

# МАТЕМАТИЧЕСКИЕ ЗАМЕТКИ

1984–2009

П. Б. Иванов

<http://unism.pjwb.org>

<http://unism.pjwb.net>

<http://unism.narod.ru>

## Множества vs. алгебра логики

Часто рассматривают классическую пропозициональную логику и теорию множеств как две реализации булевой алгебры, отождествляя объединение множеств с логическим **или**, а пересечение множеств с логическим **и**. Но это не вполне соответствует логической структуре теории множеств. Например, всякое множество можно считать объединением одноэлементных множеств, но такая трактовка допускает разные интерпретации:

- *перечисление* (сильное объединение): мы рассматриваем множество как элемент **а и** элемент **б и** элемент **с и** ... — это симультанная интерпретация множества как *актуально* целого;
- *исчерпание* (слабое объединение): множество представляется как элемент **а или** элемент **б или** элемент **с или** ... — такая интерпретация подчеркивает идею *потенциальной* целостности, предполагает, так сказать, «зондирование» множества случайным выбором того или иного элемента.

С другой стороны, перечисление алгоритмично, поскольку предполагается, что множество можно *построить* из элементов; напротив, техника исчерпания ссылается на качественную определенность множества, на некие свойства элементов, которые делают их частью целого. Ср.:

- *конвенциональная категоризация*: будем собирательно говорить про г-на Иванова и и г-жу Иванову как о семье Ивановых;
- *иллюстрация*: овощи — это... гм... морковь, огурец, лук, и всякое такое.

Аналогично, два метода определения:

- *конструктивное* (явное): числовые типы данных включают целое, вещественное, дату и время;
- *функциональное* (неявное): назовем все вещественные числа  $x < 0$  отрицательными.

8 января 2000

## Отрицательные множества

Фундаментальные математические объекты — это всегда абстракция широко распространенных способов деятельности. Целые числа восходят к процедуре многократного пересчета, когда результат не зависит от порядка перебора. Напротив, отрицательные числа выражают идею невозможности сосчитать, недостаточности, долга (а правило знаков при умножении напрямую связано с обычаем погашения и взаимозачета долгов). Потом, когда общая идея сформировалась, можно строить формальные модели — возникает теория чисел, с ее нетривиальными теоремами.

Идея множества — это абстракция включения объекта в деятельность. Когда мы приступаем к работе, мы первым делом смотрим, что для этого может пригодиться, и что может помешать. Все, что попадаете нам на глаза, мы оцениваем с этой точки зрения. Речь идет именно о качественной оценке: если годится — будем иметь в виду; если полезность не очевидна — забываем и больше не рассматриваем. Формально говорят о принадлежности элемента множеству, подразумевая возможность конструктивно это проверить. Множество — это то, что ему принадлежит, а все остальное к нему просто не относится. Другими словами, хорошо определено лишь свойство принадлежности, тогда как непринадлежность есть нечто совершенно неопределенное (просто потому, что мы не можем знать все, что имеется в этом мире, и что еще в нем появится). Однако если деятельность представляет собой часть (этап) другой деятельности, охватывающей более широкий круг объектов, свойство непринадлежности можно понимать в узком смысле, как принадлежность дополнению множества. Это уже вполне проверяемый факт.

Вообще говоря, возможны сколь угодно сложные иерархические структуры — и это отражает разнообразие человеческой деятельности. Непринадлежность по-разному определена в разных контекстах. В каждом конкретном случае она соотносится с принадлежностью чему-то другому. Математики предпочитают ограничиваться лишь чем-то конструктивно определимым («универсумом») и не обсуждать то, что лежит вне этой конструкции. Это вполне отвечает тому естественному обстоятельству, что мы всегда работаем с тем, что есть, что наличествует в нашем окружении и (хотя бы в принципе) доступно.

Однако по жизни, помимо объектов, пригодных для текущей деятельности, мы часто сталкиваемся с объектами, которые с ней несовместимы, либо практически недоступны. С тем, что есть — но не может быть использовано («запрещено»). Такие объекты необходимо связаны с деятельностью, они тоже принадлежат ей — но «отрицательным» образом. Это не элементы, а дырки, указания на то, что следует исключить из рассмотрения.

Термин заимствован из физики, где электроны и дырки успешно сосуществуют в теории атома, в физике полупроводников и во многих других практически важных областях. Некоторым образом, позитрон есть просто дырка в вакууме, возникшая при вылете электрона (и потому электроны с позитронами рождаются только парами). Оказывается, что представление о дырках ведет к (практически) интересной математике множеств.

Поскольку мы уже знаем, что дополнение играет в теории множеств роль вычитания, можно формально определить отрицательное множество как дополнение обычного множества до пустого множества, как вычитание из «нуля». Существование такого дополнения просто постулируется. Мы записываем это как  $(-)A = \emptyset \setminus A$ . Знак минус взят в скобки, чтобы подчеркнуть его операторную сущность. По определению, каждый элемент отрицательного множества есть дырка на месте соответствующего (в объектном смысле) элемента исходного множества. Очевидно,  $A \cup (-)A = \emptyset$ . То есть, будучи одновременно включенными в деятельность, элемент и дырка «аннигилируют», и у нас нет ни принадлежности, ни запрета. Симметричным образом, можно положить  $A = \emptyset \setminus (-)A = (-)(-)A$ . Так мы фиксируем логику теории; вообще говоря, для каких-то деятельностей это может быть не так.

Однако, в отличие от чисел, множества не упорядочены линейно, и потому, наряду с положительными и отрицательными множествами, могут существовать произвольные наборы элементов и дырок, которые мы продолжаем собирательно именовать множествами (или классами, если кому-то больше это нравится). Понятно, что в общем случае множество распадается на «положительную» и «отрицательную» компоненты:

$$A = A_p \cup (-)A_n,$$

где  $A_p$  и  $A_n$  — обычные множества (без дырок). В частности, они могут быть пустыми. При объединении множеств соответствующие элементы и дырки аннигилируют:

$$A \cup B = (A \cup B)_p \cup (-)(A \cup B)_n,$$

где

$$(A \cup B)_p = A_p \setminus (A_p \cap B_n) \cup B_p \setminus (B_p \cap A_n)$$

$$(A \cup B)_n = A_n \setminus (A_n \cap B_p) \cup B_n \setminus (B_n \cap A_p)$$

Нетривиальная математика возникает там, где мы переходим от сложения к умножению — от объединения к пересечению множеств. Что вычислять такие выражения формально, мы вводим обычное правило знаков:  $(+)(+) \rightarrow (+)$ ,  $(-)(-) \rightarrow (+)$ ,  $(+)(-) \rightarrow (-)$ ,  $(-)(+) \rightarrow (-)$ . Двойное отрицание уже упоминалось; по смыслу: отсутствие отсутствия есть присутствие. Точно так же осмысленны правила для смешанных знаков: что присутствие отсутствия, что отсутствие присутствия мы понимаем как отсутствие. Опять же, в деятельности это не всегда так — но во многих практически важных ситуациях оправдывается. Тогда, очевидно,

$$A \cap B = (A_p \cup (-)A_n) \cap (B_p \cup (-)B_n) = ((A_p \cap B_p) \cup (A_n \cap B_n)) \cup (-)((A_n \cap B_p) \cup (A_p \cap B_n))$$

Частный случай — «антимножество» произвольного множества:

$$(-)A = \emptyset \setminus A = (-)(A_p \cup (-)A_n) = A_n \cup (-)A_p$$

С количественной стороны обычные множества характеризуются количеством элементов, приписывая каждому элементу вклад  $+1$ . Отрицательные множества, очевидно, характеризуются количеством дырок — но каждая дырка дает вклад  $-1$ , так что мощность отрицательного множества отрицательна. Для произвольного множества надо суммировать вклады элементов и дырок, и возможны разные комбинации.

В приложениях дырки могут быть одной из возможных формализации идеи потребности. И это практически важно для таких наук как психология или экономика. С другой стороны, памятуя о том, что в физике аннигиляция сопровождается возникновением других частиц, можно развить весьма нетривиальные теории деятельности. Заметим, что в обычной теории множеств понятие элемента остается пока неопределенным, и чем это отличается от множества — большой вопрос. Традиционно считают, что ничем. И тем самым ограничивают себя, сводят все к очень узкому диапазону возможных теорий. Но можно видеть, что любой объект как элемент множества (класса) есть, по сути дела, класс всех множеств, которым он принадлежит; напротив, в качестве дырки объект есть класс множеств, которым он не принадлежит.

Любой математик сможет указать алгебраические структуры, которым это изоморфно, и делать далеко идущие выводы. Нас интересует не формальная сторона, а интуитивно ясное представление о качественно определенном математическом объекте. Один математический объект не сводится к другому, даже если они в чем-то похожи. Математическую логику можно выводить из теории множеств или наоборот — от этого они не потеряют своеобразие. Вещественные числа можно моделировать сходящимися последовательностями — но они не перестают быть особым объектом. Даже если мы обзовем их элементами поля — это лишь моделирование реального объекта алгебраической структурой, одна сторона дела. Потому что вещественные числа — не из математики, а из практики. Так же как и натуральные числа, и комплексные числа, и геометрические формы, и множества.

И вот, теперь у нас есть отрицательные множества. Идя дальше по этому пути, можно построить и комплексные множества, и пространства любой размерности... Но об этом как-нибудь в другой раз.

декабрь 1984

### Нечеткие множества и релятивистское сложение скоростей

Аксиоматический каркас для пересечения  $i(a, b)$  и объединения  $u(a, b)$  нечетких множеств задают следующим образом:

1. граничные условия:

$$\begin{aligned} i(1, 1) &= 1 \\ i(0, 1) &= i(1, 0) = i(0, 0) = 0 \\ u(1, 1) &= u(0, 1) = u(1, 0) = 1 \\ u(0, 0) &= 0 \end{aligned}$$

2. коммутативность:

$$\begin{aligned} i(a, b) &= i(b, a) \\ u(a, b) &= u(b, a) \end{aligned}$$

3. монотонность:

$$\text{if } a \leq a' \text{ and } b \leq b' \text{ then } i(a, b) \leq i(a', b') \text{ and } u(a, b) \leq u(a', b')$$

4. ассоциативность:

$$\begin{aligned} i(i(a, b), c) &= i(a, i(b, c)) \\ u(u(a, b), c) &= u(a, u(b, c)) \end{aligned}$$

Вообще говоря, объединения и пересечения нечетких множеств не идемпотентны. Любой выбор общего вида объединения и пересечения будет так или иначе нарушать какие-то свойства булевых решеток. В частности, могут не выполняться закон исключенного третьего и закон противоречия.

Как известно, условия 1–4 приводят к неравенствам

$$\begin{aligned} i(a, b) &\leq \min(a, b) \\ \max(a, b) &\leq u(a, b), \end{aligned}$$

так что  $i(a, b) < u(a, b)$  для  $a \neq b$ . Дополнение нечеткого множества чаще всего определяют как  $c(a) = 1 - a$ .

Можно предложить пример, показывающий, что аксиомы 1–4 могут быть несовместимы с правилами Де Моргана:

$$\begin{aligned} u(a, b) &= c(i(c(a), c(b))) \\ i(a, b) &= c(u(c(a), c(b))) \end{aligned}$$

Действительно, сопоставим каждому элементу  $x_i$  нечеткого множества  $A = \sum_i \mu_i x_i$  вещественное число  $\chi_i \in (-\infty, +\infty)$ , такое, что

$$\mu_i = \frac{1}{2}(1 + \text{th}(\chi_i)).$$

Когда  $\chi$  пробегает значения от  $-\infty$  до  $+\infty$ ,  $\mu$  монотонно возрастает от 0 до 1. Сконструируем дополнение  $A$  согласно общему правилу:

$$c(\mu(\chi)) = 1 - \mu(\chi) = \frac{1}{2}(1 - \text{th}(\chi)) = \frac{1}{2}(1 + \text{th}(-\chi)) = \mu(-\chi).$$

Наконец, определим пересечение двух нечетких множеств как

$$i(\mu_1(\chi_1), \mu_2(\chi_2)) = \mu(\chi_1 + \chi_2) = \frac{1}{2}(1 + \text{th}(\chi_1 + \chi_2)) = \mu_1\mu_2 / (1 - \mu_1 - \mu_2 + 2\mu_1\mu_2).$$

Это выражение удовлетворяет граничным условиям для пересечения, оно очевидным образом коммутативно и монотонно. Прямым вычислением можно показать ассоциативность:

$$i(i(\mu_1(\chi_1), \mu_2(\chi_2)), \mu_3(\chi_3)) = i(\mu_1(\chi_1), i(\mu_2(\chi_2), \mu_3(\chi_3))) = \frac{\mu_1\mu_2\mu_3}{(1 - \mu_1 - \mu_2 - \mu_3 + \mu_1\mu_2 + \mu_1\mu_3 + \mu_2\mu_3)}.$$

Так определенное  $i(\mu_1, \mu_2)$  не идемпотентно, хотя можно видеть, что оно *асимптотически* идемпотентно при  $\chi \rightarrow \pm\infty$ , что обеспечивает правильный переход к обычным (не нечетким) множествам.

Воспользуемся теперь правилом Де Моргана, чтобы определить объединение множеств через пересечение и дополнение:

$$\begin{aligned} u(\mu_1(\chi_1), \mu_2(\chi_2)) &= c(i(c(\mu_1(\chi_1)), c(\mu_2(\chi_2)))) = c(i(\mu_1(-\chi_1), \mu_2(-\chi_2))) = \\ &= c(\mu(-\chi_1 - \chi_2)) = \mu(\chi_1 + \chi_2) = i(\mu_1(\chi_1), \mu_2(\chi_2)) \end{aligned}$$

Поскольку определенные таким образом объединение и пересечение совпадают, объединение не вписывается в аксиоматический каркас, ибо нарушаются граничные условия при  $\mu_1 = 0$  и  $\mu_2 = 1$ . И, конечно же, величина  $i(\mu_1, \mu_2)$  никак не может быть меньше  $u(\mu_1, \mu_2)$ , как это следовало бы из базовых аксиом. Заметим, что при таком выборе пересечения и объединения нарушается закон исключенного третьего и закон противоречия, ибо

$$i(\mu(\chi), \mu(-\chi)) = \frac{1}{2}$$

Физический смысл этих определений — релятивистское сложение скоростей. Частицам, движущимся вперед со скоростью света, приписано значение  $\mu = 1$ , а тем, что движутся со скоростью света в обратном направлении, — значение  $\mu = 0$ . Если ограничиться при сложении скоростей только положительными значениями, объединение относится к так называемому классу Гамакера с  $\gamma = 2$ :

$$\begin{aligned}\mu(\chi) &= \text{th}(\chi_1 + \chi_2) \\ u(\mu_1(\chi_1), \mu_2(\chi_2)) &= \mu(\chi_1 + \chi_2) = \text{th}(\chi_1 + \chi_2) = (\mu_1 + \mu_2)/(1 + \mu_1\mu_2)\end{aligned}$$

Вместе со стандартным дополнением оно порождает пересечение

$$i(\mu_1, \mu_2) = \mu_1\mu_2/(1 + (1 - \mu_1)(1 - \mu_2)),$$

удовлетворяющее и аксиоматическому каркасу, правилам Де Моргана. Однако в этом случае определение дополнения оказывается «физически несовместимым» с определением объединения, поскольку галилеево правило сложения скоростей используется вместе с релятивистской формулой.

