * z_j}{\sum_{j=1}^m \alpha_j}.$$

Допустим, что разработана нейросеть с n входами и одним выходом. Каким образом такая НС может аппроксимировать базу нечетких правил? Любая совокупность нечетких продукций может рассматриваться как нелинейное соответствие, заданное таблицей определения $\{(x_k, y_k)\}$, где $k = 1, \dots, K$ — номер строки-образца в обучающей выборке, x — вектор входа, y — желаемое значение выхода, а z — значение выхода, вычисляемое нейросетью. Если определить текущую ошибку с помощью формулы $E_k = 1/2 * (z_k - y_k)^2$, то можно применить стандартный алгоритм коррекции ошибки, корректируя выход Z по следующему правилу:

$$Z(t+1) = Z(t) - \eta * \left(\frac{\partial E_k}{\partial Z} \right),$$

где η — уровень обучения, $(\partial E_k / \partial Z)$ — направление градиента снижения ошибки. Подставляя в правило формулу для средневзвешенного выхода ННС, получим:

$$Z(t+1) = Z(t) - \eta * (Z_k - Y_k) * \left[\frac{\alpha_j}{a_1 + a_2 + \dots + a_m} \right].$$

При применении стандартного алгоритма обратного распространения ошибки для настройки выхода ННС, необходимо изменить параметры функций принадлежности условных частей правил, т. е. обучение сети позволит их настроить на обучающую выборку.

Структуры гибридных систем (ГС)

Рассмотрим структуры ГС, решающих задачу управления, выделим особенности архитектуры и алгоритмов обучения для каждого конкретного типа ГС [7, 13, 15].

NNFLC — нечеткий контроллер на основе НС. Структура NNFLC приведена на рис. 4.

Структурно NNFLC — это многослойная сеть прямого распространения сигнала, причем различные слои выполняют разные функции. Опишем кратко функции слоев.

- **Слой 1:**

$$y_i^{(1)}(x) = \exp[-(x_i - c_i)^2 / 2 * \sigma_i^2].$$

Слой 1 представляет функции принадлежности, реализованные как радиальные базисные нейроны.

- **Слой 2:**

$$y_i^{(2)} = \min[y_1^{(1)}, \dots, y_n^{(1)}].$$

Слой 2 моделирует *И*-условия правил.

- **Слой 3:**

$$y_i^{(3)} = \max[y_1^{(2)}, \dots, y_n^{(2)}].$$

Слой 3 представляет собой *ИЛИ*-комбинацию правил с одинаковыми термами в консеквентах. Слой 3 выполняет разные функции в рабочем режиме и в режиме обучения. В режиме обучения слой настраивает параметры функций принадлежности выходных переменных. В рабочем режиме формирует значение выхода.

- **Слой 4:**

В рабочем режиме нейроны выполняют дефазификацию:

$$z_i^{(4)} = \sum_j w_{ji} * y_i^{(3)}$$

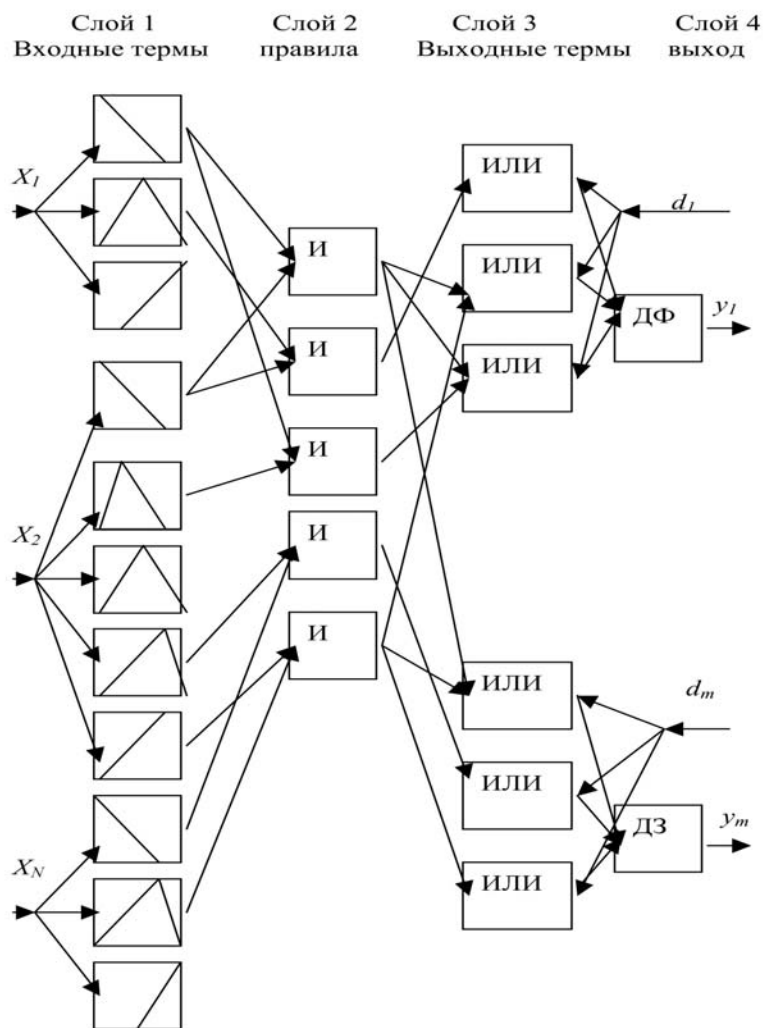


Рис. 4. Структура NNFLC

и нормализацию:

