

Формирование понятийного мышления на алгебре фурье-дуальных операций

Аннотация. Рассмотрены два этапа перехода от образного мышления к понятийному: индуктивное порождение понятий при обработке набора образов и последующее установление связей новых понятий с уже существующими. Задача решена в рамках биологически мотивированного подхода, предполагающего представление образной и понятийной информации паттернами внутренней репрезентации и их обработку. Дана модель решения обоих этапов задачи на алгебре фурье-дуальных определяющих операций. Индуктивное порождение понятий реализуется обобщением на ряде примеров – паттерн общих признаков формируется по критерию их коррелированности в наборе примеров. Представлена оценка эффективности выделения паттерна общих признаков в зависимости от числа примеров и оценок информационной емкости паттернов. Предложена архитектура нейронной сети, реализующая модель. Приведены результаты численных экспериментов на примере обращения силлогизма «Dagii».

Ключевые слова: образное мышление, понятийное мышление, индуктивное формирование понятий, фурье-дуальность, нейронные сети, паттерн внутренней репрезентации информации.

Введение

Традиционно различают две формы мышления: образное и понятийное [1, 2]. Понятийное рассматривается как более высокая форма, развивающаяся на основе образного посредством формирования интеллектуальной системой понятий и установления связей между ними в процессе обработки образов, имеющих преимущественно сенсорное происхождение.

В статье рассмотрен один из возможных методов перехода от образного к понятийному мышлению в рамках биологически мотивированного подхода – на нейросетях, реализующих алгебру фурье-дуальных операций.

В рамках нейросетевого подхода любое внутреннее представление в интеллектуальной системе (биологической или искусственной нейронной сети – НС), реализуется в виде пространственно-временных картин активности нейронных ансамблей [3]. Далее эти картины нейронной активности будем именовать паттернами внутренней репрезентации – ПВР.

Используем определение образа и понятия как крайних случаев внутренних представлений согласно [2]: образ в пространстве признаков (абстрактном в [2]) представлен унимодальным распределением с локализованным максимумом, а понятие – пологим холмом, охватывающим широкий диапазон значений только тех признаков, что существенны для понятия. Другие категории внутренних представлений информации, кроме «образ» и «понятие», в рамках настоящего рассмотрения для большей ясности изложения основной идеи не используются.

Первый шаг перехода от образного мышления к понятийному – формирование понятий. Один из методов порождения понятий, точнее – гипотез понятий, при обработке образов – индуктивный вывод [4], обеспечивающий расширение теории [5] и преобразование данных в знания [6]. Признанный ДСМ-метод [7, 8] развит в рамках формально-символьного подхода. Представляют интерес также бионические методы, поскольку мышление, как НС-обработка ПВР, характеризуется высоким параллелизмом и малой глубиной [1], а вопрос вычислительной

¹Работа выполнена при финансовой поддержке РФФИ, проект № 15-01-04111-а.

сложности для НС-подхода в принципе менее актуален, чем в ДСМ-методе [9, 10].

Реализация индуктивного вывода на НС основана на их способности к обобщению [11–14]. В статье использована модель НС со связями слоев в пространстве Фурье, описываемая алгеброй фурье-дуальных определяющих операций [15, 16]. Краткое изложение модели дадим в разделе 1.2, а сейчас отметим, что согласно результатам нейрофизиологических исследований [17–21], нейронные структуры зрительного отдела коры головного мозга реализуют преобразование Фурье. Эти результаты лежат в основе одной из двух основных на сегодня концепций анализа признаков изображений зрительной корой – гипотезы пространственных фильтров. В свете упомянутых нейрофизиологических результатов модель, основанная на преобразовании Фурье, представляется биологически мотивированной.

В статье [22] дан основанный на предложенной в [2] классификации «образ – понятие» метод индуктивного вывода на данной алгебре. Метод основан на использовании шкалы частот в пространстве Фурье как шкалы общности признаков и формировании паттерна понятия абстрагированием от частных признаков. Таким образом, в [22] пространство признаков уже не абстрактное, как в подходе [2], но реальное – это пространство весов связей НС.

Вместе с тем, предполагавшееся в [22] исходное разнесение индуцируемых образов и понятий по шкале общности признаков не всегда правомочно и не в полной мере соответствует НС-принципу неформализованности обработки. Представляет интерес реализация индуктивного вывода при отсутствии априорных критериев отнесения признаков к частным или общим. Единственным критерием должна быть частота появления признака в наборе образов [4, 5], а сами признаки не должны быть априори формализованы [23]. Последнее в идеале предполагает, что НС сама определяет признаки исходя из своей структуры связей.

Кроме того, метод [22] решает только задачу порождения индуктивной гипотезы как прототипа понятия. Необходим следующий шаг – включение нового понятия в индивидуальное знание посредством установления его связей с уже наличествующими понятиями.

Исходя из приведенных выше соображений, актуальна разработка НС-метода, свободного

от отмеченных ограничений метода [22] и реализующего оба упомянутых этапа:

- порождение индуктивной гипотезы на основе обработки набора примеров;
- самостоятельное установление связи нового понятия с уже существующими.

В статье предложена соответствующая этим требованиям модель перехода от образного к понятийному мышлению на основе нейросетевого порождения понятий на алгебре фурье-дуальных операций и установления их связей при ассоциативной обработке наборов образов.

1. Подход к задаче и модель

1.1. Формулировка задачи

Для наглядности изложения рассмотрим формирование понятийного мышления на основе индуктивного вывода на примере обращения классического силлогизма «*Darii*», описывающего дедуктивный вывод на основе общего и частноутвердительного высказываний:

Общее высказывание или правило:

Человек (любой) смертен; (1.а)

Частное высказывание или случай:

Сократ – человек; (1.б)

Частный вывод или результат:

Сократ – смертен. (1.в)

В рамках примера индуктивный вывод может быть получен из дедуктивного (1) следующей перестановкой высказываний данного силлогизма:

Частное высказывание (случай (1.б)):

Сократ – человек; (2.а)

Частное высказывание (результат (1.в)):

Сократ – смертен; (2.б)

Общий вывод (правило (1.а)):

Все люди смертны. (2.в)

Как исходный силлогизм (1), так и индуктивный вывод (2) включают три термина, данные курсивом: меньший *S* (в примере – *Сократ*), средний *M* (*человек*) и больший *P* (*смертен*). Согласно классификации «образ – понятие» [2], меньший *S* и больший *P* термины суть образы, поскольку конкретны, а средний термин – понятие, так как представляет «человека вообще». Тогда, в рамках задачи индуктивного порождения понятия меньший термин *S* определим как индексный образ, а больший *P* – как индуцируемый, подлежащий преобразованию в новое понятие абстрагированием от

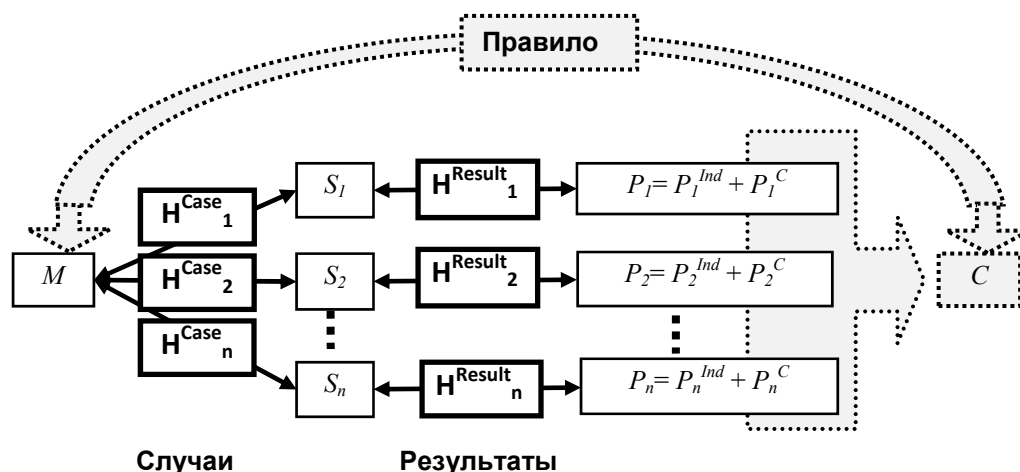


Рис. 1. Формирование понятийного мышления на основе индуктивного порождения понятий

M – наличествующее понятие (средний термин); S_i – индексные образы (меньшие термины), связанные случаями H^{Case_i} с понятием M ; P_i – индуцируемые образы (большие термины), связанные результатами H^{Result_i} с индексными образами S_i , и включающие в себя как частные P_i^{Ind} , так и общие P_i^C признаки; C – индуктивно формируемое понятие

частных (индивидуальных) признаков, существенных только для конкретного образа [2]. Далее преимущественно будем использовать эту терминологию.

Представим индуктивное порождение нового понятия и его связь с уже наличествующим в виде принципиальной схемы Рис. 1, в которой несколько выводов (2) объединены общим для них понятием – средним термином «человек» (Сократ – человек, Платон – человек, etc.).