14 мая 1997

### Логические симметрии и комплексная логика

Традиционно алгебраические представления в математической логике приписывают логическому значению **истина** число 1, а логическому значению **ложь** — число 0. При этом конъюнкция ( $\wedge$ ) и дизъюнкция ( $\vee$ ) естественно соотносятся с алгебраическим умножением и сложением соответственно. Конечно, такая алгебраическая структура отличается от обычной арифметики, но сама возможность оперировать с числами вместо особых логических значений позволяет прояснить некоторые аспекты формальной логики и подсказывает варианты ее обобщения.

Но есть и другой подход, который также может быть в чем-то полезен при построении обобщенных логик. Вместо того, чтобы представлять логическое значение ложь нулем, мы можем сопоставить ему число  $-1$ . Обозначая логическое отрицание знаком минус, мы логично получаем, что  $-1$  (**ложь**) — это **не истина**, а  $-(-1)$  равняется 1 по определению:

$$-(-a) = a,$$

где буквами  $a, b, c, \dots$  мы обозначаем все, что может принимать логические значения: константы, переменные, логические выражения (формулы), содержащие константы, переменные и другие формулы в разных комбинациях; кроме того, можно рассматривать также логические функции как сокращения для каких-то выражений, а потом перейти к переменным функциям и функциональным выражениям... В нашем контексте такая иерархичность не существенна.

Введем теперь алгебраическое умножение  $*$  так, чтобы  $1 * 1 = 1$  и  $(-1) * (-1) = 1$ , с естественным выбором:  $(-1) * 1 = 1 * (-1) = -1$ . Эта операция, очевидно, соответствует логической эквивалентности; она коммутативна и ассоциативна:

$$\begin{aligned}a * b &= b * a \\ (a * b) * c &= a * (b * c)\end{aligned}$$

но, в отличие от конъюнкции и дизъюнкции, она не идемпотентна:

$$a * a \neq a$$

что, впрочем, не вызывает особого беспокойства, поскольку обычное числовое умножение тоже идемпотентностью не страдает. Отметим также, что

$$-(a * b) = (-a) * b = a * (-b),$$

в приятном соответствии с нашими арифметическими ожиданиями.

Введенное таким образом умножение симметрично в пространстве значений истинности: замена  $1 \leftrightarrow (-1)$  приводит к той же самой таблице истинности для логической эквивалентности.

Подчеркнем, что знак равенства в этих и последующих формулах, обладая всеми свойствами эквивалентности, принадлежит логике нашего рассуждения (уровню методологии), а не логике его предмета. Алгебраические структуры этих уровней могут различаться, вплоть до способов введения логических значений. Равенство и неравенство могут использоваться только для сравнения формул предметной логики — но не внутри этих формул.

С учетом наших определений, мы можем формально исключить отрицание из теории, заменив его умножением:

$$-a = (-1) * a.$$

Это открывает многообещающее направление исследований, поскольку мы устраним весьма скользкие концептуальные моменты, вызывавшие бурные дискуссии на протяжении многих веков. С другой стороны, необходимость введения логического отрицания не очевидна в многозначных и нечетких логиках; все, на что мы можем здесь сослаться, — это старинная традиция. Тем не менее, этот подход не решает проблему до конца: одно из логических значений почему-то оказывается предпочтительнее другого — а это, по сути дела, протаскивает в логику отрицание неявным образом. В качестве альтернативы можно было бы независимым образом определить неэквивалентность:

$$a \times b = -(a * b),$$

что как бы вводит «дополнительное отрицание»:

$$\sim a = (1) * a,$$

так что в итоге мы имеем полностью уравновешенную теорию, не требующую одноместных операций. Возможные приложения такой логики подчеркивают, что в реальной жизни установление эквивалентности и поиск различий — это очень разные (противоположные) деятельности, а между ними — море промежуточных вариантов.

Очевидно, значения истинности в нашей симметризованной логике могут быть выражены через две фундаментальные операции:

$$\begin{aligned} 1 &= a * a \\ -1 &= a \times a \end{aligned}$$

для любого  $a$ . Однако, поскольку эти определения неявно содержат квантор, они принадлежат более высокому уровню иерархии, и мы не имеем права вставлять так определенные логические значения в алгебру логики, если не принять специальных мер для уравнивания теории, как в примере с двумя отрицаниями. Иначе говоря, пространство истинности глобально по отношению к предметной области: в принципе, могут быть и другие значения истинности, но каждое из них представляет предметный уровень целиком. Однако если речь идет о некоторых частных утверждениях относительно предметной области, допустимо использовать такие определения логических констант — сознавая, что полученные формулы относятся только к обычной двузначной логике.

Полностью симметричная логика может быть дополнена несимметричными операциями, которые зависят от структуры пространства истинности. Меняя местами истину и ложь, мы получим иные таблицы истинности для таких операций. Самые известные из таких операций —

конъюнкция и дизъюнкция, которые связаны друг с другом через отрицание:

$$a \wedge b = -((-a) \vee (-b))$$

$$a \vee b = -((-a) \wedge (-b))$$

По отношению к определенным двум видам логического умножения, конъюнкция и дизъюнкция могут быть охарактеризованы как «аддитивные». В логике без отрицания они становятся независимыми, подобно независимо введенным эквивалентности и неэквивалентности. Но в двузначной логике (скрыто предполагающей отрицание) мы можем симметричным образом задать их взаимозависимость выражениями вроде

$$a \wedge b = ((a * a) \times (b * b)) * (((a \times a) * a) \vee ((b \times b) * b))$$

$$a \vee b = ((a * a) \times (b * b)) * (((a \times a) * a) \wedge ((b \times b) * b))$$

Аддитивные (несимметричные) операции могут показаться более фундаментальными по сравнению с эквивалентностью и неэквивалентностью, поскольку, вроде бы, несимметричные операции всегда можно скомпоновать симметричным образом, а наоборот — не получится:

$$a * b = (a \wedge b) \vee ((-a) \wedge (-b))$$

$$a \times b = (a \vee b) \wedge ((-a) \vee (-b))$$

Формализм двузначной логики может, поэтому, быть построен, исходя из отрицания и одной-единственной несимметричной («аддитивной») операции. Но эта простота обманчива, она целиком зависит от идеи логического отрицания — что, по сути дела, компенсирует одну несимметричность другой. В каких-то случаях ограничиваться такой, урезанной логикой было бы неправильно.

Исследуя возможность обобщений, полезно обратиться к *Arithmetices Principia* Пеано. Не придираясь пока к искусственности его теоретико-множественных формулировок, подчеркнем крайне важный момент: натуральные числа (по Пеано) начинаются с выбора *единицы*, а все остальные натуральные числа представляют собой определенное количество единиц. Современные математики склонны добавлять к натуральным числам еще и псевдо-число 0; это нарушает прозрачность и последовательность теории. Исходно, Пеано включает в свой список аксиом свойства числового равенства; последующие авторы исключают эти аксиомы на том основании, что они, дескать, следуют из логики — тем самым перепутываются совершенно разные вещи, разные уровни иерархии, а это недопустимо, при всем кажущемся «изоморфизме». Арифметика натуральных чисел выводится из единицы и единственной одноместной операции, инкремента (который было бы преждевременно трактовать как сложение с единицей). Однако такая индуктивная конструкция не позволяет сразу увидеть внутренние симметрии теории — поэтому вводят двуместные операции сложения и умножения, с формальным продолжением в виде вычитания и деления — которые уже не очень вяжутся с натуральностью. Вычитание неизбежно приводит к идее отрицательных чисел — и вводит ноль как число. Деление ведет к арифметике рациональных чисел, предпосылке вещественных чисел. Такое же развитие возможно и в алгебраически переформулированной логике. Однако здесь мы обратимся к другому возможному расширению — комплекснозначной логике.

Вспомним, что пространство логических значений устроено так, что  $(-1) * (-1) = 1$ . Можно попробовать похожим образом ввести логическое значение, представляющее собой «квадратный корень» из логического значения **ложь**:  $i * i = -1$ . Да, мы не знаем, что это такое, и как это изготовить из подручных средств. Но мы можем просто постулировать его существование и обозначить буквой  $i$ . Допустимость подобных концептуализаций давно уже стала краеугольным камнем современной математики. И все же в нашей алгебраической логике эта *мнимая* единица допускает вполне очевидную интерпретацию. Действительно, равенство  $a * a = 1$  означает, что всякая логическая величина эквивалентна сама себе, — то есть, что  $a \equiv a$  есть **истина** (напомним, что здесь имеется в виду эквивалентность в предметной области, а не логика нашего описания).

Тогда равенство  $i * i = -1$  вводит в рассмотрение величину, которая *не тождественна себе самой*. Какой бы странной не казалась эта идея рационально мыслящему человеку, она отнюдь не нова, ее давно уже обсуждает философия (прежде всего гегелевская и марксизм). Таким образом, наша алгебраическая формулировка логики с комплексными значениями истинности может стать подходящим исходным пунктом для примирения научной и философской методологии.

В этой расширенной логике значения истинности могут иметь как вещественную, так и мнимую части. В зависимости от выбора «аддитивных» операций, мы получим разные правила манипулирования комплексными значениями истинности. В общем случае эти правила будут гораздо сложнее обычной комплексной арифметики. Это не отменяет соответствие как таковое. В частности, традиционная интерпретация фазы как угла поворота вектора в комплексной плоскости сразу же сопоставимо с интерпретацией логического отрицания как обращения вектора. Тут же появляется возможность различить повороты и зеркальные отражения — с тонкими различиями в логике.

Присутствие «фазовых» компонент в логических величинах открывает новые горизонты для квантовой логики, поскольку «наблюдаемая» рациональность связана с вещественными числами, а внутренние, зависящие от фазы компоненты мышления остаются в тени. Учет таких скрытых логических состояний вполне аналогичен переходу в механике от положения и импульса классической частицы к виртуальному движению во внутреннем конфигурационном пространстве, которое несопоставимо напрямую с наблюдаемыми величинами.

Подведем итоги. Симметризованная алгебраическая формулировка математической логики достаточно привлекательна, чтобы заняться исследованием ее возможных обобщений, включая многозначные, рациональные, нечеткие и комплексные логики. Особенно обнадеживает, что, с одной стороны, эта математика не требует революционных изменений и новых парадигм, а с другой, новые понятия остаются интуитивно ясными и доступными даже тем, кто не одобряет нагромождения формалистических хитросплетений. Иерархии логики еще хватает направлений нетривиального развития.

ноябрь 2003

### Точки и пределы

Традиционно, изготовление метрического пространства  $S$  выглядит примерно так: имеется некоторое множество  $B$  (которое в дальнейшем будем для краткости именовать *базой*; поскольку о топологии мы не говорим, путаницы не возникает), и мы умеем для любых двух точек  $x$  и  $y$  этого множества найти некоторое (вещественное) число  $\rho$ , которое называется расстоянием между точками при условии, что:

- (1)  $\rho(x, y) = 0$  тогда, и только тогда, когда  $x = y$
- (2)  $\rho(x, y) = \rho(y, x)$
- (3)  $\rho(x, y) \leq \rho(x, z) + \rho(x, z)$

Здесь тоже предполагается квантор «для любых» (над смыслом и осуществимостью которого мы пока не задумываемся). Вторая аксиома, по сути дела, говорит, что метрика есть функция не упорядоченной пары, а *подмножества*: порядок элементов не играет роли. Это сразу наводит на мысль, что метрика есть лишь одна из разновидностей меры: под расстоянием между точками скрывается размер того, что лежит между ними (длина или продолжительность пути). Но это, опять-таки, тема для особого рассмотрения. Последнее свойство называется правилом треугольника; именно его чаще всего модифицируют в альтернативных теориях метрики (например, для ультраметрических пространств).

Формальная математика исходит из того, что определять можно что угодно и как угодно. Осмысленности никто не требует. Разумеется, на самом деле определения вводятся не от фонаря,



а так, чтобы в конечном итоге получить то, что хочется. А что хочется — определяет практика. Поэтому полной бессмыслицы никогда нет. Но обычному человеку бывает трудно пробиться через математические дебри чаще всего потому, что смысла происходящего он не усматривает, поскольку математики обычно не склонны явным образом изложить мотивы, а даже наоборот, всячески отрешаются от практического интереса. Здесь мы попробуем обратить внимание на то, что обычно остается в тени.

Следующий этап метрического строительства — последовательности точек множества  $B$ . Формально, это отображение множества натуральных чисел в базу метрического пространства. Однако по смыслу речь идет о некотором алгоритме, позволяющем выбирать по какому-нибудь принципу одну точку базы за другой; собственно, этот принцип (совместно с указанием начального элемента) и называется последовательностью. Если нет такой направленности, упорядочения, — это уже не последовательность, а обыкновенное множество. Можно легко заметить, что само определение натуральных чисел выглядит именно так: достигли некоторого номера — можно назначить следующий. Операции над числами (арифметика) привнесены в эту последовательность извне: это лишь один из возможных способов отождествления разных подпоследовательностей. Но об этом тоже в другой раз.

А пока у нас есть последовательности точек базы, которые традиционно обозначаются индексированными буквами типа  $x_n$ , при  $n = 1, 2, \dots$  (иногда считать начинают с нуля, или с какого-то ненулевого индекса). Последовательности бывают очень разные, причем элементы с разными индексами запросто могут совпадать. В частности, вся последовательность может воспроизводить одну-единственную точку. В более общем случае — последовательность с «самопересечениями», с многочисленными петлями. Частный случай этого общего случая — периодические последовательности (замкнутые траектории). Наконец, последовательности могут быть просто случайными — и это еще один повод не считать последовательности точек базы заранее существующими: вообще говоря, их надо каждый раз честно вычислять, и не факт, что результат окажется тем же самым.

Поскольку мы интересуемся расстояниями, можно с ходу предложить два способа превратить последовательность точек базы (которые могут быть довольно сложными объектами, причем не всегда математическими) в последовательность чисел (про которые мы, вроде бы, все знаем). Первый вариант — вычислять расстояния между точками последовательности:  $\rho_{\text{in}}(n; m) = \rho(x_{n+m}, x_n)$ . Второй вариант — зафиксировать некоторую точку базы и определять расстояния до нее для каждого элемента последовательности:  $\rho_{\text{out}}(n) = \rho(x_n, x_0)$ . Обозначения подчеркивают разную природу этих величин: внутренняя и внешняя структура. Первая хорошо знакома нам из статистики: это своего рода автокорреляция. Полный набор таких функций (при всевозможных  $m$ ) достаточно хорошо характеризует последовательность как самостоятельный объект. Вторая возможность — очевидная реминисценция из векторного анализа: каждая точка базы представляется ее радиус-вектором, а если выбрать несколько опорных точек (требуемое количество зависит, в частности, от способа арифметизации базы), то мы можем задать и направление. Когда такие внешние опорные точки образуют некоторую последовательность, мы приходим к «синтезу» внешней и внутренней структур, к сравнению последовательностей:  $\rho_{y,x}(n; m) = \rho(y_{n+m}, x_n)$ , или наоборот:  $\rho_{x,y}(n; m) = \rho(x_{n+m}, y_n)$ .

Внешняя структура, вообще говоря, не вытекает из внутренней, и наоборот. Все зависит от строения базы и способа задания метрики. Но во многих практически важных случаях эти две структуры дают, по видимости, лишь два разных представления одного и того же.