$$y^{(4)}(z_i^{(4)}) = \sum_j y_j^{(3)} * \frac{w_{ji}}{\sum_j w_{ji}},$$

а в режиме обучения — это дополнительный вход, позволяющий настроить функцию принадлежности выходной переменной. Структура ННС NNFLC инициализируется по принципу формирования полной матрицы правил. Если x_i — входные переменные, $\tau(x_i)$ — количество нечетких меток (разбиений) x_i , то исходное количество правил:

$$T = \prod_{i=1}^n \tau(x_i).$$

Обучение ННС сложной архитектуры (с различными функциональными слоями) обычно происходит многоэтапно, причем, на каждом этапе используются различные алгоритмы обучения: предобучение (offline), оперативное (online), без учителя, с учителем. Общая схема обучения ННС NNFLC содержит следующие этапы:

- формирование обучающих данных;
- самоорганизующаяся кластеризация (настройка функций принадлежности);
- соревновательное обучение (алгоритм победителя);
- удаление правил;
- комбинирование правил;
- окончательная настройка параметров (тюнинг) функций принадлежности с помощью алгоритма обратного распространения ошибки.

Приведем содержательные характеристики этапов обучения. Настройка параметров функций принадлежности включает в себя определение центров c_i и ширины σ_i для функций принадлежности, представленных функциями формы:

$$y_i^{(1)}(x) = \exp[-(x_i - c_i)^2 / 2 * \sigma_i^2].$$

Алгоритм победителя выявляет c_i : $\Delta c_w(t) = \eta(t)(x - c_w(t))$, где $\eta(t)$ — монотонно убывающий уровень обучения. Настройка ширины σ_i осуществляется эвристически, например, по принципу «первого ближайшего соседа»: $\sigma_i = (\sigma_i - \sigma_w) / \lambda$, λ — параметр перекрытия. Алгоритм победителя ищет матрицу весов w_{ji} , которая оценивает качество связей левой и правой частей правил:

$$\Delta w = y_j^{(3)} * (y_i^{(2)} - w_{ji}).$$

Комбинирование правил часто целесообразно выполнять с участием эксперта. Окончательная настройка функций принадлежности выполняется с помощью алгоритма обратного распространения ошибки для функции ошибки $e_i = (y_i^{(4)} - d_k)^2$. Цепочка правил распространяет ошибку до слоя 1 с обратным роутингом. Таким образом, можно сделать вывод о том, что архитектура NNFLC может быть проинтерпретирована как система нечеткого вывода Такаджи-Суджено.

ANFIS — адаптивная НС, основанная на системе нечеткого вывода. Приведем на рис. 5 структуру ANFIS (Adaptive Network Based Fuzzy Inference System) — адаптивной НС для двух правил:

ЕСЛИ $x_1 = A_1$ И $x_2 = B_1$ ТО $y_1 = c_{11} * x_1 + c_{12} * x_2$

ЕСЛИ $x_1 = A_2$ И $x_2 = B_2$ ТО $y_2 = c_{21} * x_1 + c_{22} * x_2$.

Выход ННС формируется по формуле:

$$y = \frac{(w_1 * y_1 + w_2 * y_2)}{(w_1 + w_2)}$$

Слои ННС ANFIS выполняют следующие функции.

- **Слой 1** представлен радиальными базисными нейронами и моделирует функции принадлежности.
- **Слой 2** — это слой И-нейронов, которые моделируют логическую связку И произведением $w_i = \mu_{A_i}(x_1) * \mu_{B_i}(x_2)$.
- **Слой 3** вычисляет нормированную силу правила:

$$w_i = \frac{w_i}{(w_1 + w_2)}.$$

- **Слой 4** формирует значение выходной переменной:

$$y(x_1, x_2) = w_i y_i = w_i (c_{i1} * x_1 + c_{i2} * x_2).$$

- **Слой 5** выполняет дефазификацию:

$$y = w_1 * y_1 + w_2 * y_2.$$

Гибридная сеть архитектуры ANFIS обучается с помощью алгоритма обратного распространения ошибки.