Таким образом, задача в самом общем виде, еще не привязанном к методу ее решения, может быть сформулирована следующим образом:

Условия задачи:

1. Имеется набор запомненных интеллектуальной системой примеров частноутвердительных высказываний (в терминах Ч.С. Пирса – результатов): n пар связанных образов: индексных S_i (меньшие термины) и индуцируемых P_i (большие термины). Последние включают в себя как индивидуальные P_i^{Ind} , так и общие для всего набора P_i^C признаки.

2. Со всеми индексными образами S_i другими частноутвердительными высказываниями (в терминах Ч.С. Пирса – случаями) связано уже наличествующее понятие M (средний термин).

Требуется:

1. Найти общее P_i^C в индуцируемых образах P и сформировать новое понятие C .

2. Связать новое понятие C с уже имеющимся M – создать ранее не существовавшее новое правило, связывающее два понятия: C и M .

Предполагается: априорной информации о критериях различения общих P_i^C и индивидуальных P_i^{Ind} признаков в образах нет, единственный критерий – частота появления признака в наборе.

Поскольку и образы, и понятия суть внутренние репрезентации информации [2], то, рассматривая решение задачи в рамках нейросетевого подхода, примем, что все термины S_i , P_i (образы) и M (понятие) представлены в НС посредством паттернов [3]. P_i^{Ind} и P_i^C обозначим как субпаттерны. От механизма преобразования сенсорной информации в паттерны [24], как выходящего за рамки статьи, абстрагируемся.

Из условия отсутствия априорных критериев отнесения признаков к частным или общим следует, что статистические характеристики субпаттернов P_i^{Ind} и P_i^C не должны различаться. Последнее условие формализуем описанием паттернов как реализаций однородного в широком смысле случайного поля, т.е. стационарной случайной функции двух координат. Это положение резко ограничивает круг применимых методов, исключая из рассмотрения подходы, основанные на различиях заранее заданных характеристик паттернов из набора $\{P_i\}_{i=1}^n$.

Из Рис. 1 следует критерий отличия паттерна образа от паттерна понятия – структура связей. Первый связан только с одним паттерном образа – это структура связей типа «результат», а второй – с несколькими, эту структуру можно определить как «звезду случаев».

1.2. Подход к решению задачи: алгебра фурье-дуальных определяющих операций

Для удобства читателя кратко опишем алгебру фурье-дуальных операций, более развернутое изложение можно найти в работах [15, 16].

Определим алгебру как модель $\langle F(X), D, \odot, \oplus, O, U \rangle$, где $F(X)$ – множество элементов модели; X – универсальное множество, его элементы – x , $\odot$ и $\oplus$ – определяющие модель операции, определенные в классе треугольных норм (t-норм) и конорм [25], соответственно; D – операция, задающая дуальность определяющих модель операций; O и U – специальные константы, представляющие наименьший и наибольший элементы из $F(X)$; $[O, U]$ – решетка. Множество элементов модели $F(X)$ определим способом, формально совпадающим с определением нечеткого подмножества посредством функции принадлежности μ [26]:

$$F(X) = \{ \mu | \mu : X \rightarrow [0, 1] \}.$$

Если элементы модели – паттерны, т.е. функции двух аргументов, то элементы универсального множества – координаты (x, y) . Далее, где это возможно без потери смысла, для уменьшения громоздкости выражений будем использовать функции одной переменной x .

Операцию D определим как обобщение операции отрицания с интервала $[0, 1]$ на $[O, U]$:

$$D : [O, U] \rightarrow [O, U] \quad (3.a)$$

$$D(O) = U \quad (3.b)$$

$$D(U) = O$$

$$\forall A(x), B(x) \in F(X); \quad (3.c)$$

$$A(x) \geq B(x) \Leftrightarrow D(A(x)) \leq D(B(x)).$$

Определение (3) обобщает поэлементно определенную операцию отрицания на операции над операндами в целом, т.е. учитывает внутреннюю коррелированность информации как ее важнейший атрибут.

Операция преобразования Фурье F связывает функцию $A(x)$, удовлетворяющую условиям Дирихле, с ее фурье-образом $F(v)$, называемым также спектром:

$$F(A(x)) = F(v) = \int_{x_{min}}^{x_{max}} A(x) \exp(-j2\pi vx) dx, \quad (4)$$

где j – мнимая единица, v – частота (пространственная, если x – пространственная координата).

Преобразование Фурье удовлетворяет аксиоматическому определению операции, задающей дуальность (3). Аксиома ограниченности (3.b) в идеальном случае имеет силу в форме:

$$F(\delta(x)) = \text{Const}(v)$$

$$F(\text{Const}(x)) = \delta(v)$$

где $\delta(x)$ – δ -функция.

Для унимодальных элементов модели преобразование Фурье удовлетворяет аксиоме монотонности (3.c) в форме:

$$\forall A(x), B(x) \in F(X);$$

$$A_\alpha(x) \geq B_\alpha(x) \Leftrightarrow |F(A(x))|_\alpha \leq |F(B(x))|_\alpha,$$

где $\alpha \in [0, 1]$, A_α – α -срез A [31], $|F(A(x))|_\alpha$ – спектр амплитуд (модуль фурье-образа).

Если $A(x), B(x)$ не унимодальны, то условие невозрастания имеет силу не для них самих, но для определяемых согласно теореме Котельникова [27] их элементов $a(x), b(x)$ в форме:

$$\forall A(x), B(x) \in F(X);$$

$$a_\alpha(x) \geq b_\alpha(x) \Leftrightarrow |F(A(x))|_\alpha \leq |F(B(x))|_\alpha.$$

Приняв алгебраическое произведение в качестве поэлементно-определенной t-нормы $\odot$, получим свертку как t-конорму $\oplus$ (абстрактное сложение), фурье-дуальную произведению:

$$\begin{aligned} F(A(x) \oplus B(x)) &= F(A(x)) F(B(x)) = \\ &= F(A(x) * B(x)), \end{aligned}$$

где символ $*$ обозначает операцию свертки двух функций

$$\text{Conv}_{AB}(\eta) = A(x) * B(x) = \int_{x_{min}}^{x_{max}} A(x) B(\eta - x) dx.$$

Алгебра фурье-дуальных определяющих операций есть алгебра нечетких подмножеств, поскольку даже если исходные элементы модели – четкие, т.е. представляются отображением $X \rightarrow \{0, 1\}$, то результат свертки описывается как $X \rightarrow [0, 1]$, т.е. нечеток по определению [26].

Из определения вычитания как сложения с аддитивно противоположным элементом [28] следует, что операция корреляции $\otimes$ в алгебре имеет смысл вычитания (абстрактного) $\ominus$:

$$\begin{aligned}
 A(x) \odot B(x) &= A(x) * B(-x) = \\
 &= F(F(A(x))F(B(-x))) = \\
 &= \int_{x_{\min}}^{x_{\max}} A(x + \eta) B^*(x) dx = A(x) \otimes B(x),
 \end{aligned}$$

где астериск – символ комплексного сопряжения.

Преобразование Фурье (4), задающее дуальность, определяет диссипативность модели, включая необратимость вычислений, в том числе, в смысле операций сложения и вычитания:

$$\begin{aligned}
 A(x), B(x) &\in F(X); \\
 A(x) * B(x) &= C(x) \Rightarrow C(x) \otimes B(x) \neq A(x). \quad (5)
 \end{aligned}$$

1.3. Нейросетевая реализация алгебры

Алгебра адекватна двухслойной нейронной сети (Рис. 2) с весами связей, формируемыми по модели внешнего произведения: если паттерны $S(x, y)$ и $P(x, y)$ описывают состояние нейронных слоев S и R при обучении сети (эталонные паттерны), то веса связей в пространстве Фурье, в которых запомнены ассоциированные пары паттернов, описываются выражением:

$$H(v_x, v_y) = F^*(S(x, y))F(P(x, y)). \quad (6)$$

При предъявлении в слое S паттерна $S'(x, y)$ в слое R формируется паттерн

$$\begin{aligned}
 P'(x, y) &= F(F(S'(x, y)) \cdot H(v_x, v_y)) = \\
 &= \int_{y_{\min}}^{y_{\max}} \int_{x_{\min}}^{x_{\max}} S'(x + \eta, y + \zeta) S^*(x, y) P(x, y) dx dy = \\
 &= (S'(x, y) \otimes S(x, y)) * P(x, y) = \\
 &= (S'(x, y) \odot S(x, y)) \oplus P(x, y),
 \end{aligned}$$

т.е. сеть реализует модель гетеро-ассоциативной памяти (по определению последней).

Здесь и далее для более ясного изложения мы приняли единую систему координат для всех слоев НС и не отражаем в нотациях инверсию координат при двойном преобразовании Фурье.