Одно из важнейших понятий такого рода — сходимость. Если для любого положительного вещественного числа  $\varepsilon$  существует такое  $N$ , что  $\rho_{\text{in}}(N; m) \leq \varepsilon$  для любых  $n \geq N$  при некотором фиксированном  $m$  (обычно выбирают  $m = 1$ ), то  $x_n$  называется последовательностью Коши. С другой стороны, если для любого  $\varepsilon$  существует такое  $N$ , что  $\rho_{\text{out}}(n) \leq \varepsilon$  при любых  $n \geq N$ , говорят,

что последовательность  $x_n$  *сходится* к точке  $x_0$ , или, что точка  $x_0$  является *пределом* последовательности:  $x_n \rightarrow x_0$ . Предельные свойства последовательностей точек сводятся таким образом к соответствующим свойствам числовых рядов, абстрагируются от базы. Однако не всегда наблюдения над числами допустимо напрямую переносить на свойства нечисловых пространств: если при этом не учитывать строение базы, такой перенос может стать источником логических ошибок.

В метрических пространствах, по правилу треугольника, сходящаяся последовательность точек всегда оказывается последовательностью Коши — но не всякая последовательность Коши сходится, ибо может оказаться, что нет ни одной точки базы, расстояния до которой бесконечно убывают (стремятся к нулю). С другой стороны, по тому же правилу треугольника, если  $x_n \rightarrow x_0$  и  $y_n \rightarrow x_0$ , то  $\rho_{y,x} \rightarrow 0$ . Возникает соблазн предположить и обратное: если  $\rho_{y,x} \rightarrow 0$ , то  $x_n$  и  $y_n$  либо вместе не сходятся, либо сходятся к одной и той же точке базы. Если это так — можно смело объявлять все последовательности с  $\rho_{y,x} \rightarrow 0$  эквивалентными, и при необходимости просто пополнить базу недостающими предельными точками. В элементарной математике просто-напросто замечают, что, если расстояние между пределами эквивалентных последовательностей ненулевое, достаточно взять  $\varepsilon$  равным половине этого расстояния — и окажется, что три условия сходимости (для  $x_n$ , для  $y_n$  и для  $\rho_{y,x}$ ) нарушают правило треугольника. Следовательно, расстояние между предельными точками равно нулю — а тогда, по первому свойству метрики, предельные точки совпадают — что, вроде бы, и требовалось доказать.

Но давайте проковыряем маленькую дырочку в нерушимых стенах математической строгости и выглянем туда, наружу, где математика становится не совсем элементарной. По логике, определение предела говорит лишь, что расстояние между предельной точкой и точками сходящейся последовательности *может быть сделано* меньше любого наперед заданного числа. Но это вовсе не тождество, не равенство нулю (если речь не идет о бесконечном повторении одной и той же точки). Точно так же, совместная сходимость к двум разным точкам означает лишь, что расстояние между ними *может быть сделано* меньше любого числа — а вовсе не то, что это расстояние *равно* нулю. Другими словами, расстояние между предельными точками эквивалентных последовательностей — это предел числовой последовательности, а вовсе не готовенькое число. Оно *стремится* к нулю — но не *равно* ему. На заре математического анализа, его отцы-основатели говорили о *бесконечно малых* величинах — не отождествляя их с обычными числами. Потом статическая парадигма вытеснила понятие бесконечно малой из школьных курсов; чуть позже к нему вернулись в контексте нестандартного анализа (который, впрочем, лишь пытается иначе сформулировать все те же статические идеи). Так вот, расстояние между предельными точками эквивалентных последовательностей бесконечно мало — но не нуль. Применить к таким величинам первое свойство расстояния мы логически не вправе. Лишь в некоторых особых случаях школьное доказательство согласно с логикой — например, для изолированных точек базы.

Точно так же, и правило треугольника в исходном виде относится лишь к конечным (фиксированным) величинам, а вовсе не к процедуре перехода к пределу. В применении к бесконечно малым расстояниям, это свойство расстояния следовало бы формулировать иначе:  $\rho(x, y)$  есть бесконечно малая того же или более высокого порядка по сравнению с  $\rho(x, z) + \rho(x, z)$ .

Копнем глубже. Когда мы сравниваем две последовательности, мы выходим за рамки исходного метрического пространства в *другое* метрическое пространство — с базой в виде множества последовательностей Коши на пространстве  $S$ . Именно по отношению к этой базе определяется эквивалентность последовательностей. Тот нуль, к которому сходятся расстояния между элементами разных последовательностей, представляет иное понятие метрики, отличное от метрики в  $S$ , и потому  $x_n$ ,  $y_n$  и  $\rho_{y,x}$  просто нельзя сравнивать в рамках одного и того же правила треугольника! Это чисто логическая ошибка, подмена понятия. Не учли мы, что одним и тем же

числом (именем, ярлыком) могут обозначаться разные сущности. Расстояния в пространстве последовательностей могут вычисляться на основе расстояний на множестве их элементов — но это разные понятия, даже если они иногда совпадают в числовом выражении.

Чтобы отождествить класс эквивалентных последовательностей Коши в пространстве  $S$  с точкой базы  $B$ , требуется особая операция. Не факт, что такая операция всегда осуществима, и что она обязательно будет однозначной. Например, отождествление происходит случайным образом с разными точками в некоторой области — и можно рассматривать соответствующее распределение вероятностей. Только в особом случае, когда распределение задано  $\delta$ -функцией (которая, как известно, вовсе не функция, а функционал), возникает некое подобие точки. Точно так же, если последовательность Коши не сходится к точке базы, пополнение базы возможно далеко не всегда — поскольку такие дополнительные точки могут логически не вписываться в определение предмета теории, и для них надо строить другую теорию, с другой предметной областью.

Наконец, для продвинутых дилетантов. Совокупность последовательностей, сходящихся к некоторой точке базы  $x$  можно считать ее инфинитезимальной окрестностью. Каждая точка оказывается, таким образом, центром облака бесконечно малых отклонений — виртуальных вариаций. Для любого положительного вещественного числа  $r$ , лишь конечное число членов любой сходящейся к  $x$  последовательности оказывается за пределами сферы радиуса  $r$ . Поскольку в предметной теории сходящиеся последовательности образуются по определенному закону, можно оценить среднее количество точек  $-\varepsilon(r)$  за пределами  $r$ -сферы; знак минус взят из тех соображений, что последовательность всегда «укорачивается» на уровне  $r$ . Знатоки нестандартного анализа тут же усмотрят здесь нечто вроде ультрафильтра. Величину  $\varepsilon(r)$  можно считать мерой связи точки  $x$  с базой — как в физике определяют энергии связи электронов в атоме. Эта «энергия» отрицательна, поскольку за нулевой уровень принята точка полного отрыва точки от базы (аналог потенциала ионизации). В зависимости от строения базы (предмета теории) могут возникать разные распределения по энергиям связи; в атомах типичная картина — последовательность дискретных уровней, сходящаяся к порогу ионизации.

Одно из очевидных приложений — обобщение понятия принадлежности элемента множеству. В обычной теории элемент либо принадлежит множеству — либо нет. Теория нечетких множеств допускает введение функций принадлежности — но совершенно неясно, каковы они и откуда их брать. Здесь мы видим, что принадлежность (способ связи элемента с множеством) задана строением предметной области — допустимыми траекториями в ней. Функция принадлежности оказывается тогда свойством инфинитезимальной окрестности точки, и вычисляется как среднее некоторого оператора на ее внутреннем пространстве.

Вот мы и подошли к сути вопроса. Рассмотрение любых конструкций на множестве-базе выводит нас за его пределы, в пространство более высокого уровня. Но это означает также и возникновение внутренней структуры у каждой точки базы — ее внутреннего пространства. Так математические объекты становятся иерархиями.

Понятно, что вовсе не обязательно останавливаться на одних лишь классах сходящихся последовательностей. Можно рассматривать любые траектории приближения к предельной точке, в том числе непрерывные (например, по спирали). Можно говорить о случайных выборках приближений. Можно задавать точки неявно, как пересечения семейств множеств. Можно даже считать точки алгоритмами или физическими процессами, с их особыми симметриями («спинорными» степенями свободы). Внутреннее пространство точки может оказаться сколь угодно сложным. Пространство-база остается метрическим. И если расстояние между его точками равно нулю — точки совпадают. Но из инфинитезимальности расстояний не следует ровным счетом ничего — пока мы не указали уровень рассмотрения, не развернули иерархию каким-то вполне определенным образом. Переход от внутренней динамики в каждой точке к свойствам на уровне базы требует задания способа проецирования внутреннего пространства в базу (подобно тому, как в квантовой механике вводятся операторы наблюдаемых).

С внутренними пространствами мы сталкиваемся каждый раз при совпадении одного с другим (эквивалентность, равенство и т. д.). Количественно, это означает, что некоторая мера отличия равняется нулю. Пока речь идет о грубом сравнении, сложность сопоставляемых объектов не бросается в глаза. Но как только мы абстрагируемся от различий на некотором уровне — нам приходится иметь дело с более тонкими вариациями, залезать «внутрь» нуля. Здесь даже не нужно квантовой механики: достаточно вспомнить, что Солнечная система, во всем ее богатстве, представляется одной точкой для обитателей соседних звезд — не говоря уже о далеких галактиках.

В качестве собственно математической иллюстрации — равенство комплексных чисел. Формально, два числа равны, если расстояние между ними (модуль их разности) равняется нулю. Но если задавать комплексное число модулем и фазой, точки  $\{0, \varphi_1\}$  и  $\{0, \varphi_2\}$  — вовсе не одно и то же. Выходит, что равенство нулю расстояния — это еще не все, и надо позаботиться о том, чтобы приближались мы к нулю по одинаковым траекториям, что и обеспечит равенство фаз. Традиционно, с такими тонкостями никто не связывается, и фазу для числа нуль вообще не определяют. Следовательно, и расстояние на комплексной плоскости определено с точностью до фазы, то есть, как некоторое среднее — и сразу возникает вопрос о правомерности и способе усреднения.

Обычно первичным считают линейное алгебраическое представление комплексного числа с покомпонентным равенством. Тогда нуль — одна-единственная точка. Точно так же вводят единственную бесконечно удаленную точку (проективная геометрия) — но в реальных вычислениях приходится разбираться, по какой траектории мы обходим эти две сингулярности. Считать нуль и бесконечность числами — чистая условность. Потому что это совсем не числа, а настоящие числа так себя не ведут.

Аналогично с векторными пространствами: одно дело — покомпонентное представление, а другое — вектор как величина и направление. Какое направление у нуль-вектора?

Но почему следует исходить именно из линейной схемы? Быть может, в мире главное вращение, а вовсе не поступательное движение? Покомпонентное представление — частный случай, как прямая — частный случай произвольной траектории, вдоль которой внутреннее пространство каждой точки целиком отображается во внутреннее пространство следующей. И только в «нулевом» приближении справедлива теория обычных метрических пространств.

*август 1988*

### **Иерархическая размерность**

С тех пор, как математика (искусственно) абстрагировалась от повседневного опыта и начала заниматься исключительно формальными вопросами, мы перестали понимать, что это такое — пространство. Пространством стали называть все, что угодно. Тогда математики просто отказались от этого понятия и оставили в обиходе лишь достаточно формализованные конструкции, названия которых лишь по традиции содержат слово «пространство» — больше в смысле «многообразие». Для таких абстрактных объектов вводятся столь же абстрактные определения размерности, которые нормальному человеку равным счетом ничего не говорят. Нам предлагают поверить на слово большим ученым и тупо применять результаты их трудов, жать на кнопки, неизвестно зачем, а тем более почему. Поскольку физика теперь больше любит формулы, чем присматривается к природе, ждать объяснений от физиков — дело столь же напрасное. И пока философия не вырастет из детских штанишек и не перестанет косить под науку — от нее вразумительности ни на грош.

Но тем, у кого в душе еще осталась хоть капелька человеческой любознательности, время от времени хочется минимальной наглядности, интуиции не по поводу комбинирования иероглифов, а в действительных отношениях с окружающими нас вещами, а заодно и с другими

людьми. Наше практическое представление о пространстве — одно из таких отношений. Да, жизнь сложна, и приходится выстраивать деятельность по-разному: это приводит к мысли о возможности различно организованных пространств. В каких-то случаях можно характеризовать такие различия понятием «размерность» — но пусть это будет именно понятие, совокупность типовых приемов построения пространств разной размерности, а не просто жонглирование символами.

Давайте попробуем хотя бы схематически наметить одно из возможных решений. Пусть каждый выбор требует пространного обсуждения оснований — ограничимся хотя бы намеками на то, для чего эти основания следует подыскать.

С точки зрения человеческой деятельности, идея пространства говорит об имеющихся возможностях. Именно в этом смысле «пространственная» лексика употребляется в быту. Научный быт — ничем здесь не выделяется. Примеры из науки привычнее, поскольку философия пока не желает смотреть ни на что другое. Ладно, пусть будет наука. С прицелом на неизбежность исследования практических приемов конструирования пространств в других областях деятельности. Будем постепенно добавлять необходимые для этого инструменты.

Начнем с того, что пространство на самом деле существует. Это не произвол, не пустая фантазия — так устроен мир на самом деле. И так же мы в этом мире действуем. Но существование пространства — не того же рода, как существование материальных вещей. Пространство не существует отдельно от вещей — это именно отношение между вещами. Такое, привязанное к материи существование в философии называют идеальным. Но и вещи не существуют вне отношений между собой — и всякая материальность соотносится с чем-то идеальным (и это не обязательно пространство). В определенных условиях идеальные сущности *представлены* какими-то вещами: например, слово «пространство» само по себе — просто звук, сотрясение воздуха (или краска на бумаге, или свечение экрана); но когда мы рассуждаем о пространстве, эта вещь становится его знаком — в рамках конкретной деятельности, в пределах темы. Пространство объективно — однако в каждом конкретном случае мы подходим к этой идее с какой-то одной стороны. Поэтому и любые характеристики пространства могут относиться либо к пространству самому по себе («внутреннее устройство», независимое от наших интересов), либо к одному из возможных практических представлений (конкретной «реализации»). Речь идет не только о понятиях разных типов (уровней) — но и о расслоении каждого понятия, его многоуровневости. В соответствии с этим, мы различаем собственную («геометрическую») размерность пространства — и внешнюю («топологическую») размерность.

### *0. Точка*

По хорошему, это не то, с чего следовало бы начинать. Скорее, наоборот: итог, высшая степень абстракции. Но поскольку здесь мы лишь представляем нечто, ранее придуманное (и продуманное), — можно идти с конца.

С точки зрения конструктивной теории размерности, точка есть исходный пункт: вакуум, «нульмерное» пространство — то есть, по сути, отсутствие пространственности как таковой. Невозможность движения и действия.

Поскольку мы говорим об объективности пространства, точка есть выражение этой объективности. Пространство состоит из точек — это лишь другое выражение существования и качественной определенности. Заметим, что точка становится качественно определенной лишь по отношению к «своему» пространству — наследует его качество. Не бывает точек самих по себе — есть только разные пространства, которые мы в каких-то случаях можем свернуть в «точку», с сохранением того же качества.

### *1. Измерение*

В философии есть категория, обозначаемая словом «мера» (в отличие от меры в узко-математическом понимании). Имеется в виду сама возможность соотнести одно с другим. Когда

одно становится мерилем другого. Единицей измерения. То есть, с одной стороны, то, что мы измеряем, в каком-то аспекте качественно однородно с тем, что мы принимаем за единицу (соизмеримо с ним); с другой стороны, оно отличается от единицы, и это отличие мы называем количеством.

В отличие от точки, всякое измерение предполагает возможность движения в некоторых пределах («степень свободы»). Вот это мы и называем (одномерным) пространством. При каждом конкретном выборе единицы пространственные отношения будут выражаться разными числами (или даже вообще не числами) — но сами по себе они от этого не меняются, и столь же объективен (внутренний) порядок — в нем суть одномерности.