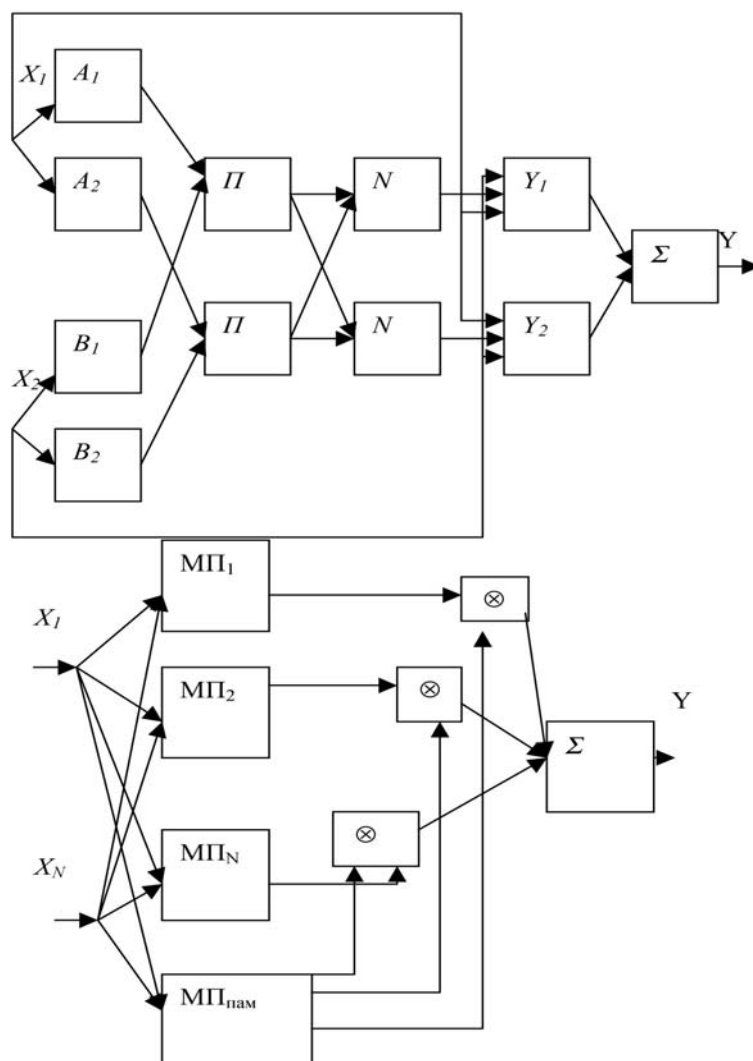


Рис. 5. Структуры ANFIS и NDNFR

NNDFR — НС для нечетких умозаключений. ННС NNDFR (Neuron Network Driven Fuzzy Reasoning) (рис. 5) способна выполнять кластеризацию в пространстве, содержащем нелинейные границы кластеров.

Обозначим через MP_i двухслойный перцептрон с одним скрытым слоем, который служит для описания желаемой кривой поверхности решений, т. е. MP_i выражает правило:

$$\text{ЕСЛИ } x_1 = A_1 \text{ И } x_{ni} = A_{ni} \text{ ТО } y_i = f_i(x_i).$$

Для обучения NNDFR предложен алгоритм «удаления правил назад». Алгоритм включает в себя следующие шаги.

- **Шаг 1.** Удаление малозначимых входных переменных. Обучаются отдельные MP_i . После отдельного обучения многослойные перцептроны объединяются и обучаются вместе. Если изменение некоторой входной переменной не влияет на сумму квадратов ошибок, то соответствующая переменная удаляется.
- **Шаг 2.** Формирование обучающей и тестовой выборок. Обучающая выборка $\{x_k, d_k\}$ делится на обучающие и контрольные образцы — паттерны.
- **Шаг 3.** Обучение «многослойного перцептрона (МП) памяти». Эвристически определяется количество кластеров (правил). МП памяти обучается выделять кластеры X . Если x_k принадлежит к A_j , то выход МП памяти $w_j(x_k) = 1$, иначе выход равен 0.
- **Шаг 4.** Обучение MP_i . Многослойные перцептроны MP_i обучаются отдельно методом обратного распространения ошибки.
- **Шаг 5.** Обратное исключение. Используется обратное исключение переменных для каждого MP_i . Общая схема рабочего режима ННС подчиняется следующим правилам:

$$\begin{aligned} \text{ЕСЛИ } X = A_1 \text{ ТО } MP_1(x_1, \dots, x_n), \\ \text{ЕСЛИ } X = A_2 \text{ ТО } MP_2(x_1, \dots, x_n), \\ \text{ЕСЛИ } X = A_n \text{ ТО } MP_{mem}(x_1, \dots, x_n). \end{aligned}$$

GARIC — интеллектуальное управление, основанное на обобщенном приближенном выводе. Архитектура GARIC (Generalized Approximate Reasoning Based Intelligent Control) имеет в своем составе две НС различного функционального назначения: НС оценки состояния и НС выбора действия. Используемые алгоритмы обучения — алгоритм усиления и обратного распространения ошибки.

Рассмотрим место GARIC в контуре управления (рис. 6). Рис. 6 и рис. 7 отражают структуру НС управления и НС оценки.