1.4. Формирование индуктивной гипотезы

Принципиальная схема (Рис. 1) на этапе индуктивного обобщения может быть рассмотре-

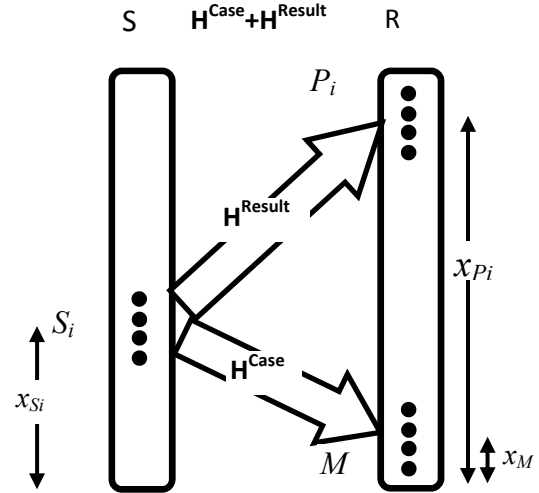


Рис.2. Принципиальная схема двухслойной нейросети

S и R – слои репрезентаций образов, S – индексных (меньших терминов), R – индуцируемых P_i и понятия M ; стрелки – однонаправленные связи случаев H^{Case} и результатов H^{Result} ; x_{S_i} , x_{P_i} , x_M – положение соответствующих ПВР при обучении сети

на как ассоциативная память, включающая две ассоциативные памяти: результатов $M \leftrightarrow S_i$ и случаев $P_i \leftrightarrow S_i$, связанные общими для них ПВР индексных образов S_i . Веса связей в предположении линейности запоминания весов формируются согласно (6): веса i -го результата:

$$\begin{aligned}
 H_i^{Result}(v_x) &= \\
 &= F(P_i(x))F^*(S_i(x))\exp(-j2\pi v_x(x_{P_i} - x_{S_i})), \quad (7)
 \end{aligned}$$

а i -го случая

$$\begin{aligned}
 H_i^{Case}(v_x) &= \\
 &= F(M(x))F^*(S_i(x))\exp(-j2\pi v_x(x_{M_i} - x_{S_i})), \quad (8)
 \end{aligned}$$

где координаты с нижними индексами описывают положение соответствующих эталонных паттернов в нейронных слоях при обучении сети.

Примем, что локализация паттернов при обучении сохраняется постоянной, т.е. $\forall i, j \in [1, n]: x_{S_i} = x_{S_j} = x_S$ и аналогично для P_i и M . Строго говоря, это допущение излишне, достаточно выполнения условия $\forall i, j \in [1, n]: x_{P_i} - x_{S_i} = x_{P_j} - x_{S_j}$, т.е. сохранения взаимного положения паттернов в обучающих парах $M \leftrightarrow S_i$ и $P_i \leftrightarrow S_i$, но оно позволит существенно упростить дальнейшие выкладки без потери общности рассмотрения.

Символы H для обозначения весов связей выбраны потому, что выражения (7) и (8) по формально-математическому определению – голограммы Фурье, вне зависимости от метода их реализации.

Веса обученной на n примерах сети:

$$H(v_x) = \sum_{i=1}^n (H_i^{Result}(v_x) + H_i^{Case}(v_x)). \quad (9)$$

Если обученной сети в слое S предъявляется индексный паттерн S_i , то в слое R формируется

$$\begin{aligned} F[F(S_i(x))H(v_x)] &= M(x+x_M) * [S_i(x) \otimes S_i(x)] + \\ &+ P_i(x+x_p) * [S_i(x) \otimes S_i(x)] + \\ &+ \sum_{i \neq j} M(x+x_M) * [S_i(x) \otimes S_j(x)] + \\ &+ \sum_{i \neq j} P_j(x+x_p) * [S_i(x) \otimes S_j(x)] \end{aligned} \quad (10)$$

Первое и третье слагаемые в (10) описывают паттерн $M_R(x)$, формируемый в области ПВР понятия M , а второе и четвертое – паттерн $P_R(x)$ в области индуцируемого образа P . При этом первое и второе слагаемые описывают восстановление эталонных ПВР в той мере, в которой автокорреляционная функция $[S_i(x) \otimes S_i(x)]$, играющая роль импульсного отклика системы, описывающего ее диссипативность (5), может быть аппроксимирована δ -функцией. Третье и четвертое накладываются на восстановленные эталонные паттерны как помеха.

Для анализа выражения (10) примем, что ПВР индексных образов в пределах последовательности $\{S_i\}_{i=1}^n$ частично коррелированы, т.е. могут быть представлены в виде:

$$\begin{cases} S_i(x) = S_i^C(x) + S_i^U(x) \\ \begin{cases} S_i^C(x) = m \cdot S_i(x) \\ S_i^U(x) = (1-m) \cdot S_i(x), \end{cases} \end{cases} \quad (11)$$

где коэффициент $m \in [0,1]$ описывает удельный вес коррелированного фрагмента в паттерне – как по площади, так и по амплитуде.

Соответственно, если S_i^U и S_j^U ортогональны, то коэффициент взаимной корреляции индексных паттернов $\rho_{ij}^S = m^2$. Тогда сумма второго и четвертого слагаемых в (10) примет вид:

$$\begin{aligned} P_R(x) &= P_i(x+x_p) * [S_i(x) \otimes S_i(x)] + \\ &+ \sum_{i \neq j} P_j(x+x_p) * [S_i(x) \otimes S_j(x)] = \\ &= P_i^C(x+x_p) * [S_i(x) \otimes S_i(x)] + \\ &+ P_i^{Ind}(x+x_p) * [S_i(x) \otimes S_i(x)] + \\ &+ \sum_{i \neq j} P_j(x+x_p) * [(S_i^C(x) + S_i^U(x)) \otimes (S_j^C(x) + S_j^U(x))] = \\ &= P_i^C(x+x_p) * [S_i(x) \otimes S_i(x)] + \\ &+ P_i^{Ind}(x+x_p) * [S_i(x) \otimes S_i(x)] + \\ &+ \rho_{ij}^S(n-1)P_j^C(x+x_p) * [S_i(x) \otimes S_i(x)] + \\ &+ \rho_{ij}^S \sum_{i \neq j} P_j^{Ind}(x+x_p) * [S_i(x) \otimes S_i(x)] + \\ &+ \sum_{i \neq j} P_j(x+x_p) * [S_i^U(x) \otimes S_j^U(x)] + \\ &+ 2 \sum_{i \neq j} P_j(x+x_p) * [S_i^C(x) \otimes S_j^U(x)] = \\ &= [1 + \rho_{ij}^S(n-1)] \cdot \{P_i^C(x+x_p) * [S_i(x) \otimes S_i(x)]\} + \\ &+ P_i^{Ind}(x+x_p) * [S_i(x) \otimes S_i(x)] + \\ &+ \rho_{ij}^S \sum_{i \neq j} P_j^{Ind}(x+x_p) * [S_i(x) \otimes S_i(x)] + \\ &+ \sum_{i \neq j} P_j(x+x_p) * [S_i^U(x) \otimes S_j^U(x)] + \\ &+ 2 \sum_{i \neq j} P_j(x+x_p) * [S_i^C(x) \otimes S_j^U(x)] = P_R^C(x) + P_R^N(x) \end{aligned} \quad (12)$$

В (12) мы приняли

$$[S_i^C(x) \otimes S_j^U(x)] = [S_i^U(x) \otimes S_j^C(x)].$$

Все слагаемые в (12) пространственно совпадают – налагаются друг на друга. Первое слагаемое описывает требуемый для решения задачи субпаттерн общих признаков, обозначенный выше как $P_R^C(x)$. Сумма второго – пятого слагаемых обозначена как помеха $P_R^N(x)$. Второе слагаемое $P_i^{Ind}(x+x_p) * [S_i(x) \otimes S_i(x)]$ описывает восстановленный субпаттерн частных признаков, его вклад в помеху не зависит от числа обучающих примеров n . Покажем, что третье – пятое слагаемые, зависящие от числа примеров n , влияние итоговой помехи $P_R^N(x)$ уменьшают с ростом обучающей выборки n .

Сравним модулированность восстановленных субпаттернов $P_R^C(x)$ и $P_R^N(x)$. В качестве интегральной оценки модулированности паттернов используем их дисперсии. Дисперсия

паттерна связана со значением его функции автокорреляции в начале отсчета $K_{ii}(0,0)$:

$$D_i = \frac{K_{ii}(0,0)}{L_x \cdot L_y}, \quad (13)$$

где L_x и L_y – размеры паттерна в предположении его прямоугольной формы.

Используем известное свойство функции корреляции суммы случайных процессов [29, 30]:

$$\begin{aligned} K\left(\sum_{i=1}^n X_i(x)\right) &= \sum_{i=1}^n K_{ii}(X_i(x)) + \sum_{i \neq j}^n K_{ij}(X_i(x), X_j(x)) = \\ &= nK_{ii}(X_i(x)) + \sum_{i \neq j}^n K_{ij}(X_i(x), X_j(x)), \end{aligned} \quad (14)$$

где $X_i(x)$ – паттерн как i -ая реализация случайного поля, K_{ij} и K_{ii} – функции взаимной корреляции i -го и j -го паттернов и автокорреляции i -го, соответственно.