Если есть несколько разных мер (и соответствующих единиц измерения), мы говорим о многомерном пространстве. О возможной взаимозависимости разных действий — см. ниже. Здесь важно подчеркнуть, что, вообще говоря, разные измерения качественно различны, и просто так друг с другом не складываются. Например (забегая чуть вперед), чтобы построить многомерное метрическое пространство, надо каким-то образом привести разные единицы к общей мере: в форме для интервала

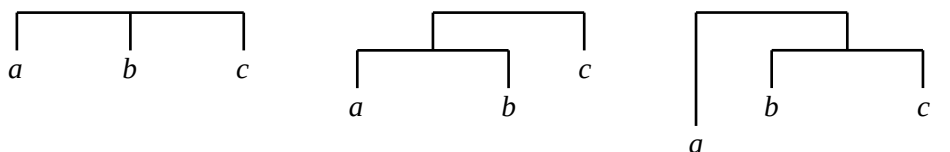
$$ds^2 = g_{ik} dx^i dx^k$$

коэффициенты  $g$  имеют (физическую) размерность

$$[\text{единица интервала}] / [\text{единица } i] / [\text{единица } k]$$

Когда мы выражаем все размеры (допустим) в метрах, это означает, что у нас есть *практическая* процедура преобразования к метрам исходных единиц, которые на самом деле называются как-то вроде «длина» («погонный метр»), «ширина», «высота» — и десятки других вариантов, в зависимости от того, что именно мы измеряем. Соответственно, полученная таким образом единая единица имеет смысл лишь постольку, поскольку есть деятельность, протекающая в пространстве именно такой размерности: например, нет смысла приводить все к долларам там, где доллар не в ходу.

В каждом приложении (в конкретной деятельности) мы представляем пространство целой положительной размерности, перечисляя его измерения в определенном порядке. В одних случаях этот порядок важен, в других нет. Само пространство никак не зависит от того, чем мы его представляем, — оно предполагает все возможные представления, и не сводится ни к одному из них. Тем не менее, набор возможностей не произволен; то общее, что есть у различных представлений одного пространства мы и называем его (геометрической) размерностью. Мы говорим, что размерность — это иерархия, допускающая развертывание разных иерархических структур (специальных представлений). Так, легко видеть, что декартово произведение пространств разной размерности  $N_1$  и  $N_2$ , явным образом некоммутативно, хотя суммарная размерность в любом случае равна  $N_1 + N_2$ . Пространство любой размерности в этой модели выглядит как совокупность всех возможных разбиений полной размерности на суммы различного числа слагаемых. Графически это выглядит как набор древовидных структур (обращений иерархии):



плюс все возможные перестановки в последовательности  $(a, b, c)$ . На практике какие-то варианты развертывания могут оказаться недоступными: например, чтобы попасть в квартиру, мы сначала должны зайти в подъезд (горизонталь), а потом подняться на лифте (вертикаль); обратный порядок предполагает способность лазить по стенам.

## 2. Связь

Традиционное в аналитической механике понятие связи позволяет конструировать пространства отрицательной размерности. Добавление измерения увеличивает размерность пространства (добавляет степень свободы); связь, наоборот, запрещает движение в определенном направлении (не обязательно по прямой) и тем самым эффективно уменьшает размерность пространства. Простейшую связь мы считаем пространством размерности  $-1$ . Комбинации связей дают связи более высокого порядка; так порождаются пространства целой отрицательной размерности.

Способ наложения связи зависит от пространства, в котором она определена и от способа его параметризации. Например, если в пространстве задана система координат, связь может быть представлена некоторым уравнением. Как и в случае (положительных) измерений, связь не зависит от способа параметризации. Однако измерения пространства выражают отношения между вещами — тогда как связь выражает отношения этих отношений; это, так сказать, идеальность более высокого уровня. Тем не менее, во многих практических ситуациях, когда важно не строение идеальности как таковой, а ее отношение к материи, различие между измерениями пространства и связями несущественно, и мы можем свободно комбинировать их, порождая пространства разной размерности.

Понятно, что, вообще говоря, пространства одной размерности могут быть устроены очень различно, в зависимости от того, в каком порядке мы добавляем измерения и налагаем связи. На практике некоторые такие последовательности могут быть просто невозможны. Но в идеальном случае, предполагая допустимость произвольных построений, размерность пространства со связями есть общая характеристика всех возможных обращений иерархии (иерархических структур). Например, с точки зрения теории атома, описание движения электрона и дырки требует решения трехчастичной задачи (включая поле атомного остатка); однако при этом атом электрически нейтрален, и его сложная структура проявляется только при близком контакте.

## 3. Проекция

Подобно связям, проекторы эффективно уменьшают размерность пространства — но делают это иначе. Они соотносят  $N$ -мерное пространство с некоторым другим пространством размерности  $Q$  (пространством компонент), которое можно считать внутренним для каждой точки исходного пространства. Измерения этого внутреннего пространства и называются проекциями; их размерность равна  $N/Q$ . В частности, внутренние компоненты одномерного пространства имеют размерность  $1/Q$ .

Таким образом можно строить пространства любой рациональной размерности. Накладывая связи не на точки исходного пространства, а на их проекции, мы получаем также пространства отрицательной рациональной размерности. Проекция элементарной связи имеет размерность  $-1/Q$ . Легко видеть, что в исходном пространстве это соответствует обычным процедурам ортогонализации, и проекция «векторов» исходного пространства на ортогональные им внутренние измерения дает ноль.

Точку исходного пространства можно построить по полному набору ее проекций. Для внутреннего пространства это соответствует построению пространства  $Q$  измерений. В терминах «декартовых степеней», мы тогда получаем естественное соотношение:

$$(\mathbf{R}^{N/Q})^Q = \mathbf{R}^N.$$

Разумеется, в реальных приложениях следует позаботиться об «ортогональности» проекций, которая в нелинейной теории достижима только локально. Но сути дела это не меняет.

Вещественные числа мы определяем как классы сходящихся последовательностей рациональных чисел (или сечения). Следуя той же логике, последовательности (иерархические структуры) пространств рациональной размерности порождают вещественные размерности.

Отметим, что вещественной здесь оказывается не топологическая (внешняя), а геометрическая (собственная) размерность пространства. В общем случае топология не вытекает из геометрии, и наоборот. В некоторых частных теориях способы построения фракталов можно сопоставить с классами пространств рациональной размерности — и тогда топологическая размерность может быть понята как геометрическая; возникает своего рода изоморфизм. Но еще раз подчеркнем: изоморфизм не есть тождество. Например, точки отрезка  $(0, 1)$  однозначно отображаются на отрезок  $(1, 2)$ , с полным сохранением структуры, — однако это не означает, что  $x = x + 1$ .

#### 4. Индекс

Индексация пространства противоположна проекции: вместо разворачивания внутреннего пространства каждой точки мы, наоборот, сопоставляем ей нечто внешнее — и тем самым эффективно увеличиваем размерность. Это внешнее — «имя» точки, навешенный на нее ярлык, который, вообще говоря, может меняться при переходе от одной системы индексирования к другой. В общем случае имя может быть чем угодно — и даже вообще не подпадать под «юрисдикцию» математики. Таковы, например, физические поля или топонимы. Однако и в математике индексированные пространства не редкость. Такова, например, все координатные системы: каждой точке пространства мы ставим в соответствие набор чисел, упорядоченный в согласии с принятым порядком пространственных измерений. Это элементарное индексное пространство размерности 1 — «вектор». С другой стороны, мы можем сопоставить точке не вектор, а матрицу, «тензор ранга 2». Компоненты тензора мы перечисляем двумя индексами, и если исходное пространство имеет размерность  $N$ , компоненты тензора образуют пространство размерности  $N^2$ . Очевидно, индексирование  $k$  индексами соответствует возведению размерности базового (конфигурационного) пространства в степень  $k$ .

Казалось бы, мы могли бы определить возведение в степень как многократное умножение:

$$N^k = N \cdot \dots \cdot N \text{ (} k \text{ раз)}.$$

Для размерности квадрат определятся, вроде бы, как

$$\mathbf{R}^{N^2} = (\mathbf{R}^N)^N \sim \mathbf{R}^N \times \dots \times \mathbf{R}^N \text{ (} N \text{ раз)},$$

и тоже можно строить длинные цепочки. Проблема в том, что присутствующее здесь многоточие не элементарная операция, а «квантор»: ее нельзя определить в терминах исходного объекта — это другой уровень логики. По сути дела, предполагается некоторая деятельность, процесс в базовом пространстве. Иногда этот процесс можно (до какой-то степени) алгоритмизировать; чаще он понимается неформально — и это неисчерпаемый источник все новых математических структур. Учитывая, что повторение целое число раз есть лишь частный случай (коль скоро мы взялись обсуждать пространства произвольной размерности), такое «простейшее» определение нас не устраивает; поэтому мы с самого начала признаем возведение в степень операцией особого рода, не сводящейся, вообще говоря, к умножению. Однако для небольшого целого числа индексов, при некоторой системе индексирования, возможно установить соответствие (изоморфизм) между полученными разными путями пространствами и соблюсти «принцип соответствия».

Легко видеть, что индексирование воспроизводит привычные свойства степени:

$$1^k = 1, N^1 = N.$$

Действительно, если каждый индекс может принимать только одно значение, то у объекта с любым числом индексов будет только одна компонента; если индекс только один, то число компонент равно размерности базового пространства.

Каждый индекс пробегает измерения базового пространства в определенном порядке. Как уже говорилось, это соответствует одному из возможных обращений иерархии, способу ее разворачивания. То же самое справедливо и в общем случае, когда конструирование пространства



включает наложение связей и проекции. Последовательность применения «конструкторов» есть полный аналог пространственного измерения по отношению к набору индексов. И точно так же, на индексы могут быть наложены связи, и возможно развертывание внутренних пространств. Количество индексов (или компонент индексного пространства) тогда выражается любым рациональным (в пределе вещественным) числом. Так определяется произвольная вещественная степень размерности.

Индексное пространство размерности 0, очевидно, задает скалярное поле — числовую функцию на базовом пространстве. Ясно, что любая размерность в нулевой степени дает ноль. Будем также считать, что любая степень нульмерного пространства порождает индексное пространство без компонент. Здесь, однако, есть альтернативная возможность: можно строить разные бескомпонентные объекты, которые могут иметь сколько угодно индексов — но любая их комбинация ни на что не ссылается и ни от чего не зависит. Это вполне аналогично тому, как если бы мы различали комплексные числа нулевой (или бесконечной) амплитуды с разными фазами (что, разумеется, не принято в большинстве современных математических теорий).

Степень  $(-1)$  размерности  $N$  есть связь порядка  $N$  в индексном пространстве, эквивалентная пространству размерности  $(-N)$ . Для тензоров можно представить связь наглядно как нижний (ковариантный) индекс, в отличие от контравариантных индексов для измерений индексного пространства. Наложение такой связи — просто свертка связи с одним из верхних индексов, так что общее количество индексов уменьшается на единицу, в полном соответствии с интуитивно ожидаемой картиной. Однако возможны и связи другого рода, непредставимые непосредственно в виде степени какого-либо пространства: в частности, фиксация одной из компонент объекта-степени (или некоторой линейной комбинации компонент) есть связь размерности  $-1$  на индексном пространстве. В общем случае связываются значения многих компонент; по отношению к базовому пространству связи в индексном пространстве являются *симметриями*. Они не меняют размерности базового пространства, но могут накладывать ограничения на динамику (вспомним еще раз о различии геометрической и топологической размерности). Когда таких ограничений достаточно много (не меньше размерности), симметрии становятся связями.

Квадрат связи (пространства размерности  $-1$ ) по смыслу есть связь, наложенная на связь (ограничение связи); эффективно это соответствует добавлению степени свободы. Таким образом, для размерностей  $(-1)^2 = 1$ .

### 5. Ветвление

Индексирование (возведение размерности в степень) есть переход от одного уровня иерархии размерности к другому, объекты которого структурированы иначе, чем исходное пространство. Для приведения их к «пространственному» виду требуется особая операция — перечисление компонент. Такое упорядочение, разумеется, можно проводить по-разному. Принципиально оно не отличается от перечисления измерений обычного пространства — однако необходимость «снятия» степени, перехода от многоиндексной величины к единственному индексу (более высокого порядка) остается, и это не чисто формальная, а практическая операция, связанная с выбором предметной области. При наличии связей, такой переход можно сравнить с каноническими преобразованиями в аналитической механике.

Вообще говоря,

$$\mathbf{R}^{(N^k)^{1/k}} \neq \mathbf{R}^N.$$

Каждое возведение в степень (хотя бы и дробную) есть переход на более высокий уровень иерархии, который не сводится к нижележащим, поскольку возможны разные обращения иерархии — иерархические структуры. Различные «прообразы» пространства-степени можно назвать ветвями (листами, репликами) базового пространства. Это объекты более высокого уровня, которые имеют разную размерность, но при этом «изоморфны» друг другу. Например, двухиндексное пространство размерности 1 может быть получено как квадрат одномерного

пространства — или квадрат связи. В этом случае выделяются две ветви (два обращения иерархии пространства-степени). Если индексное пространство имеет более сложное строение, ветвей будет больше (при вещественной степени — вплоть до бесконечности).

Особый интерес представляют ветви для связи на индексном пространстве ранга 2 (связи на компоненты квадратной матрицы). Такая связь имеет геометрическую размерность  $-1$ , а поиск ветвей формально требует извлечения квадратного корня. Так мы приходим к понятию мнимой размерности  $(+i)$  и мнимой связи  $(-i)$  — и к теории пространств произвольной комплексной размерности, которые в простейшем случае представляются в виде «декартова» произведения вещественной и мнимой части.

1984

### Квантовая теория множеств

Основные понятия классической теории множеств неформальны: множество, элемент, принадлежность и непринадлежность, тождество, равенство и неравенство. В зависимости от принятых правил употребления появляются различные классические теории; но формальная аксиоматика — это не определение, а всего лишь (выражаясь физическим языком) «связь», ограничение на возможные конструкции, не снимающее многозначности базовых терминов. Когда заходит речь об аналогиях с физикой, приходится выбирать одну из возможных интерпретаций; это нормально для всякой науки, и не вызывает особых проблем, пока мы не придаем ни одной из возможных абстракций преимущества перед другими, помним, что у каждой теории есть своя область применимости. С другой стороны, мы никак не ограничены в выборе возможных моделей — и потому смело развиваем сколь угодно экзотические идеи в надежде, что для чего-нибудь и они сгодятся.

Традиционные математические объекты по своей сути статичны: предполагается, что они просто даны — а математик лишь изучает свойства этой данности. Это вполне подобно тому, как мы относились к миру до XX века, в классической физике. В этом контексте, множество есть само по себе, и оно никак не зависит от наших манипуляций с ним. Кто-то его для нас изготовил, и мы спокойно работаем, зная, что никуда оно от нас не убежит. Соответственно, если есть несколько множеств, мы можем делать с ними что угодно, и результат будет одинаков, для всех и всегда.

Основное свойство всякого множеств — наличие элементов. Возможно, элементы есть и у чего-то еще, кроме множеств, — но то, у чего нет элементов, заведомо множеством не является. В частности, термин «пустое множество» — не более чем сокращение для текста вроде: не существует множества, такого, что... Можно придать этому понятию и объектное содержание (выбирая одну из возможных моделей), однако множеством он не станет все равно.