ННС GARIC реализует схему вывода Цукамото. НС управления обучается с помощью алгоритма обратного распространения ошибки. Сложнее обучить НС оценки, которая тренируется алгоритмом усиления:

$$y_j^n(t+1) = f \left[\sum_i w_{ij}(t) * x_i(t+1) \right],$$

$$f(z) = [1 + \exp(-z)]^{-1},$$

$$\sigma(r, r_p) = r - r_p,$$

$$\sigma(r, r_p) \approx r - r_p + \gamma * r_p,$$

$$0 \leq \gamma \leq 1,$$

где $\sigma(r, r_p)$ — разница реального и «внутреннего» r . Случайная модуляция y выполняется колоколообразной, случайно распределенной функцией с центром y и разбросом $\sigma(r, r_p)$, что необходимо для покрытия входного пространства. НС оценки представляет собой линейный многослойный перцептрон. Матрицы B и C изменяются стратегией «награда-штраф».

$$\Delta b_i(t) = \eta * \sigma^{(t+1)} * x_i^t,$$

$$\Delta c_j(t) = \eta * \sigma^{(t+1)} * y_j^t.$$

Матрица W обучается упрощенным алгоритмом обратного распространения ошибки:

$$\Delta w_{ij}^{(t)} = \eta * \sigma^{(t+1)} * \text{sign}[c_j^{(t)}] * y_j^t * (1 - y_j^{(t)}) * x_i^{(t)}.$$

Нечеткая сеть Fuzzy Net (FUN) Нечеткая нейронная сеть FUN (рис. 7) предложена для управления перемещениями мобильного робота. В этой сети осуществляется структурное и параметрическое обучение — последовательный стохастический поиск в пространствах правил и входов. Обучение