Тогда, если $X_i(x)$ коррелированы и их размеры (площади паттернов и длины сигналов) постоянны, то из (12) и (13) следует:

$$D\left(\sum_{i=1}^n X_i(x)\right) = n^2 D_{ii}(X_i(x)). \quad (15)$$

Дисперсия требуемого для решения задачи паттерна, описываемого первым слагаемым в (12), согласно (14) будет зависеть от числа обучающих примеров n и коэффициента корреляции индексных образов ρ_{ij}^S следующим образом:

$$\begin{aligned} D_R^C &= D\left([1 + \rho_{ij}^S(n-1)] \cdot \{P_i^C(x+x_p) * [S_i(x) \otimes S_i(x)]\}\right) = \\ &= [1 + \rho_{ij}^S(n-1)]^2 D\{P_i^C(x+x_p) * [S_i(x) \otimes S_i(x)]\} = \\ &= (1 + \rho_{ij}^S(n-1))^2 D_{iR}^C, \end{aligned} \quad (16)$$

где D_{iR}^C – дисперсия члена в фигурных скобках, т.е. i -го субпаттерна в слое R с учетом ее изменения относительно дисперсии эталонного субпаттерна P_i^C в результате свертки с импульсным откликом схемы $[S_i(x) \otimes S_i(x)]$.

При $\rho_{ij}^S = 1$, т.е. полной коррелированности индексных образов, D_R^C зависит от числа примеров n квадратично, согласно (15).

Детальный анализ дисперсии помехи, описываемой остальными слагаемыми в (12), сопряжен с появлением и учетом большого числа ковариационных членов, существенно загромождающих итоговое выражение, но дающих

относительно небольшой вклад, которым на практике в силу нестрогости выполнения для большинства реальных паттернов гипотезы эргодичности [29] можно пренебречь. Поэтому ограничимся рассмотрением двух крайних случаев: полной коррелированности и, напротив, некоррелированности индексных образов.

При полной коррелированности индексных образов дисперсия помехи принимает вид:

$$\begin{aligned} D_R^N &= D\left(\sum_{i \neq j} P_i^{Ind}(x+x_p) * [S_i(x) \otimes S_i(x)] + \sum_{i \neq j} P_j^{Ind}(x+x_p) * [S_j(x) \otimes S_j(x)]\right) = \\ &= D\sum_i P_i^{Ind}(x+x_p) * [S_i(x) \otimes S_i(x)] = \\ &= nD_{iR}^{Ind} + n(n-1)D_{ijR}^{Ind} \end{aligned} \quad (17)$$

где $D(P_i^{Ind}(x+x_p) * [S_i(x) \otimes S_i(x)]) = D_{iR}^{Ind}$, а D_{ijR}^{Ind} – ковариация i -го и j -го субпаттернов частных признаков.

Отсюда в первом приближении зависимость отношения дисперсий восстановленных в слое R субпаттерна общих признаков и пространственно совпадающего с ним субпаттерна помехи от числа примеров n может быть оценена следующим выражением:

$$\begin{aligned} V(n) &= \frac{D_R^C}{D_R^N} = \frac{D_{iR}^C}{D_{iR}^{Ind}} \frac{(1 + \rho_{ij}^S(n-1))^2}{n \left(1 + (n-1) \frac{D_{ijR}^{Ind}}{D_{iR}^{Ind}}\right)} = \\ &= \frac{D_{iR}^C}{D_{iR}^{Ind}} \frac{(1 + \rho_{ij}^S(n-1))^2}{n(1 + \rho_{ij}^{Ind}(n-1))}, \end{aligned} \quad (18)$$

где ρ_{ij}^{Ind} – коэффициент корреляции i -го и j -го субпаттернов индивидуальных признаков индуцируемых образов.

Введем корреляционную оценку информационной емкости субпаттернов [31]:

$$\Omega^{Ind} = \frac{L_x^{Ind} L_y^{Ind}}{\pi r_R^2}, \quad (19)$$

где r_R – радиус корреляции с учетом влияния импульсного отклика. Тогда, как показано в [31]:

$$\rho_{ij}^{Ind} \approx \sqrt{\frac{2\kappa}{\Omega^{Ind}}},$$

где κ – коэффициент, зависящий от функции корреляции поля (процесса). Отсюда имеем:

$$V(n) = \frac{D_R^C}{D_R^N} \approx \frac{(1 + \rho_{ij}^S (n-1))^2}{n \left(1 + (n-1) \sqrt{\frac{2\kappa}{\Omega^{Ind}}} \right)}. \quad (20)$$

Зависимость отношения (20) от числа обучающих примеров n растет с ростом оценки информационной емкости субпаттернов индивидуальных признаков Ω^{Ind} . Зависимость отношения (20) от оценки информационной емкости субпаттернов именно индивидуальных признаков отражает механизм решения задачи – сглаживание с ростом n субпаттернов индивидуальных признаков как все более ровного фона, на котором выявляются субпаттерны общих признаков.

При полной некоррелированности индексных образов выражение (16) теряет зависимость от числа наложенных голограмм:

$$D_R^C = D_{iR}^C.$$

Дисперсия помехи, напротив, эту зависимость сохраняет, так как принимает вид:

$$\begin{aligned} D_R^N &= D \left(\sum_{i \neq j} P_i^{Ind} (x + x_p) * [S_i(x) \otimes S_i(x)] + \right. \\ &= D_{iR}^{Ind} + D \left(\sum_{i \neq j} P_i (x + x_p) * [S_j(x) \otimes S_j(x)] \right) + 2D_{12} \end{aligned} \quad (21)$$

где D_{12} – ковариация первого и второго слагаемых в первой строке (21).

Поскольку число максимумов кросскорреляционной функции, участвующих в формировании $P_i (x + x_p) * [S_j(x) \otimes S_j(x)]$, равно Ω^{Ind} , то в первом приближении можно принять:

$$D(P_i (x + x_p) * [S_j(x) \otimes S_j(x)]) \approx D_{iR}^{Ind}.$$

Отсюда видно, что если при $n=1$ отношение дисперсий

$$V(n) = \frac{D_R^C}{D_R^N} = \frac{D_{iR}^C}{D_{iR}^{Ind}},$$

то при $n=2$, даже без учета ковариационного члена $2D_{12}$:

$$V(n) \approx \frac{D_{iR}^C}{2D_{iR}^{Ind}},$$

т.е. снижается. Дальнейшее увеличение числа примеров n в силу (14) может вести к некоторому росту этого отношения, но при любом $n > 1$ оно будет меньше 1.

Таким образом, при частичной коррелированности индексных паттернов обучающей последовательности $\{S_i\}_{i=1}^n$ дисперсия активированного в слое R субпаттерна общих признаков растет быстрее, чем дисперсия субпаттерна, соответствующего субпаттернам частных признаков. При этом математическое ожидание обоих субпаттернов как суммы случайных полей (процессов) растет одинаково – пропорционально n [30]. Иными словами, общая модулированность субпаттерна общих признаков с ростом n сохраняется, а субпаттерна частных признаков – уменьшается, он превращается в равномерный фон. Следовательно, отношение дисперсий $V(n)$ может служить оценкой эффективности выделения субпаттерна общих признаков относительно субпаттерна частных по критерию их модулированности.

Биологические системы эффективно решают задачу выделения общих фрагментов в последовательности образов [32]. Возможно, что метод ее решения основан на реализации преобразования Фурье (4) нейронными структурами рецептивных полей [17-21].

Рассмотрение суммы первого и третьего членов в (10), т.е. восстановленного паттерна M_R , здесь опустим, так как и анализ, и результат аналогичны вышеприведенному с тем, что в силу отсутствия в паттерне M различающихся фрагментов выражения существенно упрощаются.

1.5. Формирование правила

Выше дана модель первого шага перехода от образного мышления к понятийному – индуктивного формирования паттерна нового понятия. Ниже покажем решение второго этапа.

Для формирования правила необходимо сформировать и записать связи восстановленных в слое S паттернов: наличествующего понятия M и индуктивной гипотезы (12). Строго говоря, эти паттерны уже связаны посредством связей (9). Но использование этих связей предполагает промежуточный шаг – активацию в слое S паттерна S^C , т.е. построение ассоциативной цепочки «понятие → образ → понятие». С точки зрения реализации «чистого» понятийного, т.е. абстрактного, мышления желательно установление правила как прямой ассоциации «понятие → понятие», свободной от промежуточных ассоциаций с конкретными образами.