Множества бывают разные. Когда множество маленькое, мы можем перечислить все его элементы — охватить их одним взглядом, или хотя бы указать способ перебора всех элементов за конечное (то есть, укладывающееся в рамки данной деятельности) время. Для очень больших множеств такое в принципе невозможно; в лучшем случае, мы можем указать, на что это похоже: отрезок прямой, совокупность функций, или еще что-нибудь. Иногда и этого сделать не удастся; кое-кто даже отказывается считать такие необъятные объекты множествами. Но допустим, что мы каким-то способом заполучили такого монстра. Перед нами две проблемы: с одной стороны, если мы по жизни столкнулись с чем-то достаточно определенным, встает вопрос, не является ли оно элементом данного множества, — то есть, должна существовать эффективная процедура определения принадлежности чего угодно нашему (пусть даже очень большому) множеству. С другой стороны, если мы хотим, чтобы это было множеством, нам надо уметь предъявить публике хотя бы один его элемент — то есть, мы знаем практический способ запустить руку внутрь множества и вытащить («выбрать») нечто, ему заведомо принадлежащее (причем именно

как элемент, а не подмножество). Обе эти процедуры могут оказаться весьма нетривиальными; собственно, их выработкой как раз и занимается любая из областей классической физики: всякая наука в процессе своего развития определяет свой предмет.

О принадлежности в классической теории множеств говорить можно только если мы явно конструируем множество — то есть исходим из уже готовой предметной области и лишь ограничиваем (прямо или косвенно) право ее элементов принадлежать данному множеству. Никакой универсум принципиально неопределим в рамках теории множеств; имеющиеся на этот счет предложения всегда содержат логический круг: заранее предполагается то, что мы хотим в итоге получить.

Выборки тоже дело не простое. Но если все же мы как-то договорились о технологиях, множества по отношению к процедуре выборки могут вести себя по-разному, и это не произвол математика, а веление той предметной области, к которой его математика прилагается. Традиционная теория множеств имеет дело с такими совокупностями, в которых не может быть больше одного объекта каждого типа. Если мы вытащили при выборке элемент — получается другое множество, в котором такого элемента уже нет. Множество как система распадается на две части: элемент и остаточное множество (которого в каких-то случаях может и не оказаться). Такие превращения нам хорошо известны из физики и химии. Разумеется, можно обойтись и без человеческого вмешательства, если множества в рамках какой-то теории взаимодействуют и обмениваются элементами.

Противоположный вариант — несколько элементов одного типа в одной совокупности, которую мы должны были бы назвать уже иначе: например, «набор» — вместо «множество». В тех случаях, когда существует эффективная процедура определения количества элементов одного типа, набор оказывается множеством; вообще говоря, это не так. Для набора-множества есть разные представления. Например, двухуровневая структура, в которой на низшем уровне все элементы различимы, а на верхнем — объединены в классы эквивалентности. Соответственно, полное количество элементов на нижнем уровне — и количество классов эквивалентности на высшем. В статистическом представлении — мы говорим о том, что каждый элемент принадлежит множеству с какой-то вероятностью: количество элементов одного типа отнесено к полному количеству элементов. Для больших множеств статистическое описание может оказаться предпочтительным (или даже единственно возможным). Промежуточные представления имеют дело со статистическими весами произвольного вида и соответствующими статистическими суммами; конкретный выбор мы делаем из чисто практических соображений.

Однако в любом случае классическая теория имеет дело с «готовыми» множествами, исследовать которые можно сколь угодно долго. Любая структурная перестройка для классического наблюдателя выглядит как «сингулярность», «катастрофа», или «фазовый переход» — конец одного и начало другого. Нас интересуют плавные изменения в зонах непрерывности; на их стыках новые устойчивые образования успевают вполне сформироваться от одного акта «измерения» к другому.

Квантовый эксперимент отличается от классического прежде всего вмешательством наблюдателя в поведение наблюдаемой системы: сначала мы готовим то, что собираемся наблюдать — а потом уже соображаем, что, собственно, мы приготовили. Классическая физика обычно следует той же схеме — но здесь этап подготовки никак не связан с этапом измерения, они разнесены в пространстве, во времени, или еще как-нибудь; грубо говоря, приготовленная система живет достаточно долго и успевает забыть о тех, кто ее приготовил, до того как кто-то другой начнет ее изучать. Квантовую систему мы употребляем сразу же, иногда в самом процессе ее приготовления. Похожее различие существует между кинокартиной — и театральной постановкой, между перепиской — и живой беседой. Разумеется во всем есть свои градации, и различие квантовых множеств от классических существует лишь на одном уровне иерархии, при определенном способе ее развертывания.

Квантовая динамика целиком помещается в классические точки сингулярности:

классическое движение до и после — для квантовой теории лишь начальное и конечное состояние, асимптотика. Самое интересное внутри — но непосредственно наблюдать мы его не можем (не превращая тем самым в классическое движение) и должны догадываться по косвенным признакам, сопоставляя структуры на входе и на выходе.

Квантовое множество (или, вообще говоря, набор) — это как-то приготовленная система в определенном состоянии, которое мы абстрактно представляем «вектором» состояния  $|A\rangle$ . Возможные в данной схеме приготовления множества все вместе мы метафорически объединяем в некоторый универсум — «конфигурационное пространство». Чтобы выяснить, принадлежит ли элемент  $a$  множеству  $A$ , мы вычисляем число  $\langle a|A\rangle$ , квадрат модуля которого дает нам степень принадлежности элемента  $a$  множеству  $A$  (статистический вес). Пользуясь векторной метафорой, мы можем сказать, что  $\langle a|$  представляет некий «функционал» над пространством множеств данного уровня. Совокупность таких функционалов и задает предметную область теории. Иначе говоря, это как раз то, что мы в результате своей деятельности хотим получить, ее продукт.

Одноэлементное множество, содержащее только элемент  $a$ , можно обозначить как  $|a\rangle$ ; при этом  $\langle a|a\rangle = 1$  (в общем случае единица может быть не числом, а чем-то вроде  $\delta$ -функции — то есть, еще одним функционалом), а для любых других элементов  $b$  из предметной области теории  $\langle b|a\rangle = 0$ .

Здесь пока не видно существенных отличий от классической теории множеств (наборов): переход от вероятностей к амплитудам ничего не меняет сам по себе, это лишь своего рода «замена переменных», «подстановка» — смена точки обозрения; на выходе вполне классические веса элементов, и польза от нововведения не очевидна (а иногда и сомнительна). Так оно и будет, пока мы говорим об уже приготовленном множестве, на которое мы не оказываем никакого влияния, а только лишь измеряем степени принадлежности. Кинематическая картина всегда формируется на уровне макроскопического (классического) наблюдателя — поскольку в конце всего нам нужны совершенно практические вещи, пригодные к повседневному употреблению. Все различие — в характере динамики: классические множества взаимодействуют на уровне наблюдаемых, квантовые — на уровне амплитуд.

Поскольку речь идет о математике (науке о структурах вообще), динамика не фигурирует в теории непосредственно, она должна быть представлена особыми структурами. В квантовой парадигме, мы ассоциируем всякие изменения состояния (движение в конфигурационном пространстве) с «операторами». С одной стороны, речь идет о преобразовании множеств — то есть, об операциях над множествами в заданном универсуме; с другой стороны, всякое сопоставление одного множества с другим есть тоже своего рода преобразование — и потому различные отношения между множествами также представимы операторами — хотя, возможно, иного типа.

В частности, отношение принадлежности элемента множеству требует переосмысления. Казалось бы, простая вещь: мы не имеем право непосредственно сопоставлять объекты разного типа (разных уровней). Можно сравнивать элементы с элементами, а множества с множествами; для сопоставления элементов с множествами надо как-то привести их к общему типу. Традиционная запись  $a \in A$  — лишь сокращение для последовательности нетривиальных процедур, каждая из которых требует определенных условий осуществимости. Большинство этих условий явно не оговаривается, и математик в любом рассуждении заранее предполагает, что его мир устроен достаточно хорошо, чтобы получился именно тот результат, который предполагается получить. Логические затруднения приводят к тому, что некоторые математики вообще отказываются говорить об элементах — и сравнивают только множества; однако на само деле это лишь заматывание мусора под ковер: те же трудности обязательно появятся в других местах, в другом виде.

В квантовой теории множества представлены «векторами состояния», а элементы — «функционалами». Различие сразу бросается в глаза. Сопоставить одно другому — дело непростое, и здесь возможны разные технологии (и, соответственно, весьма разные теории). Однако в любом случае есть два взаимно дополнительных подхода: либо мы элементы преобразуем в множества и сравниваем множества — либо, наоборот, множества сводим к элементам. Третий путь, когда и элементы, и множества приводят к чему-то синтетическому, означает на деле выход за рамки собственно теории множеств.

Первый способ опирается на особую операцию над множествами — проекцию; по смыслу это выделение «части» множества (в классической теории — подмножества). Для выделения единичного элемента  $a$  соответствующий оператор записывают как  $|a\rangle\langle a|$ , так что проекция изображается как  $|a\rangle\langle a|A\rangle$  — что выглядит, как будто элементу (функционалу) сопоставлено некоторое одноэлементное множество (вектор); таким образом, по отношению к данной предметной области множество предстает линейной комбинацией одноэлементных множеств:

$$|A\rangle = \sum |a\rangle\langle a|A\rangle = \sum |a\rangle\psi_a(A).$$

То же самое множество можно отнести к другой предметной области и получить какое-то иное представление:

$$|A\rangle = \sum |b\rangle\langle b|A\rangle = \sum |b\rangle\psi_b(A).$$

Формально это «условие полноты» записывают как

$$\sum |a\rangle\langle a| = 1, \quad \sum |b\rangle\langle b| = 1,$$

однако надо помнить, что далеко не всегда пространства элементов  $a$  и  $b$  совпадают — это вовсе не переход от одного «базиса» к другому, а рассмотрение того же самого с какой-то другой (иногда противоположной) стороны. Например, граф можно представить как множество вершин, соединенных стрелками — или как множество стрелок разделенных вершинами; управлять компьютером можно либо с клавиатуры, в командной консоли, — либо мышью, в графическом интерфейсе. Суть дела при этом не меняется — а выглядит очень по-разному.

Заметим, что объем базиса, вообще говоря, никак не связан с мощностью множества. Например, в атомной физике одни и те же величины можно вычислять либо интегрированием по состояниям непрерывного спектра — либо суммированием по дискретному базису. В логике базис из двух элементов **T** и **F** используется для оценки самых разных высказываний — их предметное содержание и смысл для логической оценки не существенны. Одну и ту же задачу можно решить надежным, но громоздким способом — или нестандартно и элегантно.

Вообще говоря, продукт деятельности отличается от ее объекта (исходных материалов и технологий). Однако бывают ситуации, когда и объект, и продукт рассматриваются узко, лишь в одном из возможных аспектов — и различие снимается. Так, в рыночной экономике материальное и духовное производство предстает движением меновой стоимости, а в структуре научной теории дедукция переходит от одних истинных суждений к другим.

На практике удобнее работать в ортогональном базисе. Но, как и в классической теории, это вовсе не обязательно: присутствие одного элемента может (в какой-то мере) предполагать присутствие другого. В квантовой теории ортогональность означает, что каждое одноэлементное множество является собственным вектором некоторого фиксированного оператора.

Переход от одного базиса к другому формально записывается как

$$|A\rangle = \sum |a\rangle\langle a|A\rangle = \sum |b\rangle\langle b|\rho|a\rangle\langle a|A\rangle.$$

То есть,

$$\psi_b(A) = \sum \rho_{ba}\psi_a(A).$$

В простейшем случае, когда базисы  $a$  и  $b$  определены в одном и том же пространстве, «оператор

плотности»  $\rho$  может обращаться в единицу — и речь идет о переходах между эквивалентными представлениями. В общем случае требуется привести одно пространство к другому, чтобы сделать их сопоставимыми — способ такого приведения зависит от конкретной предметной области и характера задач.

Трактовать множества как совокупность элементов — значит уметь явно конструировать их. Поскольку в классической теории конструирование отделено от наблюдения, появляется иллюзия одновременного присутствия всех элементов множества: они все вместе попадают в поле зрения и один элемент никак не выделяется на фоне других. Квантовая теория множеств представляет операции добавления или изъятия элемента соответствующими операторами: запись  $a^+ |A\rangle = |a; A\rangle$  говорит о том, что после действия «оператора рождения»  $a^+$  на множество  $A$  образуется новое множество, скорее всего (но, как указано ниже, вовсе не обязательно) содержащее элемент  $a$ , — то есть, мы ожидаем, что  $\langle a | a; A \rangle \neq 0$ . Обратная операция: действуя «оператором уничтожения»  $a^-$  на множество  $|a; A\rangle$ , мы ожидаем восстановления множества  $A$ . Таким способом можно явно конструировать любые конечные расширения исходного множества  $A$ , которое в данном случае играет роль «вакуума» — фонового состояния, по отношению к которому мы рассматриваем все остальное. Может появиться соблазн взять в качестве вакуума пустое множество — и строить теорию абсолютным образом. Но пустое множество — это не множество, и мы не имеем права обращаться с ним как с обычными множествами (и в частности, действовать на него операторами, определенными для множеств). Точно так же в современной физике вакуум — лишь условность, уровень отсчета. Встречая в физическом тексте формулу вроде  $a^- |0\rangle = 0$ , мы должны понимать это как идиому, условное обозначение для совокупности наложенных на физическую систему связей; во многих случаях оператор уничтожения частицы может рассматриваться как оператор рождения античастицы:  $a^- |0\rangle = |\bar{a}\rangle$ , и возможны состояния с одновременным присутствием нескольких частиц и античастиц (вроде электрон-позитронной пары, или системы свободный электрон + дырка при ионизации атомов). Точно так же,  $a^- |A\rangle = |\bar{a}; A\rangle$  (множество с «дыркой», анти-элементом); как именно следует реализовать добавление или уничтожение элемента — зависит от предметной области. Например, добавление элемента не обязательно связано с его возникновением, а лишь увеличивает количество элементов того же вида (наподобие добавления электрона к атому или иону); точно так же, уничтожение элемента — уменьшает «вес» элементов этого вида в множестве: то есть, элемент и дырка тут же «аннигилируют», не оставляя в теории существенных следов. В простейшем случае, когда множество может содержать только один элемент данного вида, операторы рождения и уничтожения идемпотентны:  $a^+ a^+ = a^+$ ,  $a^- a^- = a^-$ .

Последовательное действие нескольких операторов рождения/уничтожения порождает многоэлементные множества:

$$a^- c^+ b^+ a^+ |A\rangle = |\bar{a}, c, b, a; A\rangle.$$

Порядок действий может иметь значение, и лишь в простейших случаях множество  $\{\bar{a}, c, b, a\}$  равно множеству  $\{c, b\}$ .

Когда квантовые элементы «интерferируют» внутри множества (а здесь могут быть самые разные возможности), состояние  $|a; A\rangle$  уже нельзя представить как  $|a\rangle |A\rangle$ . Лишь очень простые множества представляются произведением одноэлементных множеств (собственных векторов некоторого оператора) — такие состояния (а также их невырожденные линейные комбинации) называют «чистыми». При наличии «полного» базиса, мы можем «перечислить» элементы исходного множества:

$$|A\rangle = \sum |z\rangle \langle z | A \rangle,$$

и тогда

$$|a; A\rangle = a^+ |A\rangle = \sum |b\rangle \langle b| \rho a^+ |z\rangle \langle z| A\rangle = \sum |b\rangle \rho_{bc} \langle c| a; z\rangle \psi_z(A).$$

Так присоединение элемента к множеству из многих элементов сводится к объединению двухэлементных множеств: новый элемент  $a$  поочередно связывается с каждым из элементов исходного множества. Это очевидным образом соотносится с аналогичным разложением для традиционных множеств:

$$\{a\} \cup A = \bigcup_{c \in A} \{a, c\},$$

однако в квантовой теории требуется еще и учет интерференции различных путей виртуального перехода от одного состояния к другому.