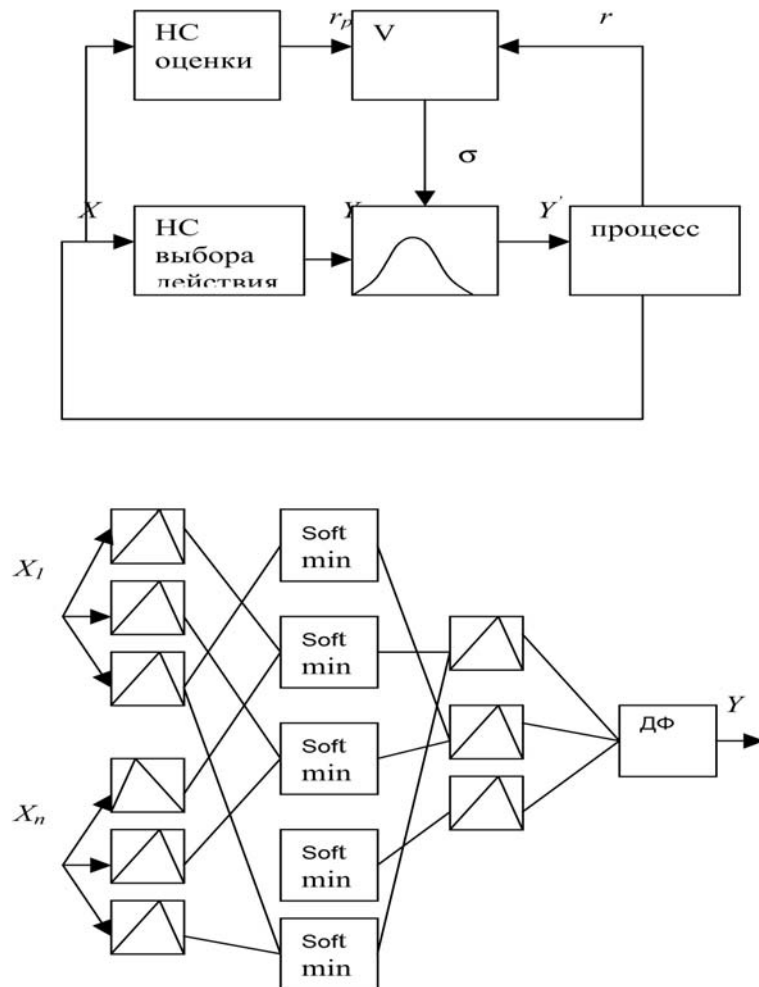


Рис. 6. Нейро-нечеткий контроллер GARIC.НС управления

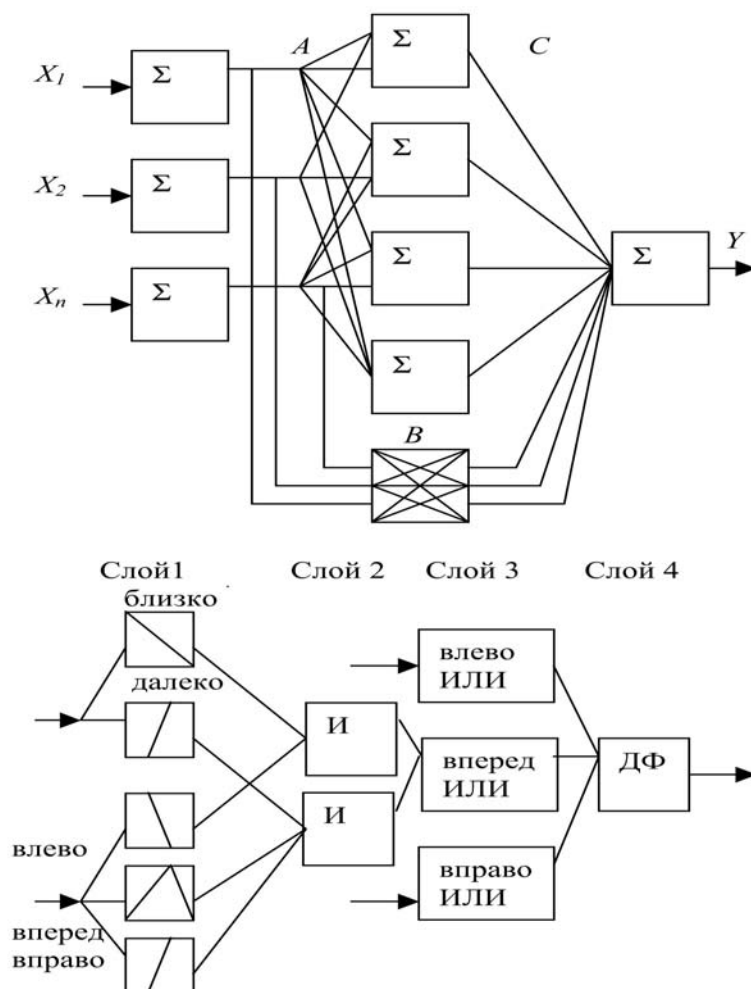


Рис. 7. Структура НС оценки GARIC

правил построено следующим образом. Случайно выбранная связь правил до и после слоя 2 изменяется и проверяется сеть на улучшение ценовой функции. Плохие изменения отменяются, а хорошие сохраняются. Параметры функций принадлежности итерационно изменяются от некоторого случайного значения: $D_p^{(t)} - D_p(t-1) = -\eta * D_p^{(t-1)}$, где $\eta = 10^{-3} \dots 10^{-2}$ — фактор сходимости.

Таким образом, анализ конкретных архитектур нечеткой нейронной сети позволяет сделать следующий вывод. Наряду с классическими нейронами, являющимися пороговыми суммирующими элементами, должна включать в себя «логические» нейроны, моделирующие логические связи. Фундаментальная задача объединения методов перцепции (восприятия) и логического (абстрактного) мышления не могла не отразиться на уровне структуры ГС. В нечеткой нейронной сети любой из известных архитектур, слои нейронов становятся специализированными. Существуют слои, выполняющие распознавание, и слои, вычисляющие логические формулы и импликации (правила ЕСЛИ ... ТО). Причем конкретная схема нечеткой нейронной сети зависит от решаемой задачи.

Нечеткий нейронный контроллер. Рассмотрим особенности применения нечетких нейронных сетей на примере нечеткого регулятора для стиральной машины, то есть построим теперь fuzzy-неуго контроллер. Выберем для демонстрации технологии нечетких нейронных сетей архитектуру ANFIS. Основа интеграции нейронных сетей и систем нечеткого вывода заключается в том, что оба метода представляют нелинейное отношение в пространстве входов и выходов. Важная задача нечеткого моделирования — это настройка функций принадлежности, являющаяся по существу задачей оптимизации. Для ее решения используются как нейронные сети, так и генетические алгоритмы. Самый простой подход к настройке функций принадлежности заключается в том, что выбирают определенную параметризованную функцию формы и подбирают ее параметры на основе обучения нейронной сети. Рассмотрим простой контроллер с тремя нечеткими

правилами вывода в базе знаний:

R_1 : ЕСЛИ «количество белья» = «много»
 И «температура воды» = «высокая»
 И «загрязненность» = «высокая»
 ТО «длительность» = «высокая» ;
 R_2 : ЕСЛИ «количество белья» = «много»
 И «температура воды» = «высокая»
 И «загрязненность» = «низкая»
 ТО «длительность» = «низкая» ;
 R_3 : ЕСЛИ «количество белья» = «мало»
 И «температура воды» = «низкая»
 И «загрязненность» = «низкая»
 ТО «длительность» = «низкая» .

Зададим следующие функции формы для правил. Пусть выражению «количество белья» = «мало» соответствует функция $L_1(x)$, а «количество белья» = «много» функция $H_1(x)$:

$$L_1(x) = \frac{1}{(1 + \exp(b_1(x - c_1)))},$$

$$H_1(x) = \frac{1}{(1 + \exp(-b_1(x - c_1)))},$$

$$L_1(x) + H_1(x) = 1.$$

Аналогично для выражения «температура воды» = «высокая» будем использовать функцию $L_2(t)$, а для «температура воды» = «низкая» функцию $H_2(t)$:

$$L_2(t) = \frac{1}{(1 + \exp(b_2(t - c_2)))},$$

$$H_2(t) = \frac{1}{(1 + \exp(-b_2(t - c_2)))},$$

$$L_2(t) + H_2(t) = 1.$$

Для выражения «загрязненность» = «высокая» определим функцию $L_3(z)$ и для «загрязненность» = «низкая» функцию $H_3(z)$:

$$\begin{aligned} L_3(z) &= \frac{1}{(1 + \exp(b_3(z - c_3)))}, \\ H_3(z) &= \frac{1}{(1 + \exp(-b_3(z - c_3)))}, \\ L_3(z) + H_3(z) &= 1. \end{aligned}$$

Для выходов «длительность» = «высокая» и «длительность» = «малая» аналогично определим функции

$$\begin{aligned} L_4(y) &= \frac{1}{(1 + \exp(b_4(y - c_4)))}, \\ H_4(y) &= \frac{1}{(1 + \exp(-b_4(y - c_4)))}, \\ L_4(x) + H_4(x) &= 1. \end{aligned}$$

Система нечеткого вывода по Цукамото представлена на рис. 8.