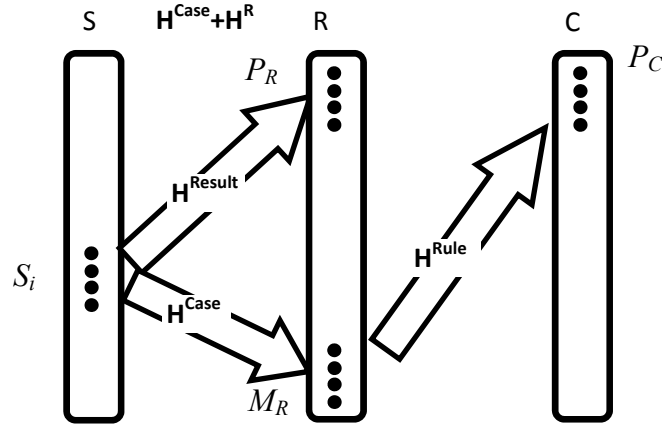


Рис.3. Принципиальная схема трехслойной нейросети, развивающей двухслойную НС (Рис.2) для формирования правила

С-слой репрезентации паттерна нового понятия P_C ; стрелки – связи слоев

Для этого добавим к двухслойной НС (Рис. 2), использованной для формирования индуктивной гипотезы, еще один слой С, получив трехслойную НС, представленную на Рис. 3.

Если паттерны M_R и P_R (10) в слое R активированы i -м индексным паттерном S_i , то веса связей слоев $R \leftrightarrow C$ в пространстве Фурье согласно (6)

$$\begin{aligned} H_i^{Rule} = & F(S_i(x))F(M_R(x))F^*(S_i(x))F^*(S_i(x)) \cdot \\ & \cdot F^*(P_R(x))F(S_i(x))\exp(-j2\pi v_x(x_M - x_P)) + \\ & + F^*(S_i(x))F^*(M_R(x))F(S_i(x))F(S_i(x)) \cdot \\ & \cdot F(P_R(x))F^*(S_i(x))\exp(j2\pi v_x(x_P - x_M)) \end{aligned} \quad (22)$$

Если отношения дисперсий $V(n)$ для восстановленных в слое R паттернов M_R и P_R достаточны для того, чтобы можно было пренебречь членами (12), описывающими помеху P_R^N для P_R (и аналогичными членами для M_R), то при предъявлении связям (22) паттерна понятия M (в слое R), второе слагаемое в (22) восстановит в слое С паттерн нового понятия P_C

$$\begin{aligned} P_C &= F[F(M(x))H^{Rule}(v)] = \\ &= P_R^C(x + x_P)^* \left[\begin{aligned} & (M(x) * (S_i(x) \otimes S_i(x))) \otimes \\ & (M(x) * (S_i(x) \otimes S_i(x))) \end{aligned} \right] \end{aligned} \quad (23)$$

Аналогично, при предъявлении в слое R паттерна нового понятия P_C , первое слагаемое в (22) восстановит в слое С паттерн ранее существовавшего понятия M .

1.6. Потери информации при порождении понятия

Выражение в квадратных скобках в (23) описывает импульсный отклик системы, т.е. оценивает её диссипативность. Его влияние проявляется в увеличении радиуса корреляции r_R восстановленного паттерна понятия (23) относительно радиуса корреляции эталонного субпаттерна общих признаков P^C . Увеличение радиуса корреляции, в свою очередь, ведет к снижению оценки информационной емкости порождаемого понятия

$$\Omega^C = L_x^C L_y^C / \pi r_R^2$$

относительно информационной емкости субпаттерна общих признаков P^C , на основе которого это понятие сформировано.

Минимизация импульсного отклика и, соответственно, потерь информации, как видно из (23), возможна в пространстве Фурье на обоих этапах работы сети:

1. На этапе обучения НС примерами – выбором условий записи связей (7) и (8). В частности, методом «разбеливания» совместных спектров, т.е. выполнением условий:

$$\begin{aligned} F(S_i(x))F^*(S_i(x)) &= Const(v) \\ F(P_i(x))F^*(P_i(x)) &= Const(v), \end{aligned} \quad (24)$$

Тогда $(S_i(x) \otimes S_i(x)) = \delta(x)$ и потерь информации не происходит.

2. На этапах формирования правила (22) и предъявления понятия M (23) разбеливанием его спектра амплитуд (модуля фурье-образа), т.е.

$$|F(M(x))| = \text{Const}(v). \quad (25)$$

Совместное выполнение (24) и (25) обеспечивает равенство импульсного отклика дельта-функции и, согласно (22) и (23), связь и взаимное восстановление паттернов понятий без потерь информации (без потерь мелких деталей в паттернах).

2. Моделирование

Моделировалась работа НС в предположении линейности динамического диапазона регистрирующих сред для записи весов связей. Предполагалось невыполнение свойства инвариантности к сдвигу преобразования Фурье, что соответствует регистрации связей нейронных слоев в объеме. Моделировались оба этапа задачи:

- индуктивное порождение паттернов гипотезы понятия;
- установление связи нового понятия со старым и восстановление паттерна нового понятия паттерном ранее наличествовавшего.

Численные эксперименты были проведены для двух вариантов:

1. Для выполнения требования на отсутствие иных, кроме частоты появления в обучающей выборке, критериев отнесения признаков к общим или частным, паттерны терминов M , S_i и P_i представлялись реализациями одного стационарного случайного процесса;

2. Для наглядности обрабатывались изображения, иллюстрирующие моделируемый пример.

Поскольку структуры биологического мозга имеют объемную организацию, то при моделировании предполагалось, что свойство инвариантности преобразования Фурье к сдвигу силы не имеет.

2.1. Моделирование на примере паттернов как реализаций стационарного случайного процесса

Паттерны представлены реализациями стационарного случайного процесса с экспоненциальным спектром амплитуд и случайным спектром фаз с нулевым матожиданием и дисперсией 2π , радиусом корреляции по уровню $0.5 \tau=11$ отсчетов. Длины всех паттернов

составляли 1024 отсчета, т.е. оценка информационной емкости паттернов $\Omega = L/r = 93$. Индуцируемые паттерны P_i моделировались для двух вариантов:

– пространственного разнесения субпаттернов P_i^C и P_i^{Ind} – в этом случае их длины были равны $L=512$ отсчетов, т.е. информационная емкость обоих субпаттернов оценивалась $\Omega^C = \Omega^{Ind} = L^C/r = L^{Ind}/r = 46.5$;

– пространственного наложения субпаттернов P_i^C и P_i^{Ind} , их длины $L^C=512$ и $L^{Ind}=1024$, т.е. информационная емкость субпаттернов: общих признаков $\Omega^C = L^C/r = 46.5$ и частных признаков $\Omega^{Ind} = L^{Ind}/r = 93$.

Субпаттерн общих признаков P_i^C , одинаковый для всех реализаций, занимал боковые области каждой реализации P_i (от 0 до 255 и от 768 до 1024 отсчетов), а субпаттерн P_i^{Ind} при отсутствии наложения – центральную область (256 – 767) отсчетов.

Моделировалась запись весов связей: линейная и с вариантами разбеливания для исключения потерь информации – выполнением условий (24) и (25).

На Рис. 4 дан один из паттернов P_i , использовавшийся для обучения сети, и паттерны, формируемые согласно (23) при выполнении условий (24) и (25). Обобщение проведено на пятнадцати примерах при пространственном разнесении и наложении субпаттернов P_i^C и P_i^{Ind} и полной коррелированности индексных паттернов.

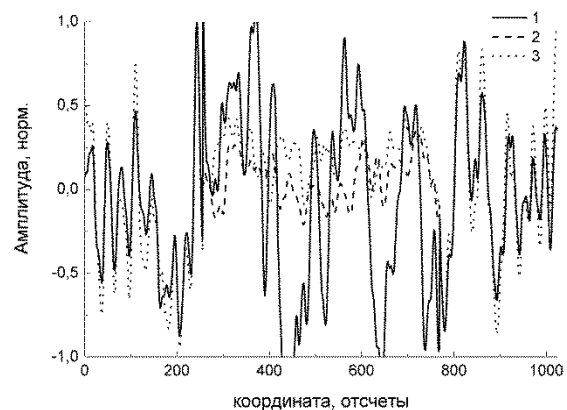


Рис. 4. Паттерны: 1 – индуцируемый P_i (первый пример при обучении, второй, третий – восстановленные по результатам обучения на 15 примерах); 2 – с пространственным разнесением P_i^C и P_i^{Ind} ; 3 – с их наложением

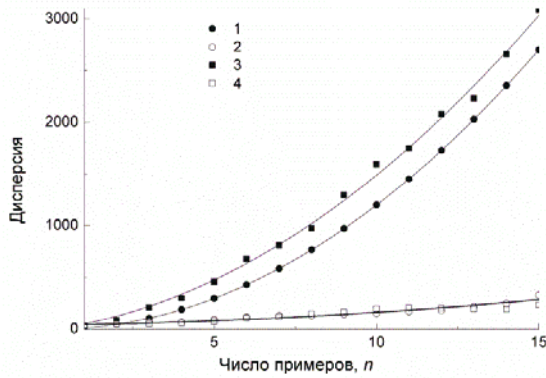


Рис. 5. Зависимости дисперсий восстановленных паттернов P_R^C (1, 3) и P_R^{Ind} (2, 4) от числа обучающих примеров n для 1, 2 — пространственного разнесения и 3, 4 — наложения P_i^C и P_i^{Ind}