Поскольку операторы рождения и уничтожения мы сопоставляем не множествам, а элементам, объединение двух произвольных множеств, вообще говоря, не определено. Но, например, если два множества получены из одного и того же исходного множества («базы», «вакуума»), появляется возможность рассматривать композицию соответствующих операторов порождения как оператор порождения объединения. Это подобно тому, как один и тот же атом может иметь различные степени ионизации. Точно так же в арифметике мы порождаем натуральные числа из единицы при помощи единственной операции — инкремента. Сумма натуральных чисел — это особый объект, привнесенный в теорию извне, и определенный, вообще говоря, на другом уровне, как класс изоморфизмов.

В некоторых случаях возможно определить объединение множеств, полученных расширением разных баз; это соответствует переходу от атомной к молекулярной физике, когда в образовании целого участвуют лишь «валентные» электроны и дырки, а соответствующие атомные остатки считаются изолированными (то есть, вакуум объединения множеств равен произведению их баз). Здесь также бывают многочисленные вариации. Аналогом классического объединения будет нечто вроде ковалентной связи, когда все валентные электроны в одинаковой степени принадлежат каждому атому в связи. Другой вариант — нечто вроде ионной связи, когда элементы одного множества связываются с «дырками» другого. Возможны также «водородные связи» — когда реального объединения нет, но для разных множеств используется общий базис.

Обобщение конечных расширений на бесконечные (счетные и несчетные) не представляет особого труда. С практической точки зрения это означает переход на следующий уровень, рефлексивность деятельности: вместо выполнения единичных операций мы конструируем сами эти операции по единому правилу. Как и любую иерархию, порождение множеств можно свернуть в «точку» — и развернуть в сколь угодно развитую иерархическую структуру.

Все возможные «взаимодействия» между множествами (теоретико-множественные операции и отношения) выражаются через комбинации операторов рождения и уничтожения. Каждая конкретная теория опирается на некоторый набор базовых (элементарных) взаимодействий, а все остальное может быть сведено к их различным комбинациям. Поскольку в конечном итоге нас интересует вполне определенный продукт, мы приводим результат действия каждого оператора к единому базису; то есть, разные комбинации элементарных операций (последовательности взаимодействий) могут привести к возникновению одного и того же (в смысле практической неразличимости) множества, а интерференция этих виртуальных процессов приводит к особенностям глобальной организации (спектра) множества-продукта, вплоть до невозможности получить какие-то множества из исходного (правила отбора).

Тут самое время задуматься о сравнении множеств. Что значит — «одно и то же»? Равенство элементов — вопрос практический, оно определено строением предметной области. Что такое равенство множеств — вопрос непростой. Традиционно полагают, что два конечных множества равны тогда, и только тогда, когда они состоят из одних и тех же элементов. Даже такое, вроде бы, интуитивно ясное определение на самом деле опирается на весьма серьезные допущения о предметной области и способах перебора элементов (начиная с самой возможности

перебора). Для бесконечных множеств трудностей еще больше: надо успеть за конечное (или даже бесконечно малое) время перебрать все элементы одного множества и проверить наличие их в другом, и наоборот. Сразу вспоминается квантовая механика, где тонкие процессы взаимодействия «упрятаны» в одну макроскопическую точку, а на выходе мы имеем готовую интегральную оценку — спектр. Однако и в классической теории можно представить себе простую схему сравнения, в которой одно из множеств превращается в фильтр, своего рода барьер, на который накатывается поток элементов другого множества: совпадающие элементы поглощаются (или отражаются назад) — и если на выходе (по другую сторону барьера) ничего нет, можно утверждать, что рассеянное множество заведомо меньше множества-фильтра, это его подмножество. Обращая ситуацию (меняя поток и фильтр местами), проверяем обратное отношение: если на выходе опять пусто — значит, два множества равны.

Понятно, что это лишь одна из возможных реализаций; но такая мысленная конструкция выпукло высвечивает главное: сравнение множеств неизбежно опирается на многочисленные предположения о динамике, о характере взаимодействия. Меняем постановку эксперимента — и рискуем получить другой результат. Ситуация не нова: например, в математике сосуществуют различные определения размерности — и приходится устанавливать соответствия. Точно так же, вместо равенства бесконечных множеств мы говорим лишь об их равномощности или, самое большее, об изоморфизме. Но всякому нормальному человеку понятно, что, например, звучание фонемы физически отличается от письменного знака, который условно сопоставлен ей в МФА, и одно в другое никак не превратить: два множества лишь подобны, но не равны. Точно так же, четные натуральные числа — вовсе не то же самое, что нечетные, хотя одно множество однозначно сопоставимо с другим. Нотную запись мелодии можно воспроизвести на скрипке, на фортепиано, на органе... — однако в оркестре эти инструменты далеко не «изоморфны».

Квантовая парадигма вносит свои коррективы, добавляя «встроенную» неопределенность, «частичную» принадлежность. Однако процедура «фильтрации» одного множества другим переносится в квантовую теорию множеств без изменений; более того, превращение множества в фильтр превращается в простой формальный трюк: надо все операторы рождения элемента для одного из множеств заменить на соответствующие операторы уничтожения — так, чтобы получившиеся «дырки» («антиэлементы») аннигилировали с элементами другого множества, и тогда на выходе мы сразу получим спектр различий. То есть, для множеств

$$\begin{aligned} b^+ a^+ |A\rangle &= |b, a; A\rangle \\ y^+ x^+ |Z\rangle &= |y, x; Z\rangle \end{aligned}$$

результат сравнения дан амплитудой

$$\langle y, x; Z | b, a; A \rangle = \langle Z | x^- y^- b^+ a^+ | A \rangle = \sum \langle Z | \nu \rangle \langle \nu | x^- y^- b^+ a^+ | \mu \rangle \langle \mu | A \rangle = \sum \bar{\psi}_\nu(Z) D_{\nu\mu} \psi_\mu(A).$$

В каких-то допущениях, равенство элементов  $a = x$ ,  $b = y$ , превращает это выражение в  $\langle Z | A \rangle$ ; при равенстве баз, это просто единица. Разумеется, в более сложных конструкциях равенство единице не гарантирует равенства множеств; но если оказывается, что виртуальные переходы до такой степени нейтрализуют друг друга, — значит, одно из множеств может быть эффективно преобразовано в другое, и в практическом плане они равны. Рассматривая это выражение как функцию от каких-либо «макроскопических» параметров, можно получить полноценный спектр, и внутренние компенсации проявятся в его структуре (например, как резонансы).

Квантовая теория множеств не просто расширяет классическую — она фактически предлагает пучок конкретных реализаций, в зависимости от конкретных задач. Точно так же в физике «теории всего» допускают много разных «ландшафтов». Выбор решения — не произвол, решает практика. Такая, практически ориентированная математика перестанет быть всего лишь игрой ума и станет по-настоящему осмысленной.

ноябрь 1985



### Абстрактные картинки

Наука математика (как и всякая наука вообще) сводит наши обыденные проблемы к формальностям. И это правильно. Потому что иначе приходилось бы заниматься только повторением пройденного, и на творческие задумки уже не хватало бы сил. Но пока человек отличается от пчелы, ему недостаточно знать, как действовать, — ему еще надо бы понять, как это выглядит в целом, со стороны, из другой вселенной... Это называется: интуиция. Когда оно есть — даже формальности становятся излишними, и дело движется вперед семимильными шагами.

Но люди разные. Интуиция одного может оказаться совершенно неприменима к другому. Кому-то вполне достаточно представлять себе последовательность действий — «дао», процессуальность как таковую. А всякие там картинки — это как результат. Другим подобная алгоритмичность претит: им, наоборот, что-нибудь нарисованное, чтобы со всех сторон посмотреть, а потом уже решать. Одни в музыке ловят мелодию — другим ближе гармония. Третий путь — инструментальная интуиция, практическое чутье, умение на лету выстроить, нарисовать, сыграть... Для этого и слово другое придумали: талант.

У математиков интуиция особого рода. Порядок у них выражается числом — а картинками заведует геометрия. Однако числа постепенно превращаются в абстрактные структуры, которые вполне возможно представлять себе геометрически; наоборот, геометрия растворяется в гомеоморфизмах и сводится к числовым (топологическим) инвариантам. И талант математика перерастает в логику, умение усматривать главное, абстрагироваться от мелочей.

По жизни, нам как раз мелочи интереснее всего. То, что можно постепенно разглядывать. Но сами по себе — мелочи превращаются в хаос, и это скучно. Поэтому важно прилепить их к какой-нибудь абстракции, чтобы чувствовать в зыбях твердость и не расстраиваться зря. Если же абстракция к тому же еще и наглядна — тут никакие бури не страшны, и занудный штиль можно перетерпеть.

Вот и давайте пофантазируем на тему визуализации чисел.

Числа, как известно, бывают разные — тут все как у людей. Начинается, натурально, с простого перечисления, подсчета. Раз, два, три, четыре, пять — дальше можно не считать. Почему? Да потому что уже есть общее представление о процессе. То есть, нарисовалась некое направление, вдоль которого мы равномерно шагаем.

Рациональные числа при таком раскладе нормально увидеть как парочку независимых направлений — и по каждому свои шаги. Возникает нечто вроде плоской решетки. Тут черт приносит какого-то работягу с дурацким вопросом: а что, если мы будем шагать разными шагами в одном и том же направлении? Сумеем прийти к чему-то общему или нет? Оказывается, все зависит от размера шагов. В каких-то случаях можно договориться, в других — ни в какую. Это называется соизмеримость. Или рациональность. Когда все рациональны — можно навести порядок, и снова все числа представляются точками на прямой, и всегда есть возможность сказать, чего больше, и чего не хватает. В чем отличие от натуральности? В том, что приходится каждый раз согласовывать наши мерки, и одно и то же называется по-разному, в зависимости от того, чем его измерять.

Прекрасно. Вот, мы договорились меж собой — и все довольны. Нет, не все. Остаются всякие там несоизмеримые, иррациональные... И прямо как в жизни: таких намного больше, чем пробившихся в рациональную элиту! Их так много, что игнорировать никак не получается. Именно они — основная масса, вещество той самой прямой, на которой рациональные размещаются просторно и с комфортом; потому и название: вещественные. Или действительные. На каждого рационального вещественных по соседству больше, чем всех рациональных вместе взятых. В любой, сколь угодно малой окрестности. С другой стороны, всякое вещественное число с обеих сторон можно сколь угодно плотно зажать между парочкой рациональных — и тем самым восстановить на вещественной прямой полный порядок.

В мире ограниченных возможностей никаких других полностью упорядоченных числовых систем больше нет. Любой порядок представим все той же картинкой: сплошная линия — которую всегда можно вообразить себе совершенно прямой, если закрыть глаза и плавно ехать вдоль. Для этого есть свое название: непрерывность. Колдобины и неуместные препятствия непрерывность нарушают — но если их немного, общее впечатление сохранится. В итоге вещественные числа становятся мерой всему, и любое расстояние можно выразить вещественным числом. Где-то там, за гранью нашего мира имеются и другие, бесконечные числа. Например, целое число  $\omega$  (точнее,  $\omega_0$ ), которое больше любого другого целого (а значит, и любого вещественного) числа. А за ним следует число  $\omega_1$ , и между ними бесконечность (гипер)целых (или гипердействительных) чисел, которые прекрасно упорядочиваются — но бесконечно велики. И так далее. Чтобы не бегать далеко и держать себя в рамках, у человека серьезного есть конкретная задача: движемся мы не просто так, а из точки  $A$  в точку  $B$ . Или в числовых обозначениях: от нуля до единицы. Из всей прямой нас в каждом конкретном случае интересует лишь отрезок  $[0, 1]$ . Тогда при любой длине шага мы таки дойдем до заявленной цели за конечное число шагов. А чтобы не проскочить — всегда можно поставить стенку в конце. То есть мы рассматриваем отрезок вместе с его границей; кому хочется, может называть это сегментом, или замкнутым многообразием.

Как выясняется, точек на отрезке  $[0, 1]$  ничуть не меньше, чем на всей числовой прямой — так что, ограничивая себя, мы ничего не потеряем. Соответственно, и рациональных чисел там очень много. Сколько? Мы толком не знаем — но опять-таки, кто мешает придумать название? Пусть количество целых или рациональных чисел называется «алеф-нуль» ( $\aleph_0$ ), а количество вещественных — обозначим словом «континуум» ( $\aleph_1$ ). При этом есть все основания полагать, что «кардинальное число» алеф-нуль строго меньше, чем континуум. Насчет того, что между ними — есть разные мнения; предположить существование чего-то такого мы можем, но добыть и пощупать пока не удалось. Зато мы точно знаем, что есть количества и побольше континуума!

Действительно, представим себе, что каждую точку отрезка  $[0, 1]$  мы хотим пометить каким-нибудь числом из того же диапазона. Понятно, что сделать это всегда возможно (поскольку количества точек в «области определения» и в «области значений» одинаковы). Тогда мы говорим, что на отрезке задана функция. Некоторые точки, в принципе, могут называться одинаково — такая функция будет необратимой: по имени мы не сможем указать, к чему имя относится. Дело, как говорится, житейское. Зато мы можем взять всех однофамильцев вместе и сказать, что они образуют некую общность — подмножество сегмента  $[0, 1]$ . Так вот, оказывается, что количество таких подмножеств намного больше количества вещественных чисел; соответственно, мощность множества всех ограниченных функций на отрезке больше мощности континуума. И можно придумать еще громаднее. Но зачем? Для начала нам бы представить себе, как выглядит нечто, не помещающееся на прямой...

Казалось бы, а почему не повторить трюк с размещением точек на плоскости в пределах вещественной оси? — как мы переходили от целых чисел к рациональным. Но вспомним: такое «свертывание веера» не увеличило полного количества точек. Мощности множества целых и рациональных чисел выражаются кардинальным числом алеф-нуль. Точно так же, вещественных точек на плоскости ровно столько же, сколько на вещественной прямой — континуум. Однако направление мысли достаточно разумное. Только надо пойти чуть-чуть дальше.

Мы легко можем представить себе квадрат и куб. Обобщение на большее число измерений мыслится как нечто в том же духе. И многие свойства привычного в быту пространства прямо переносятся на абстрактные пространства высших измерений. Разумеется, в каких-то случаях экстраполяции не срабатывают: поиск таких подковырок (а это прекрасная реклама!) — любимое занятие математиков. Но в целом у нас есть интуитивное представление о пространствах высокой размерности в виде некой абстрактной картинки. Мы знаем, что упорядоченная пара точек (по крайней мере, достаточно близко расположенных) задает вектор в пространстве, и у вектора есть длина и направление. Длина — это уже нечто одномерное, и обращаться с этим мы

умеет. Многомерные углы воображение не отрабатывает — но два вектора всегда лежат в одной плоскости, и плоский угол между векторами определить легко по их скалярному произведению. Чуть больше воображения — и можно представлять себе всякие тела, в том числе многомерные. Плавают эдакое нечто внутри гиперкуба.

Измерения многомерного пространства можно упорядочить, перенумеровать. Переход к другой нумерации не меняет геометрии — однако описывать какие-то свойства в одном представлении может быть проще, чем в другом. Последовательность измерений называется ориентацией пространства. Поскольку же мы перечисляем измерения извне, наблюдая пространство в целом как бы со стороны, не всякие геометрические объекты одной размерности можно перевести друг в друга непрерывным образом с сохранением их собственной ориентации.