Для четких значений «количество белья», «температура воды» и «загрязненность»: A_1, A_2, A_3 определим релевантность (силу) правил α_i , как показано на рис. 8:

$$\begin{aligned} \alpha_1 &= H_1 \cap H_2 \cap H_3, \\ \alpha_2 &= H_1 \cap H_2 \cap L_3, \\ \alpha_3 &= L_1 \cap L_2 \cap L_3. \end{aligned}$$

Выходы по каждому из правил определяется с помощью обратных функций принадлежности правых частей правил.

$$\begin{aligned} Y_1 &= H_4^{-1}(\alpha_1), \\ Y_2 &= H_4^{-1}(\alpha_2), \\ Y_3 &= L_4^{-1}(\alpha_3). \end{aligned}$$

Общий выход из системы нечетких правил определяется как

$$y_0 = \frac{(\alpha_1 * Y_1 + \alpha_2 * Y_2 + \alpha_3 * Y_3)}{(\alpha_1 + \alpha_2 + \alpha_3)}.$$

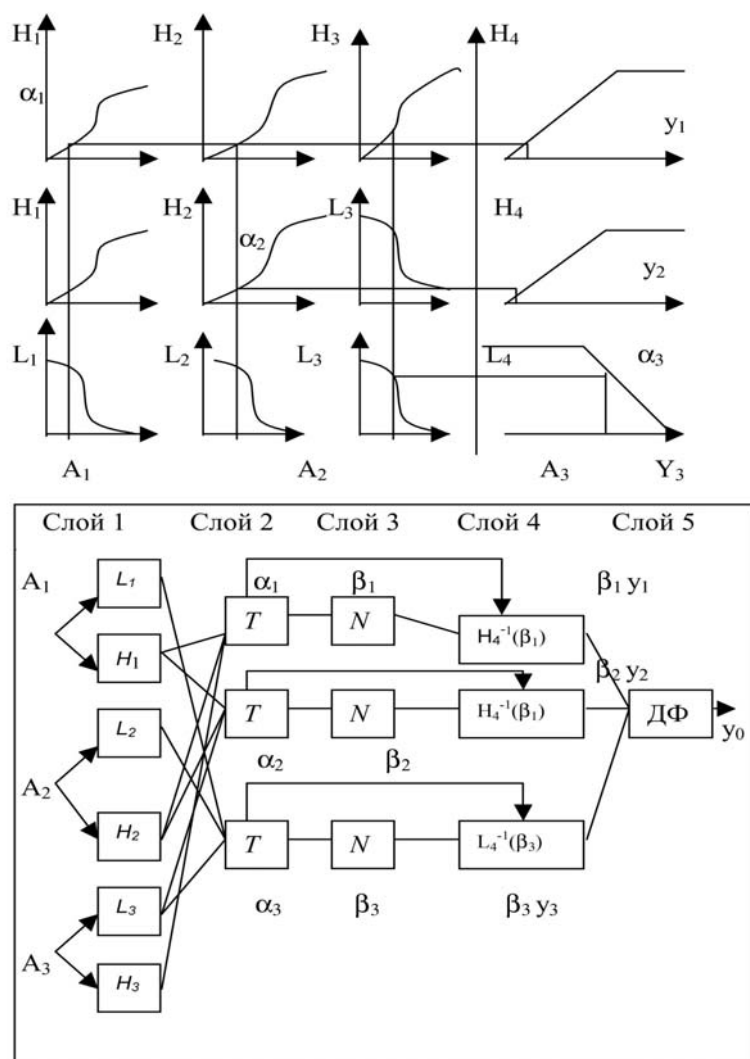


Рис. 8. Нечеткий вывод и структура нейро-нечеткого контроллера

Далее построим нечеткую нейронную сеть, идентичную системе нечеткого вывода и обучим функции принадлежности анцедента и консеквента правил (рис. 8):

- **Слой 1.** Выходы узлов — это степени, с которыми заданные входы удовлетворяют функциям принадлежности, ассоциированным с этими узлами.
- **Слой 2.** Каждый узел вычисляет силу правила. Выход верхнего нейрона $\alpha_1 = H_1 \cap H_2 \cap H_3$, выход среднего нейрона $\alpha_2 = H_1 \cap H_2 \cap L_3$, а выход нижнего: $\alpha_3 = L_1 \cap L_2 \cap L_3$. Все узлы помечены T , так как можно выбрать любую T -норму для моделирования логического И. Узлы этого слоя называются узлами правил.
- **Слой 3.** Каждый узел помечен N , чтобы показать, что узлы нормализуют силу правил

$$\beta_i = \frac{\alpha_i}{(\alpha_1 + \alpha_2 + \alpha_3)}.$$

- **Слой 4.** Выход нейронов — это произведение нормализованной силы правила и индивидуального выхода соответствующего правила.

$$Y_1 = H_4^{-1}(\alpha_1),$$

$$Y_2 = H_4^{-1}(\alpha_2),$$

$$Y_3 = L_4^{-1}(\alpha_3).$$

- **Слой 5.** Одиночный выходной нейрон вычисляет выход сети

$$y_0 = \beta_1 * Y_1 + \beta_2 * Y_2 + \beta_3 * Y_3.$$

Алгоритмы обучения для нечеткой нейронной сети контроллера.