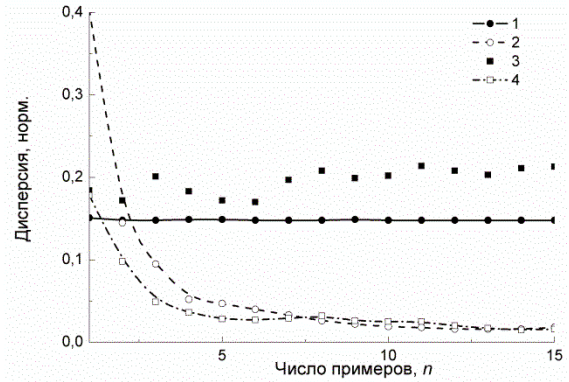


Рис. 6. Зависимости нормированных на максимальную амплитуду дисперсий от числа примеров (обозначения кривых как на Рис. 5)

На Рис. 5 даны зависимости дисперсий восстановленных субпаттернов P_R^C и P_R^{Ind} в зависимости от числа обучающих примеров n для вариантов пространственного разнесения и наложения P_i^C и P_i^{Ind} . На Рис. 6 даны зависимости дисперсий этих паттернов от числа примеров n , нормированные на максимальную амплитуду — такая нормировка соответствует согласованию амплитуды паттерна с рабочим участком активационной функции нейрона.

Аппроксимация кривых на Рис. 5 для случая без наложения:

$$D^C = -0.56 + 0.29N + 25.157N^2$$

$$D^{Ind} = 14.77N + 0.25N^2,$$

т.е. фактически квадратичная для P_R^C , согласно (15), и линейная для P_R^{Ind} , согласно (14). Для случая с наложением P_i^C и P_i^{Ind} зависимость дисперсии от числа примеров сохраняет преимущественно квадратичный характер для P_R^C и существенно линейный для P_R^{Ind}

$$D^C = -80.37 + 85.22N + 22.99N^2$$

$$D^{Ind} = -4.26 + 22.13N - 0.44N^2.$$

Таким образом, из Рис. 4-Рис. 6 видно, что эффект обобщения на примерах работает и в варианте пространственного наложения субпаттернов P_i^C и P_i^{Ind} . Эффективность выделения общих признаков растет с ростом числа обучающих примеров согласно (20).

Для подтверждения роли коррелированности индексных образов на Рис. 7 даны примеры для разных коэффициентов корреляции индексных образов S_i при пространственном раз-

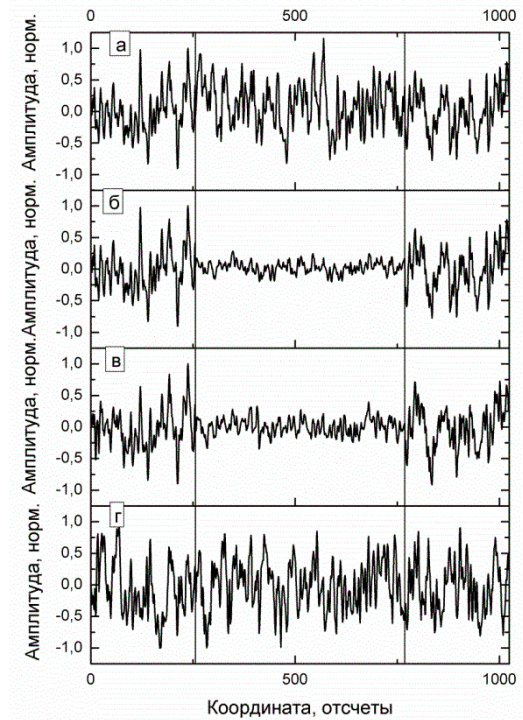


Рис. 7. Паттерны: а — индуцируемый P_I (первый пример при обучении), б — P_R , восстановленный при полной коррелированности индексных паттернов $\rho_{ij}^S = 1$; в — P_R , восстановленный при частичной коррелированности индексных паттернов $\rho_{ij}^S = 0.25$; г — P_R , восстановленный при полной некоррелированности индексных паттернов $\rho_{ij}^S = 0$

несении субпаттернов P_i^C и P_i^{Ind} (обучение велось также на 15-ти примерах).

Рис. 4-Рис. 7 наглядно иллюстрируют работоспособность модели:

- при полной коррелированности индексных паттернов (Рис. 4 и Рис. 7 б) в восстановленном паттерне P_R^C (от 0-го до 255-го и от 768-го до 1024-го отсчетов) как модулированность, так и форма реализации сохраняются практически постоянными, паттерн P_R^C – неискаженный эталонный субпаттерн P_i^C ; в области P_R^{Ind} (от 256-го до 767-го отсчетов) модулированность существенно уменьшилась, максимальные и минимальные значения амплитуды приблизились к среднему значению, при этом восстановленный паттерн P_R^{Ind} ни одному из эталонных паттернов набора $\{P_i^{Ind}\}_{i=1}^n$ не соответствует.

- аналогичные результаты наблюдаются при пространственном наложении субпаттернов P_i^C и P_i^{Ind} с тем отличием, что восстановленный паттерн P_R^C не в точности совпадает с эталонным P_i^C ;

- при частичной коррелированности индексных паттернов (Рис.7 в) $\rho_{ij}^S = 0.25$ метод работает с несколько меньшей эффективностью: можно заметить как снижение отношения дисперсий, так и искажения восстановленного паттерна P_R^C относительно эталонного P_i^C ;

- при полной некоррелированности индексных паттернов (Рис.7.г) $\rho_{ij}^S = 0$ метод не работает.

2.2. Моделирование наглядного примера

Моделировался наглядный пример на основе силлогизма «Darii»:

Случаи: Эпикур (S_1) – человек (M),
Сократ (S_2) – человек (M), etc.

Результаты: Эпикур (S_1) – смертен (P_1),
Сократ (S_2) – смертен (P_2), etc.

Правило: Человек вообще (M) смертен (P^C).

Число обучающих примеров $n=5$, использованные паттерны и их связь в виде ассоциативной сети, представлены на Рис. 8. Субпаттерн общих признаков P^C представлял собой стили-

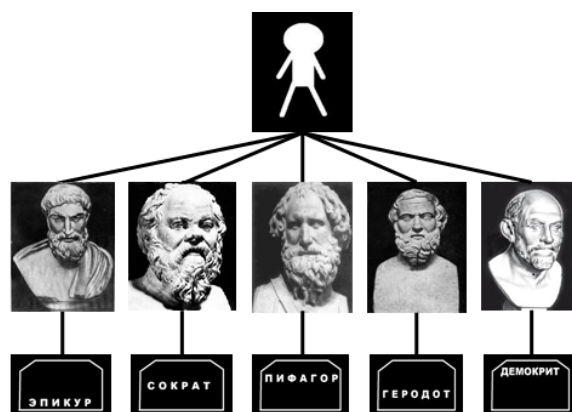


Рис. 8. Пример ассоциативной сети и паттернов, использованных в численном эксперименте

зованное изображение гробницы (шестиугольник), а субпаттерны частных – надписи на ней, выполненные печатным шрифтом. В связи с тем, что индексные образы (скульптурные изображения) очень слабо коррелированы между собой, в эксперименте формирование паттерна результата индуктивного обобщения P_R^C выполнялось по модели (10) – (20) не в слое R , а непосредственно в слое C . Эта модификация не затрагивает принципиальной сути метода, так как механизм работы тот же самый, основан на свойстве корреляционной функции (дисперсии) суммы (14).

Результаты представлены на Рис. 9. На Рис. 10 даны теоретические и экспериментальные (измеренные) зависимости дисперсий от числа примеров n .

Особенность данного примера в том, что достаточно велика только информационная емкость индексных паттернов, а паттерны, индуцируемые и понятия имеют, очевидно, низкую информационную емкость. Более того, субпаттерны частных признаков (надписи) явно имеют периодические составляющие и частично коррелированы (одинаковые буквы и их элементы в разных надписях). Но слабая взаимная коррелированность индексных образов обусловила и слабую кросскорреляционную (ковариационную) составляющую в восстановленных

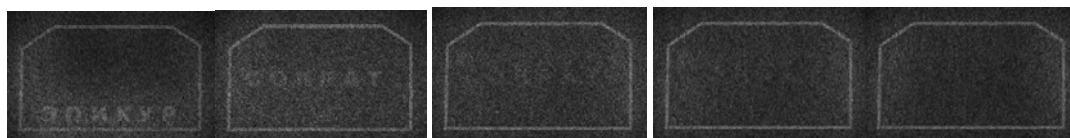


Рис. 9. Результаты численного эксперимента – восстановленный паттерн (гипотеза понятия) в зависимости от числа примеров (1 – 5)

паттернах. Вместе с тем, в примере, строго говоря, не выполнено требование на отсутствие априорных критериев отнесения признаков к общим или частным, так как стилизованное изображение гробницы по своим характеристикам отличается от надписей.