Теперь обратимся к пространству функций, переводящих отрезок  $[0, 1]$  в себя. Допустим, каждая точка отрезка соответствует измерению многомерного пространства, а результат вычисления функции в точке  $x$  — координата по соответствующему измерению. Тогда любую функцию можно геометрически представлять себе как точку гиперкуба с очень большим числом измерений. Да, их континуум. Но от количества измерений суть пространственности как таковой нисколько не страдает: у нас в распоряжении все те же геометрические образы плоскости и трехмерного тела, и все остальное — по аналогии.

Заметим, что в физике пространства сколь угодно большого числа измерений — далеко не новость. Вектора состояния в квантовой механике могут иметь бесконечно много компонент, свободно перемешивая дискретный и непрерывный спектр с чем-то вообще заоблачным. И никому это по жизни не мешает. На самом деле, там еще круче: есть внутренние, «спинорные», измерения. Но пока не будем замахиваться на такие высоты.

Можно долго и увлеченно забавляться выяснением, какие классы функций соответствуют различным геометрическим объектам: дискретным множествам («кристаллам»), кривым, плоскостям, телам. Например, главная диагональ гиперкуба задает семейство постоянных функций на отрезке  $[0, 1]$ . Сразу ясно, что все подмножества отрезка  $[0, 1]$  расположены в вершинах гиперкуба. Действительно, каждое подмножество задано соответствующей характеристической функцией, принимающей только два значения: либо 0, либо 1. А последовательность нулей и единиц длины континуум как раз и задает одну из вершин. В частности, при «естественной» ориентации гиперпространства, когда последовательность координат совпадает с отрезком  $[0, 1]$ , пустое множество логично попадает в начало координат: это последовательность из одних нулей. Противоположная ему (самая удаленная) вершина — соответствует континууму единиц, представляющему отрезок целиком. Точно так же, любые другие семейства родственных функций можно увидеть внутри гиперкуба. И делать из этого интуитивные выводы.

Геометрия становится собственно геометрией, когда в пространстве можно что-то измерить. В обычных (евклидовых) пространствах работает многомерный вариант теоремы Пифагора. Можно по аналогии говорить о норме функции как расстоянии от начала координат:

$$\|f\|^2 = \int [f(x)]^2 dx$$

Соответственно, расстояние между функциями определяется как длина вектора-разности:

$$\Delta^2 = \int (f_1(x) - f_2(x))^2 dx$$

В частности, расстояние между противоположными вершинами гиперкуба равно 1, а расстояние между любыми «соседними» вершинами (образующими конечномерный гиперкуб) равно 0. Это естественно дополняется понятием угла между функциями:

$$\cos \theta = \int f_1(x) f_2(x) dx / \|f_1\| \|f_2\|$$

В каком смысле определены эти интегралы — мы пока уточнять не будем. Для каждого определения есть своя геометрическая интерпретация.

Некоторые вещи в таком определении выглядят очень естественно. Например, норма постоянной функции просто равняется ее значению. Расстояние между функциями  $f(x) = a$  и  $f(x) = b$  равно  $|a - b|$ , а косинус угла между ними всегда равен единице: вполне ожидаемая параллельность. Квадраты нормы для функций  $f(x) = x$  и  $f(x) = 1 - x$  одинаковы и равны  $1/3$ ; расстояние между ними равно той же величине, а косинус угла равен  $1/2$  (угол в  $60^\circ$ ). Эти две функции принадлежат классу унитарных перестановок на отрезке  $[0, 1]$  (изменение ориентации системы координат). Понятно, что такую «перепутанную» функцию всегда можно упорядочить по значению и тем самым вернуть к виду  $f(x) = x$ ; следовательно, норма любой перестановки равна  $1/3$  — при значительном разбросе взаимных расстояний и углов.

Это возвращает нас к вопросу об определении интегралов. Чтобы вычислить интеграл для произвольной функции, мы сначала упорядочиваем функцию по значению — и потом вычисляем площадь под полученной кривой. Поскольку речь идет об ограниченных функциях, дополнительных технических проблем тут не возникает: любой интеграл оказывается числом от 0 до 1. Разумеется, есть свои маленькие хитрости. Например, если непрерывная функция не является взаимно-однозначной, одинаковые значения попадают в «соседние» точки отрезка, и соответствующую меру (длина элементарного интервала  $dx$ ) придется удваивать, утраивать и т. д. Эквивалентное представление: изменяется «плотность» размещения точек на отрезке. Но это не разрушает геометричности как таковой.

Теперь понятно, как работать с характеристическими функциями подмножеств. После переупорядочения такая функция превращается в «ступеньку»: сначала нули, потом единицы. Норма функции определяется как длина отрезка с единицами — а это есть в точности количество элементов подмножества (его мера). Расстояние между характеристическими функциями есть мера их объединения за вычетом меры пересечения.

Разумеется, у бесконечностей свои капризы. Например, расстояние между постоянными функциями 0 и 1 (противоположные вершины гиперкуба) по общей формуле равняется 1. Но мы же знаем, что диагональ единичного квадрата имеет длину  $\sqrt{2}$ ! А здесь получается, что все стороны равны 1, а диагональ тоже единица! Чудеса.

На самом деле все в норме. По-пифагоровски, диагональ  $N$ -мерного единичного куба равна  $\sqrt{N}$ . При  $N \rightarrow \infty$  мы получили бы непонятность вида  $\sqrt{\infty}$ . Чтобы нормально работать с такими конструкциями, можно все длины нормировать: в конечномерном случае делить на  $\sqrt{N}$  — а в непрерывном случае это превращается в плотность последовательности измерений пространства. Разумеется можно было бы нормировать и двумерный квадрат: для его диагонали мы просто выбираем другую единицу измерения, и в этих новых единицах длина диагонали равна 1, в полном соответствии с бесконечномерным случаем. Решение достаточно логичное: в физике мы таким образом выбираем «естественные» единицы — или, например, привязываем все скорости к скорости света, и это удобно для описания наблюдаемых симметрий.

Куда более существенная проблема — неоднозначность норм, которые определены лишь с точностью до множества меры нуль (интеграл по которому равен нулю). То есть, мы определяем расстояние между семействами функций, а не между конкретными функциями. Для математиков это в порядке вещей: у них всегда так. Человеку просто хотелось бы иметь такое определение расстояния, которое интуитивно соответствовало бы визуальному разделению пространственных точек. Например, окружность (или сфера) с центром во внутренней точке — функции, равноудаленные от некоторой заданной функции. Но отличия меры нуль определяют бесконечное семейство функций, и соответствующие точки всюду плотно распределены по гиперкубу. Как-то не вяжется с геометрической наглядностью. Особенно, если использовать метрику для определения понятий окрестности, открытых и замкнутых классов функций и т. д. Любые топологические построения становятся отнюдь не интуитивными.

Конечно, можно заметить, что класс функций ни при каком раскладе невозможно сделать полностью упорядоченным, впихнуть в действительную ось. На то она и кардинальность выше

континуума. Можно пойти по пути самоограничения и рассматривать только функции достаточно простого вида (например, диффеоморфные постоянной). Тогда близость в гиперкубе совпадет с метрической близостью. И можно рассматривать траектории в пространстве функций, и плавно переходить от одного к другому... Однако многие интересные функции окажутся тогда за бортом. В частности, перестановки и характеристические функции подмножеств. Но есть еще вариант: вместо одного пространства на всех — расслоенное пространство, иерархия классов функций, внутри каждого из которых расстояние будет геометрически определенным, а все вместе можно при необходимости агрегировать в нечто общее. Например, произвольная функция представляется (прямой) суммой компонент: подсистемы конечного множества точек, бесконечные дискретные подсистемы, компоненты мощности континуум (кусочно-связные области). Соответственно, вместо одного — получаем три расстояния, или больше — по числу компонент (уровней). То есть, вместо тупого игнорирования вкладов от множеств меры нуль, мы учитываем их особым образом, как особенности. В другой формулировке — речь идет о сингулярных мерах (типа  $\delta$ -функции). Поскольку соответствующие подпространства ортогональны, естественно определить общее расстояние через сумму квадратов (нормированных) расстояний по отдельным компонентам. Или еще как-нибудь. Геометрия будет ненамного сложнее — а для интуиции простор.

май 1992

### Предметная теория множеств

Традиционная теория множеств ничего не говорит об их элементах. Как правило предполагается, что элементами множества могут быть только другие множества, а множество в качестве элемента ничем не отличается от собственно множества. Такой подход страдает излишней общностью, а неопределенность базовых понятий приводит к парадоксам. Для их преодоления приходится отказываться от изначальной всеобщности, различать множества и классы. Остается сделать еще один шаг: признать, что элемент и множество — не одно и то же. Одно вместо другого нельзя подставлять безоговорочно, и смешение объектов разного уровня в рамках одного рассуждения (формулы) — логическая ошибка.

На практике любая наука применима лишь в пределах своей предметной области. Абстрактный формализм лишь следует организации предмета. Иначе теория оказывается просто бессмысленной — и никому не нужна. Прикладные науки не нуждаются в чрезмерной строгости: им нужнее практические рецепты на каждый день. Разумеется, приложения бывают разные, и некоторые из них могут показаться очень абстрактными; это не меняет сути дела: мы изучаем вполне определенный предмет.

Можно явным образом положить в основу теории некоторую предметную область. Для математики не важно, какую именно: главное, что она есть, и работать мы можем только с реально доступными объектами. Для теории множеств, это некоторый универсум  $U$ , базовый уровень теории. На следующем уровне появляются собственно множества — наборы элементов из универсума  $U$ . В отличие от традиционного подхода, предметные множества не могут непосредственно принадлежать универсуму: это объекты иной природы. Одноэлементное множество  $\{x\}$  — не то же самое, что элемент  $x$ . Отношение принадлежности  $a \in A$  или непринадлежности  $a \notin A$  связывает два уровня иерархии. Мы можем обычным образом записывать небольшие множества, перечисляя их элементы в фигурных скобках:  $\{a, b, c, \dots\}$ . При этом все элементы принадлежат универсуму и не могут быть множествами. На таких условиях нет проблем с формированием множеств по общему признаку: если элементы обладают некоторым предметным свойством  $\varphi$  (в соответствии с природой предмета и логикой теории), они могут рассматриваться вместе как принадлежащие одному множеству. Запись  $\{x \mid \varphi(x)\}$  допустима для любых предметных свойств; более того, всякое множество определяет некоторое

предметное свойство, общее для всех его элементов. Когда элементы четко отличаются от множеств, рефлексивные условия, вроде  $x \in x$ , просто невозможны.

Вообще говоря, не все комбинации объектов из универсума  $U$  представляют допустимые множества. Для разных предметных областей возможны специфические ограничения. Если на универсум не наложено никаких связей, каждому единичному объекту  $x$  соответствует одноэлементное множество  $\{x\}$ . Для конечного универсума  $U$  можно было бы говорить о множестве, содержащем все вообще объекты  $x$  из  $U$ ; однако такого множества может и не быть. Например, когда объекты не постоянно присутствуют в  $U$ , — или какие-то из них не могут принадлежать одному множеству (обладают несовместимыми свойствами). Мы предполагаем также, что всякое множество содержит хотя бы один элемент; обычное представление о пустом множестве  $\emptyset = \{\}$  в предметной теории возможно лишь метафорически, в качестве сокращения для фразы «не существует множества, такого, что...» Множество, по определению, есть то, у чего имеются элементы. Когда элементов нет — это уже не множество, а нечто иной природы, что надо изучать отдельно.

Если одно множество равно другому ( $A = B$ ) — это просто одно и то же, по-разному обозначенное. Разные способы построения могут приводить к одному набору элементов. Равенство множеств — это предметное равенство, одинаковость свойств.

Одно множество может быть (собственным) подмножеством другого:  $A \subset B$ . Мы говорим о сужении множества  $B$ , выделении части  $A$ . Заметим, что речь всегда идет о конкретных объектах из универсума, которые присутствуют в обоих множествах, либо только в одном из них. Независимо от количества элементов (мощности), множество  $B$  будет шире множества  $A$ , если в  $B$  есть элементы, отсутствующие в  $A$ .

На уровне множеств можно обычным образом определить операции объединения  $A \cup B$  и пересечения  $A \cap B$ , а также дополнение одного множества до другого  $B \setminus A$ . Однако формальное построение ничего не говорит о существовании соответствующего множества. Например, объединение двух множеств может отсутствовать при наличии связей, несовместимости свойств элементов. Пересечение множеств не может быть пустым; в противном случае мы говорим, что множества не пересекаются. Точно так же, дополнение существует только для собственного подмножества:  $A \subset B$ . Другими словами, набор реализуемых в теории множеств операций определяется особенностями предмета; и наоборот, из существования теоретико-множественных структур мы делаем выводы о строении универсума.

Совокупность реализуемых предметных множеств может стать универсумом для следующего уровня иерархии, на котором мы строим множества из множеств. По отношению к предмету, это уже не множества, а классы — сущности существенно иного уровня, которые не могут включать объекты из универсума, и сопоставлять одно с другим напрямую нельзя.

Если построение множеств связано преимущественно с предметными свойствами, классы в большей степени определяют логику теории, способ рассмотрения. Разные теории одного и того же будут различаться на уровне классов.

В простейшем случае, когда возможны любые комбинации свойств, уровень классов обладает характерной структурой. Для каждого одноэлементного множества, можно построить класс всех пересекающихся с ним множеств; другими словами, мы рассматриваем все совместимые с выделенным объектом из универсума свойства как класс, совокупность множеств, содержащих соответствующий элемент. Поскольку всякое множество принадлежит классу каждого из своих элементов, уровень классов полностью покрывается такими, элементарными классами. Таким образом, классы однозначно сопоставлены объектам универсума, и уровни иерархии переходят один в другой. Это называется обращением иерархии: элемент принадлежит множеству — но и множество принадлежит элементарному классу. Такая трехуровневая схема и представляет собой формальную теорию, математическую модель предметной области.

Еще один способ развертывания иерархии опирается на возможность сделать произвольное множество универсумом (базой) для множеств следующего уровня. Такие множества можно было бы назвать «внутренними» — в противоположность классам как «внешним» множествам. Внутренние множества, очевидно, соответствуют подмножествам базы; однако это не одно и то же: они принадлежат разным уровням. На этом примере можно легко видеть, что набор внутренних множеств ограничивается реально имеющимися подмножествами базы; это один из способов формализации понятия связи. Разумеется, существуют также классы внутренних множеств. При достаточно богатом универсуме теория допускает весьма сложные иерархические структуры.

Пусть теперь у нас есть два различных универсума  $U_1$  и  $U_2$ . Каждый из них порождает свою теоретико-множественную иерархию. Нельзя произвольно комбинировать компоненты разных иерархий, даже одного уровня. Это как если бы мы вдруг начали складывать миллиметры с килограммами. Требуется объединить два универсума в один — и только тогда можно строить единую теорию. А объединять можно по-разному. Две основные парадигмы: последовательное и параллельное соединение, временное и пространственное, внутреннее и внешнее.