Пусть задана четкая обучающая выборка

$$\{(x_1, t_1, z_1, y_1), \dots, (x_k, t_k, z_k, y_k)\},$$

где x, t, z — входные условия «количество белья», «температура воды» и «загрязненность», а y — длительность стирки. Ошибку на k -ом образце определим как обычно $E^k = \frac{1}{2} * (O_i - Y_i)^2$. Можно настроить параметры

функций принадлежности в ходе обучения нечеткой нейронной сети, используя следующие соотношения как правила изменения весов в алгоритме обратного распространения ошибки:

$$\begin{aligned} b_4(t+1) &= b_4(t) - \eta * \frac{\partial E}{\partial b_4}, \\ c_4(t+1) &= c_4(t) - \eta * \frac{\partial E}{\partial c_4}, \\ &\dots \\ c_1(t+1) &= c_1(t) - \eta * \frac{\partial E}{\partial c_1}. \end{aligned}$$

Нечеткие нейронные сети с генетической настройкой Настройка функций принадлежности с помощью нейронной сети или ГА устраняет принципиальную слабость теории нечетких систем — субъективность функций принадлежности. Если настроить функцию принадлежности с помощью нейронной сети, то окончательная форма функции будет аппроксимацией обучающей выборки. С помощью ГА, т.е. стохастической оптимизации также можно настроить функции принадлежности. Используем традиционный градиентный метод для обучения параметров левой и правой частей нечетких правил. Покажем, как можно настроить параметры функции формы.

Генетической нечеткой системой называют нечеткую систему, функции принадлежности и база правил которой спроектирована с помощью генетического алгоритма. Для применения ГА необходимо выполнить следующие действия:

- выбрать единицу кодирования, т.е. хромосому;
- уточнить эволюционные операторы рекомбинации, мутации и селекции;
- сформировать функцию оптимальности (fitness function, performance index).

В настоящее время ГА используют либо для настройки функций принадлежности (базы данных), либо для формирования базы правил (базы знаний), либо для одновременного формирования и функций принадлежности, и правил. В соответствии с объектами оптимизации выбирают единицу кодирования — хромосому. Для настройки функций принадлежности за хромосому выбирают одно правило (мичиганский подход), для настройки

базы правил за хромосому выбирают вариант базы правил (питтсбургский подход, подход итеративного обучения правил) [12, 17]. В соответствии с кодированием уточняются правила генерации новых хромосом.

Функция оптимальности представляет собой механизм нечеткого вывода, который для каждого варианта базы правил строит либо управляющее воздействие для нечеткого контроллера, либо экспертное заключение для диагностической экспертизы. Как видно из описания механизма нечеткого вывода, все нечеткие правила вносят вклад в окончательный результат, то есть правила сотрудничают. Но при отборе правил (хромосом) ГА сохраняет только правила, внесшие максимальный вклад в общий результат, т. е. хромосомы конкурируют. Говорят о проблеме «конкуренции и кооперации» в генетических нечетких системах (a competition vs. a cooperation). Решение проблемы в каждом конкретном случае строится эвристически, например, при подходе итеративного обучения правил используют два этапа оптимизации. На первом шаге правила конкурируют за право войти в базу правил, а на втором — взаимодействуют при формировании общего результата.

Системы генетического проектирования нечетких нейронных сетей

Рассмотрим возможности генетических вычислений, как средства структурной оптимизации нечетких нейронных сетей. Генетические вычисления можно применить на этапе проектирования нейронной сети, как предиктора временных рядов, в том числе можно проектировать и нечеткую нейронную сеть. Возможности нейронных сетей интерполировать значения временных рядов широко используют для оценки поведения макроэкономических показателей, в том числе индексов деловой активности или уровня ценных бумаг. Задача построения нейронного предиктора связана с принятием решений на разных этапах проектирования. Необходимо исследовать входные данные и решить, по какому отрезку входных данных рационально делать предсказания, сколько точек в будущем разумно предсказать. Пространство перебора решений огромно для развитых рынков ценных бумаг, накопивших статистические сведения за десятки лет. Необходимо принять решение и о типе нейронной сети: многослойный персептрон, радиально базисная сеть, вероятностная нейронная сеть, сеть регрессии и т. д. Для выбранного типа НС необходимо определить количество скрытых слоев, количество нейронов в них, виды функций активации и другие. Для каждой сети необходимо выбрать алгоритм обучения и его

параметры, например, использование моментов (уровень моментов), уровень обученности, целевой уровень накопленной ошибки и т.д. Для нечеткой нейронной сети необходимо определить архитектуру сети, параметры функций принадлежности. В результате пространство решений при формировании нечеткого нейронного предиктора становится необозримым для исследователя и целесообразно применить ГА как средство эволюционного проектирования.