Заключение

Представлена модель перехода от образного к понятийному мышлению на основе индуктивного формирования понятий при обработке образов. Модель развивает ранее показанные возможности алгебры фурье-дуальных операций по реализации на ассоциативной памяти монотонной и немонотонной логик в классе нечетко-значимых в плане реализации правдоподобных рассуждений. Модель соответствует критерию биологической мотивированности в части обработки информации, представленной паттернами внутренней репрезентации, и обучаемости вместо формального программирования. Судя по результатам нейрофизиологических исследований [17-21], есть основания полагать и возможное соответствие положенной в основу модели алгебры фурье-дуальных операций реальным когнитивным механизмам, работающим в живых системах. Модель свободна от требований на априорную информацию о критериях различения частных и общих признаков образов, в том числе, от ограничения на разнесение частных и общих признаков по шкале их общности.

Вместе с тем, данная модель предполагает устойчивость взаимной пространственной локализации паттернов внутренней репрезентации наличествующего понятия и индуцируемых образов в рамках обучающего набора примеров. Поскольку наличествующее понятие и индуцируемые образы (большие термины) связаны между собой индексными образами (меньшими терминами), то модель позволяет реализовать запоминание и восстановление временной последовательности индексных образов, характерной для биологической памяти и используемой живыми системами для анализа и поиска закономерностей в цепочках событий [32].

Модель соответствует тезису когнитивного подхода об обусловленности свойств модели обработки информации свойствами ее материального носителя [33]. Эта обусловленность конкретизирована зависимостью оценки ин-

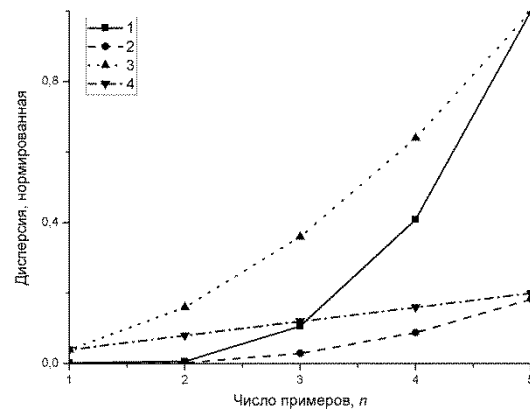


Рис. 10. Теоретические и экспериментальные зависимости дисперсии от числа примеров n

1, 3 – дисперсия ПВР общих признаков экспериментальная (1) и теоретическая (3), 2, 4 – дисперсия ПВР индивидуальных признаков экспериментальная (2) и теоретическая (4)

формационной емкости индуктивно порождаемого понятия от характеристик сенсорных трактов, условий записи весов связей и свойств регистрирующих сред, ведущих к дополнительной фильтрации. Дополнительная фильтрация при обучении и восстановлении ведет к изменению радиуса корреляции запомненных и восстановленных паттернов и, тем самым, изменению их информационной емкости.

Авторы считают приятным долгом выразить благодарность проф. Кузнецову О.П., проф. Фоминых И.Б. и с.н.с., к.т.н. Ройзензону Г.В. за обсуждения и критические замечания, способствовавшие формированию развиваемого подхода, а также Рецензенту за критические замечания, направленные на улучшение подачи материала.

Литература

1. Кузнецов О.П. Быстрые процессы мозга и обработка образов // Новости искусственного интеллекта. 1998. №2. <http://raai.org/library/ainews/1998/2/DISTR.ZIP>
2. Голицын Г.А., Фоминых И.Б. Нейронные сети и экспертные системы: перспективы интеграции // Новости искусственного интеллекта. М.: АИИ, 1996, №4. С.121-145.
3. Борисюк Г.Н., Борисюк Р.М., Казанович Я.Б., Иваницкий Г.Р. Модели динамики нейронной активности при обработке информации мозгом – итоги десятилетия // Успехи Физических наук. 2002. Т.172. №10. С.1189-1214.
4. Hayes B.K., Heit E., Swendsen H. Inductive reasoning // Wiley Interdisciplinary Reviews: Cognitive Science. 2010. V. 1, Issue 2. P. 278-292.

5. Вагин В.Н., Головина Е.Ю., Загорянская А.А., Фомина М.В. Достоверный и правдоподобный вывод в интеллектуальных системах. Изд. 2-е – М.: Физматлит. 2008.
6. Барский А.Б. Логические нейронные сети. М.:Бином, 2007. 351 с.
7. Финн В.К. Индуктивные методы Д.С. Милля в системах искусственного интеллекта. Часть 1 // Искусственный интеллект и принятие решений. 2010. №3. С. 3-21.
8. Финн В.К. Индуктивные методы Д.С. Милля в системах искусственного интеллекта. Часть 2 // Искусственный интеллект и принятие решений. 2010. №4. С.14-40.
9. Забейайло М.И. О некоторых оценках сложности вычислений в ДСМ методе. Часть I // Искусственный интеллект и принятие решений. 2015. №1. С.3-17.
10. Забейайло М.И. О некоторых оценках сложности вычислений в ДСМ методе. Часть II // Искусственный интеллект и принятие решений. 2015. №2. С.3 – 17.
11. Patrick Cavanagh. Holographic and trace strength models of rehearsal effects in the item recognition task // *Memory & Cognition* 1976. Vol. 4. Issue 2. P.186-199
12. Rosenblatt F. Principles of neurodynamics // Washington: Spartan Books, 1962.
13. Minsky M., Papert S. Perceptrons // Cambridge: MIT Press, 1969.
14. Anshelevich V.V., Amirikian B.R., Lukashin A.V., Frank-Kamenetskii M.D. On the Ability of Neural Networks to Perform Generalization by Induction // *Biological Cybernetics*, 1989. V. 61. P. 125 – 128.
15. Павлов А.В. Математические модели оптических методов обработки информации // *Известия Академии Наук. Серия: Теория и системы управления*. 2000. №3. С.111-118.
16. Павлов А.В. Об алгебраических основаниях голографической парадигмы в искусственном интеллекте: алгебра Фурье-дуальных операторов // V-я Международная научно-практическая конференция «Интегрированные модели и мягкие вычисления в искусственном интеллекте», 28-30 мая 2009 г., Коломна. Труды конференции.- М.Физматлит, 2009. Т.1, с.140-148.
17. Глезер В.Д., Иванов В.А., Щербач Т.А. Ответ рецептивных полей зрительной коры кошки на сложные стимулы // *Физиологический журнал СССР*. 1972. Т. 58. № 3.
18. Глезер В.Д., Иванов В.А., Щербач Т.А. Исследование рецептивных полей нейронов зрительной коры кошки как фильтров пространственных // *Физиологический журнал СССР*. 1973. Т. 59. № 2.
19. Глезер В.Д., Дудник К.Н., Куперман А. М., Лушина Л. И., Невская А.А., Подвитин Н.Ф., Праздникова И.В. Зрительное опознание и его нейрофизиологические механизмы. — Л.: Наука, 1975.
20. Глезер В.Д. Зрение и мышление. СПб.: Наука, 1993. 341 с.
21. Глезер В.Д. О роли пространственно-частотного анализа, примитивов и межполушарной асимметрии в опознании зрительных образов // *Физиология человека*. 2000. Т. 26. №5. С. 145-150.
22. Павлов А.В. Реализация правдоподобных выводов на нейросетях со связями по схеме голографии Фурье // *Искусственный интеллект и принятие решений*. 2010. №1. С.3-14.
23. Кобринский Б.А. Значение визуальных образных представлений для медицинских интеллектуальных систем // *Искусственный интеллект и принятие решений*. 2012. №3. С.3-8.
24. Marr D. Vision: a computational investigation into the human representation and processing of visual information. Freedman Hill. 1982. 397 p.
25. Mesiar R., Pap E. Different interpretations of triangular norms and related operations // *Fuzzy Sets and Systems*. 1998. V.96. P.183 – 189.
26. Аверкин А.Н., Батыршин И.З. и др. Нечеткие множества в моделях управления и искусственного интеллекта / Под ред. Д.А. Поспелова. М.: Наука. 1986.
27. Котельников В. А. О пропускной способности «эффира» и проволоки в электросвязи // *Успехи физических наук*, 2006. № 7. С. 762-770.
28. Dubois D., Prade H. Fuzzy numbers: an overview // in: Ed. by J.C.Bezdek, *Analysis of Fuzzy Information*.- Boca Raton, FL, 1987. V.1. P. 3-39.
29. Яглом А.М. Корреляционная теория стационарных случайных функций. Л.: Гидрометеоиздат, 1981. 280 С.
30. Вентцель Е.С. Теория вероятностей: Учеб. для вузов, 6-е изд. стер. М.: Высш. шк. 1999. 576 с.
31. Шубников Е.И. Отношение сигнал/помеха при корреляционном сравнении изображений // *Оптика и спектроскопия*. 1987. Т.62. №.2. С. 450-456.
32. Foster D.J., Wilson M.A. Reverse replay of behavioural sequences in hippocampal place cells during the awake state // *Nature*, 2006. V. 440. PP. 680-3.
33. Кузнецов О.П. Когнитивная семантика и искусственный интеллект // *Искусственный интеллект и принятие решений*. 2012. №4. С.32-42.