Примером такого средства может служить программная система Neuro-Forecaster GENETICA (NFGA Copyright 1993–1998, NIBS Pte Ltd, 62 Fowlie Road, Republic of Singapore 428501.) Программа NFGA — это программа поиска и оптимизации НС, основанная на генетических алгоритмах, созданная для поиска лучшей комбинации входных данных, лучшего периода прогноза (т. е. количества прогнозируемых значений) для данной проблемы.

Заключение

В настоящее время рассмотренные технологии «вычислительного интеллекта» являются успешными прикладными технологиями. Самыми яркими событиями были [1,4] запуск в 1987 году системы управления новым метро в г.Сендай около Токио, достижение японского экспорта изделий с fuzzy logic к 1991 году цифры в 25 млрд. долл. США. Это, в первую очередь, товары культурно-бытового назначения — фотоаппараты, видеокамеры, стиральные машины, холодильники, пылесосы, микроволновые печи и многое другое. Международный научно-исследовательский институт (LIFE — Laboratory for International Fuzzy Engineering Research), созданный в Японии располагал в 1993 году бюджетом в 64 млн. долл. Таким образом, «вычислительный интеллект» — это успешная электронная и программная индустрия.

Литература

1. Абдусаматов Р. М., Беркинблит М. Б., Фельдман А. Г., Чернавский А. В. Моторика и интеллект // Интеллектуальные процессы и их моделирование. — М.: Наука, 1987.
2. Аверкин А. Н. и др. Нечеткие множества в моделях управления и искусственного интеллекта / Под ред. Д. А. Поспелова. — М.: Наука, 1986.

3. *Аверкин А. Н., Федосеева И. Н.* Параметрические логики в интеллектуальных системах управления. – М.: Вычислительный центр РАН, 2000.
4. *Васильев В. И., Ильясов Б. Г.* Интеллектуальные системы управления с использованием генетических алгоритмов. Учебное пособие. – Уфа: УГАТУ, 1999.
5. *Заде Л. А.* Понятие лингвистической переменной и его применение к принятию приближенных решений. – М.: Мир, 1976.
6. *Клир Дж.* Системология. Автоматизация решения системных задач: Пер. с англ. – М.: Радио и связь, 1990.
7. *Кофман А.* Введение в теорию нечетких множеств. – М.: Радио и связь, 1982.
8. *Мелихов А. Н., Берштейн Л. С., Коровин С. Я.* Ситуационные советующие системы с нечеткой логикой. – М.: Наука, 1990.
9. *Наместников А. М., Ярушкина Н. Г.* Эффективность генетических алгоритмов для задач автоматизированного проектирования // Известия РАН. Теория и системы управления. – 2002. – № 2.
10. *Скурихин А. Н.* Генетические алгоритмы // Новости искусственного интеллекта. – 1995. – № 4.
11. *Ярушкина Н. Г.* Нечеткие нейронные сети // Новости искусственного интеллекта. – 2001. – № 2–3.
12. *Ярушкина Н. Г.* Гибридные системы, основанные на мягких вычислениях: определение, архитектура, возможности // Программные продукты и системы. – 2002. – № 3.
13. *Bothe H.-H.* Fuzzy Neural Networks. – Prague: IFSA, 1997.
14. *Deichmann H.* Linguistische Systeme und ihre Anwendung. – Darmstadt: Fachhochschule Darmstadt, 1996.
15. *Fuller R.* Fuzzy systems.
URL: <http://www.abo.fi/~rfuller/>
16. *Goldberg D. E.* Genetic Algorithms in Search, Optimization and Machine Learning. – Addison-Wesley Publishing Company Inc., 1989.
17. *Herrera F., Magdalena L.* Genetic fuzzy systems. – Prague: IFSA, 1997
18. *Holland Y.* Adaptation in Natural and Artificial Systems. – University of Michigan, Press, Ann Arbor, 1975.
19. *Schweizer B., Sklar A.* Associative functions and abstract semigroups // Publ. Math., No. 10, 1963.
20. *Zadeh L. A.* Fuzzy Sets // Information and Control, **8** (1965), 338–353.
21. *Zadeh L.* What is soft computing? // Soft Computing, 1997, **1**, pp.1–2.

Надежда Глебовна ЯРУШКИНА, доктор технических наук, заведующая кафедрой Информационные системы Ульяновского государственного технического университета. Область научных интересов — теория нечетких множеств и систем, искусственные нейронные сети, генетические алгоритмы, гибридные системы. Автор 2 монографий и более 150 научных публикаций.

НАУЧНАЯ СЕССИЯ МИФИ–2004

НЕЙРОИНФОРМАТИКА–2004

VI ВСЕРОССИЙСКАЯ
НАУЧНО-ТЕХНИЧЕСКАЯ
КОНФЕРЕНЦИЯ

ЛЕКЦИИ
ПО НЕЙРОИНФОРМАТИКЕ
Часть 1

Оригинал-макет подготовлен Ю. В. Тюменцевым
с использованием издательского пакета $\text{\LaTeX}$ 2_ε
и набора PostScript–шрифтов PSCyr

Подписано в печать 25.11.2003 г. Формат 60 × 84 1/16
Печ. л. 12, 5. Тираж 200 экз. Заказ №

*Московский инженерно-физический институт
(государственный университет)
Типография МИФИ
115409, Москва, Каширское шоссе, 31*