Павлов Александр Владимирович. Доцент университета ИТМО. Окончил Ленинградский институт точной механики и оптики в 1980 году. Доктор физико-математических наук. Автор более 100 печатных работ. Область научных интересов: искусственный интеллект, обработка образов, голография. E-mail: pavlov@phoi.ifmo.ru

Кочетков Павел Вадимович. Магистрант университета ИТМО. Окончил бакалавриат университета ИТМО в 2015 году. Область научных интересов: искусственный интеллект, обработка образов, голография. E-mail: pavlov@phoi.ifmo.ru

Algebra of fourier-dual operations: conceptual thinking generation

A.V. Pavlov, P.V. Kochetkov

Abstract. Two steps of conceptual thinking on the figurative one basis generation have been considered: the first step is inductive concepts generation by set of examples processing, and the second one is connections of the generated concepts with the yet existed ones establishing. The task has been considered in the framework of biologically inspired approach, i.e. both figurative and conceptual forms of information are represented and processed by the patterns of inner representation. A model based on the algebra of Fourier-dual operations for the both steps of the task solving has been proposed. Inductive concepts generation is implemented by sub-pattern of common features in the generalized on a set of examples pattern revealing under the criteria of the sub-patterns mutual correlation. A measure for revealing efficiency in dependence on both: a number of examples in the learning set and the sub-pattern of individual features information measures is proposed. Three layered neural network for the model implementation has been proposed. Computer simulations in the framework of “Dari” syllogism inversion results have been demonstrated.

Keywords: Figurative thinking, Conceptual thinking, Associative memory, Induction, Inductive concepts generation, Fourier-duality, Correlation, Neural networks, Pattern of inner representation

References

1. Kuznetsov O.P. Bystrye processy mozga i obrabotka obrazov // Novosti iskusstvennogo intellekta. 1998. №2. <http://raai.org/library/ainews/1998/2/DISTR.ZIP>
2. Golitsyt G.A., Fominyh I.B. Neironnye seti i ekspertnye systemy: perspektivy integratsii // Novosti iskusstvennogo intellekta. 1996, №4. P.121-145.
3. Borisyuk G N, Borisyuk R M, Kazanovich Ya B, Ivanitskii G R Models of neural dynamics in brain information processing — the developments of ‘the decade’ // Phys. Usp. 451073–1095 (2002); DOI: 10.1070/PU2002v045n10ABEH001143.
4. Hayes B.K., Heit E., Swendsen H. Inductive reasoning // Wiley Interdisciplinary Reviews: Cognitive Science.2010. V. 1, Issue 2. P. 278-292.
5. Vagin, V.N., Golovina, E.Yu., and Zagoryanskaya, A.A., Dostoverniy i pravodopodobnyi vyvod v intellektual’nykh sistemakh (Reliable and Probable Conclusion in Intelligent Systems), Vagin, V.N. and Pospelov, D.A., Eds., Moscow; FIZMALT, 2008..
6. Barskii A.B. Logicheskie neironnye seti.- M.: Binom, 2007. 351 P.
7. V. K. Finn J.S. Mill’s inductive methods in artificial intelligence systems. Part II // Scientific and Technical Information Processing December 2011, Volume 38, Issue 6, pp 385-402.
8. V.K. Finn J.S. Mill’s inductive methods in artificial intelligence systems. Part II // Scientific and Technical Information Processing December 2012, Volume 39, Issue 5, pp 241-260.
9. M.I. Zabezhailo. To the Computational Complexity of Hypotheses Generation in JSM-method. Part I. // Iskusstvennyy intellekt i prinyatie resheniy. 2015. №1. P.3 – 17.
10. M. I. Zabezhailo To the Computational Complexity of Hypotheses Generation in JSM-method. Part I // Iskusstvennyy intellekt i prinyatie resheniy. 2015. №2. P.3 – 17.
11. Patrick Cavanagh Holographic and trace strength models of rehearsal effects in the item recognition task // Memory & Cognition 1976. Vol. 4. Issue 2. P.186-199
12. Rosenblatt F. Principles of neurodynamics // Washington: Spartan Books, 1962.
13. Minsky M., Papert S. Perceptrons // Cambridge: MIT Press, 1969.
14. V.V. Anshelevich, B.R. Amirikian, A.V. Lukashin, and M. D. Frank-Kamenetskii On the Ability of Neural Networks to Perform Generalization by Induction // Biological Cybernetics, 1989. V. 61. P. 125 – 128.
15. Pavlov A.V. Mathematical Models of Optical Methods in Data Processing// Journal of Computer and Systems Sciences International. 2000. №3. P.441 – 448.
16. Pavlov A.V. Ob algebraicheskikh osnovaniyakh golograficheskoi paradigmy v iskusstvennom intellekte: algebra Fourier-dual’nykh operatorov // V mezhdunarodnaya nauchno-prakticheskaya konferentsiya: Integrirovannyye modeli i myagkie vychisleniya v iskusstvennom intellekte, 28-30.05.2009. Kolomna. Trudy konferentsii. – M.Fizmatlit,2009. T.1, P.140-148.
17. Glezer V.D., Ivanov V.A., Sherbach T.A. Otvet retseptivnykh polei neironov zritel’noi kory koshki na sloznye stimuly // Phys. Journal USSR. 1972. V.58. №3.
18. Glezer V.D., Ivanov V.A., Sherbach T.A. Issledovanie retseptivnykh polei neironov zritel’noi kory koshki kak filtrov prostranstvennykh chastot // Phys. Journal USSR. 1973. V.59. №2.
19. Glezer V.D., Dudnik K.N. et al. Zritel’noe opoznanie I ego neirofiziologicheskie mekhanizmy. L.: Nauka, 1975.
20. Glezer V.D. Zrenie I myshlenie. SPb.: Nauka, 1993. 341 P.

21. Glezer V.D. The Role of Spatial-Frequency Analysis, Primitives, and Interhemispheric Asymmetry in the Identification of Visual Images // *Human Physiology*. 2000. V.26. №.5. P.636 – 640.
22. Pavlov A.V. Realizatsia pravdopodobnykh vyvodov na neurosetiyah so svyaziami po scheme holographii Fourier // *Iskusstvennyy intellekt i prinyatie resheniy*. 2010, №1. P.3-14.
23. B. A. Kobrinskii The significance of visual-image presentations for medical intelligent systems // *Scientific and Technical Information Processing* December 2013, V. 40, Issue 6, pp 337-341.
24. Marr D. Vision: a computational investigation into the human representation and processing of visual information. Freedman Hill, 1982. 397 p.
25. Mesiar R., Pap E., Different interpretations of triangular norms and related operations // *Fuzzy Sets and Systems*. 1998. V.96. P.183 – 189.
26. Averkin A.N., Batyrshin I.Z., et.al. Nechetkie mnozhestva v modelyakh upravleniya i iskusstvennogo intellekta / Under ed. by Pospelov D.A. – M.: Nauka, 1986.
27. V.A. Kotelnikov On the transmission capacity of 'ether' and wire in electric communications // *Physics-Uspekhi (Advances in Physical Sciences)* 49 736–744 (2006) DOI: 10.1070/PU2006v049n07ABEH006160
28. Dubois D., Prade H. Fuzzy numbers: an overview // in: Ed. by J.C.Bezdek, *Analysis of Fuzzy Information*.- Boca Raton, FL, 1987. V.1. P. 3 – 39.
29. Yaglom. A.M Korrelyatsionnaya teoriya stacionarnykh sluchaynykh funktsiy // L.: Gidrometeoizdat, 1981. 280 P.
30. Ventsel E.S. Teoriya veroyatnostey. M.: Vyshaya shkola, 1999.
31. Subnikov E.L. Signal to noise ratio under images correlation comparison // *Optics and Spectroscopy*. 1987. V.62. i.2. P. 450 – 456.
32. Foster D.J., Wilson M.A. Reverse replay of behavioural sequences in hippocampal place cells during the awake state // *Nature*, 2006. V. 440. PP. 680-3.
33. O.P.Kuznetsov Cognitive semantics and artificial intelligence // *Scientific and Technical Information Processing*, December 2013, Volume 40, Issue 5, pp 269-276.

Alexander V.Pavlov. Docent of the ITMO University. Graduated from the Leningard Institute of Fine Mechanics and Optics in 1980. Dr. Sc. in Physics and Mathematics. Has more than 100 publications in indexed journals. Area of interests: Artificial Intelligence, Cognitive Systems, Information Processing, Pattern Recognition, Holography.
E-mail: Pavlov@phoi.ifmo.ru

Pavel V.Kochentkov. Undergraduate student of the ITMO University. Bachelor degree 2015 (ITMO Univ.). Area of interests: Artificial Intelligence, Cognitive Systems, Information Processing, Holography.