При последовательном синтезе мы получаем универсум, который называется (прямым, или декартовым) произведением:  $U = U_1 \times U_2$ . Объектами такого универсума будут упорядоченные пары объектов из  $U_1$  и  $U_2$ :  $x = \langle x_1, x_2 \rangle$ . Здесь существенно, какой из исходных универсумов идет первым, а какой после него. Даже если эти два универсума совпадают, в прямом произведении они выступают в разном качестве: как первый и второй. В каждой из этих позиций могут быть свои связи, а также ориентированные связи на возможные сочетания объектов в паре. «Декартов квадрат»  $U^2$  отличается от числового возведения в степень; в частности, в паре  $\langle x, x \rangle$  объект  $x$  в первой позиции — не то же самое, что объект  $x$  во второй: один и тот же объект здесь рассматривается с разных сторон. Например, он может быть взят в разные моменты времени — и существенно измениться от одного к другому. Формально, мы вполне можем взять за основу противоположный порядок. Однако в этом случае и выводы теории будут звучать иначе: например, если при одном порядке  $x_1$  «предшествует»  $x_2$ , — противоположный порядок предполагает другую терминологию: мы говорим, что  $x_2$  «следует за»  $x_1$ . Синтез универсумов объективен; он предполагает практическое использование объединенного универсума, умение реально различать место в последовательности.

Множества при таком синтезе оказываются множествами упорядоченных пар, с учетом наложенных связей. Каждый элемент множества содержит компоненты от каждого из исходных универсумов. Легко видеть, что, при наличии множества всех элементов для начального и конечного универсума по отдельности, множества над объединенным универсумом возможно трактовать как подмножества обычного декартова произведения множеств. Однако мы знаем, что такие «глобальные» множества могут не существовать, а декартовы произведения существующих по отдельности множеств не обязательно порождают те же самые множества пар, особенно при наличии ориентированных связей.

Последовательный синтез соответствует выделению у объекта разных сторон, с которыми мы на практике можем иметь дело по отдельности. Например, температура и давление газа; ширина реки и ее глубина; или общество в разные эпохи. Параллельный синтез, напротив, соединяет универсумы внешним образом, как две «параллельных» реальности (например, как разные области одного пространства, или компоненты и фазы многочастичной физической системы). В этом случае мы говорим о (прямой) сумме, или суперпозиции:  $U = U_1 + U_2$ . Такая сумма коммутативна: нам важен сам факт присутствия объектов, а не их порядок. В достаточно развитых иерархиях возможны также «взвешенные» суперпозиции, когда элементы разных универсумов учитывают с разными коэффициентами. Объекты разных универсумов сосуществуют одновременно, мы можем использовать в работе и те, и другие. Тогда каждое множество  $A$  над объединенным универсумом будет прямой суммой множеств, построенных над каждым универсумом по отдельности:  $A = A_1 + A_2$ . От объединения множеств это отличается тем,

что элементы из разных универсумов не смешиваются, они участвуют в теоретико-множественных построениях по отдельности. В частности, нельзя говорить о полном количестве элементов множества: есть количество элементов одного и другого универсума, два числа вместо одного. Не всякие комбинации (прямые суммы) множеств, надстроенных над  $U_1$  и  $U_2$  по отдельности, будут множествами объединенного универсума: существование результата следует устанавливать особо, исходя из имеющихся связей.

Последовательное объединение универсумов по отношению к множествам является внутренним: каждый элемент множества расщепляется на две компоненты. Напротив, параллельное объединение связывает множества внешним образом, не меняя их элементы.

Если  $A = A_1 + A_2$  и  $B = B_1 + B_2$ , объединение множеств происходит покомпонентно:  $A \cup B = A_1 \cup B_1 + A_2 \cup B_2$ . Поскольку же в предметной теории не бывает пустых множеств, в каждом множестве над составным универсумом должны присутствовать обе компоненты. Это вполне аналогично тому, как декартово произведение содержит только пары элементов: позиции в кортеже не могут быть пустыми. Поэтому пересечение  $A \cap B = A_1 \cap B_1 + A_2 \cap B_2$  существует только при одновременном наличии общих элементов у  $A_1$  и  $B_1$ , и у  $A_2$  и  $B_2$ . Заметим, что ни то, ни другое пересечение не обязано существовать по отдельности как множество над  $U_1$  и  $U_2$  соответственно; однако в качестве компонент составного универсума они оказываются возможными.

В общем случае универсум  $U$  может быть многокомпонентным, с разными типам соединения компонент. В этом случае особенно важно дополнить формальные построения анализом существующих ограничений. Корректность рассуждений в такой теории связана не только с ее логикой, но и со строением предмета. Например, допустим, что  $S$  — универсум симптомов,  $C$  — индивидуальные особенности пациента,  $M$  — универсум схем терапии; тогда множества над универсумом  $(S + C) \times M$  представляют возможные методики лечения болезней. Понятно, что далеко не всякие комбинации имеют смысл. В качестве связи здесь можем выступать динамика заболевания, которую нам важно направить к выздоровлению, а не наоборот. Процесс лечения представляется в этой схеме последовательностью множеств, когда методы лечения регулярно подстраиваются под изменения в состоянии больного.

Для множеств составного универсума обычным образом определяются внутренние множества и классы (внешние множества). Помимо элементарных классов, здесь появляются также разного рода проекции: классы множеств, совпадающих по одной из компонент. Например, можно рассматривать параметрическое семейство  $A = A_1 + *_2$ , понимая под  $*$  произвольное множество над универсумом  $U_2$ . Класс существующих над  $U$  сумм такого вида (проекция на  $A_1$ ) определяет некоторый класс над  $U_2$ , который можно было бы назвать смежным классу  $A$ ; вообще говоря, он может не существовать в иерархии над  $U_2$  сам по себе. Пересечение этого класса с собственными классами  $U_2$  образует границу семейства  $A$  в  $U_2$ . Точно так же определяются классы декартовых проекций — семейства множеств с элементами вида  $\langle x_1, * \rangle$  или  $\langle *, x_2 \rangle$ . Как и в случае элементарных классов, иерархия здесь замыкается на себя, поскольку полученные структуры оказываются подобны каким-то из уже имеющихся в теории объектов.

Разумеется, можно рассматривать и другие структуры на уровне множеств (логические структуры), и порождаемые ими «классы эквивалентности». Однако сами по себе формальные конструкции не становятся компонентами теории: для этого нужно еще и предметное наполнение. Сопоставимость классов с компонентами универсума как раз и дает критерий приемлемости формальных процедур, ограничивает область их применимости.

Итак, всякая предметная теория множеств: 1) опирается на некоторый, не обязательно формально определенный универсум; 2) надстраивает над универсумом уровень множеств; рассматривает классы множеств; 3) идентифицирует классы с объектами универсума. Строение предмета определяет структуру теории.

август 2006



### Кардинальная иерархичность

Всем известна канторовская иерархия бесконечных кардинальных чисел: следующий уровень возникает при рассмотрении множества всех подмножеств бесконечного множества. Можно пронумеровать эти бесконечности целыми числами (ординалами): дискретной бесконечности соответствует уровень 0, континуум — бесконечность уровня 1, всевозможные функции над континуумом — уровень 2, и т. д. Разумеется, мыслим и выход за пределы этой иерархии — как некие бесконечные ординалы, включая несчетные.

В этих терминах, знаменитая континуум-гипотеза утверждает, что не существует уровня мощности между нулем и единицей. Утверждение сильное, и в рамках аксиоматической теории множеств его нельзя ни доказать, ни опровергнуть. Можно показать, что такая существенная дискретность связана с двужначностью логики (способом конструирования подмножеств). Однако никто не запрещает нам несколько расширить идею множества, сохраняя традиционные представления как предельный случай. Существует обобщение, позволяющее конструктивно показать, что континуум-гипотеза в обобщенной теории неверна.

Действительно, обратим внимание, что в традиционной теории элементы объединяются в множества внешним образом, как противопоставленные друг другу. Каждый элемент соответствует одной и той же счетной единице, которая для бесконечных множеств бесконечно мала — но все же сохраняет качественную определенность, общую для всех элементов (именно поэтому мы имеем право их пересчитывать). Поэтому канторовская иерархия — это ряд внешних бесконечностей.

Допустим теперь, что элементы множеств — не просто счетные единицы, а каждый из них обладает внутренним строением. Например, каждой точке (элементу множества) соответствует некоторое внутреннее пространство, со своей внешней мощностью. Вообще говоря, устройство внутренних пространств у разных элементов не совпадает. Однако, если мы хотим изучать достаточно целостные объекты, можно полагать, что принцип разворачивания внутреннего пространства у всех точек одинаков — и их внутренние пространства, как минимум, одинаковой мощности; во многих практически важных случаях (например, при описании механического движения) допустимо потребовать и одинаковости топологии.

В простейшем случае внутреннее пространство дискретно (нечто вроде спинорных компонент). Это уровень 0 в иерархии внешних мощностей. Однако точно так же каждая точка может внутренне представляться некоторой непрерывной областью — зоной, множеством уровня 1. На практике такие ситуации не редкость. Например восприятие чистого тона некоторой высоты (логарифм частоты звука) представляет его как некоторое распределение частот в окрестности хорошо выраженного максимума. Отсюда следует, что возможные наборы воспринимаемых как различные тонов образуют зонные структуры (звукоряды, лады, созвучия), каждый элемент которых вовсе не точка, а континуум возможных отклонений от центра зоны.<sup>1</sup> Как определять мощность такого множества? С одной стороны, оно дискретно, а с другой — это набор континуальных интервалов.

Определим теперь мощность двухуровневого множества (со стандартным внутренним пространством каждого элемента) как пару кардинальных чисел  $(K_1, K_2)$ , где  $K_1$  и  $K_2$  задают уровни мощности верхнего (внешнего) и нижнего (внутреннего) соответственно. В частности, внутреннее пространство может отсутствовать — в этом случае оно эффективно состоит из одного элемента, и его уровень мощности равен нулю. Таким образом, чисто дискретное множество представляется иерархической мощностью  $(0, 0)$ , обычный континуум имеет мощность  $(1, 0)$ , а дискретная структура с континуальным внутренним пространством имеет мощность  $(0, 1)$ . Порядок на множестве иерархических мощностей логично установить

<sup>1</sup> L. V. Avdeev and P. B. Ivanov, "A Mathematical Model of Scale Perception", *Journal of Moscow Physical Society*, 3, 331–353 (1993).

лексикографически: сначала по первому показателю, потом по второму. Очевидно,  $(0, 0) < (1, 0)$ . Однако столь же естественно  $(0, 0) < (0, 1) < (1, 0)$ . Таким образом, существуют иерархические множества с мощностью между дискретным и континуумом.

Разумеется, процесс можно продолжать и дальше: элементы нижнего уровня также могут быть внутренне сложными. Ограничиваясь пока первыми двумя канторовскими уровнями, получим кардинальное число иерархического множества как  $(b_1, b_2, b_3, \dots)$ , где  $b_k$  равны либо 0, либо 1. Легко видеть, что это эквивалентно двоичной записи некоторого вещественного числа в интервале  $(0, 1)$ , — и существует бесконечно много множеств промежуточной мощности.

Конечно, строение внутренней иерархии может быть достаточно сложным. Например, в нем воспроизводится обычная иерархия бесконечных ординалов. Особый интерес представляют ситуации, когда пространства разных уровней изоморфны: иерархия в этом случае уже не останется простой древовидной структурой, а может содержать циклы и петли. Вычисление мощности таких множеств — особый интересный вопрос.

июль 1994

### Качество отрицания

Как учит философия, главное назначение человека разумного — связывать мир воедино, соединять то, что без нас никаким образом соединиться не может. Поэтому обычные в математике сопоставления одного с другим — часть разумной работы. По жизни часто оказывается, что навыки из одной отрасли производства прекрасно работают в другой, — математический язык для фиксации таких фундаментальных схем бывает очень удобен и потому безусловно важен. Однако трезво оценивать имеющиеся возможности тоже небесполезно. Поскольку мы лишь часть мира — сколь угодно полное представление его в наших абстракциях остается всего лишь частичным, и ни одна из формальных конструкций не вправе претендовать на безоговорочную универсальность. В частности, если удалось связать пару-тройку математических объектов, вплоть до полной практической неразличимости, это никоим образом не означает, что они не пожелают в какой-то ситуации вспомнить о своих различиях, и потребуют другой математики, — при всем уважении к традициям.

С другой стороны, разнообразие частных возникает на фоне исконной уникальности мира: никаких других миров нет и (следовательно) быть не может. В человеческой практике (включая математическую) это означает скрытое присутствие зачатков будущего в прошлом и настоящем: если мы до чего-то додумались, оно существовало и в прежних представлениях — но до поры не привлекало внимания. То есть, типовая технология научного открытия — взять что-нибудь до боли знакомое, о чем приличные люди в обществе не разговаривают, — и указать маленькие тонкости, которые при случае станут исполнены великого смысла. Кому лень оформлять такие открытия — могут ограничиться предварительными замечаниями, полностью в русле общеизвестного — но с философскими экскурсами.

Вот и давайте возьмем нечто совершенно заурядное и посмотрим, нельзя ли углядеть в нем недооцененный творческий потенциал.

Если еще не рожденному младенцу в утробе матери задать вопрос: *сколько будет два минус три?* — он не задумываясь ответит: *конечно же, минус один!* Ленивый философ тут же подбрасывает вопрос на засыпку: *а что такое минус один?* Младенец чешет пуповиной в затылке и мычит: *ну, это число такое... — отрицательное, — вроде единицы, только с минусом.* Тут мы догадываемся, что беседуем с будущим математиком: люди обыкновенные, как правило, лишены подобной кристальности понимания. Многие сочли бы задачку на вычитание некорректно поставленной: нельзя из меньшего вычесть большее! Другие согласились бы со школьным ответом (минус один) — но с добавлением, что таких чисел на самом деле нет, это чистая условность, в расчете на то, чтобы когда-нибудь пристроиться к достаточно положительному

числу — и уменьшить его в разумных пределах. В естественных языках отрицательность прочно ассоциируется с дурным и неправильным, с болезненным извращением, — чего в норме быть не должно. Наконец, есть древняя профессия, представители которой глубоко убеждены: никаких отрицательных чисел нет, а есть реальные количества по разным счетам, каждый из которых представлен парой: *у нас есть* — *нам нужно*; или, иначе: *мы должны* — *нам должны*. Соответственно, различаются активные и пассивные счета, оценка состояния — или долгов.

Глупо выяснять, кто из них ближе к истине: правы все — но каждый в пределах своего опыта. И каждому своя математика. В донаучную эпоху признавали, что любая вещь обладает букетом разных качеств — а в пределах одного качества можно говорить о количественных различиях. То есть, сначала: есть или нет, — а уж если есть, тогда сколько. Отсутствие мы обозначаем нулем; для наличного выбираем единицы измерения — и все сводим к принятой шкале; по каждому измерению, число **1** ссылается на выбранный масштаб. Античному геометру заранее ясно: у отрезка есть длина, и это не зависит от того, как и в каком направлении мы его откладываем. Точно так же, есть площади фигур и объемы тел; когда чего-то измеримого нет — понятия длины, площади или объема просто неуместны, и говорить о них — логическая ошибка.

Однако люди живут не одним мгновением: у них есть прошлое и будущее. О прошлом — память, про будущее — планы. И можно подойти к вопросу по-человечески эмоционально, сожалея, что былой славы уже нет, а чего-то очень желательного — пока нет. Эту грусть мы и обозначаем в математике знаком — (минус).

Важный момент: по сравнению с простой констатацией отсутствия, отрицательность — серьезный шаг вперед. Не ноль, а таки наличие — но идеально: в субъекте, или в других телах (как след или возможность). Понятно, что такое бытие несколько отличается от настоящего, и смешивать их в одну кучу, вообще говоря, неправильно. Единицы измерения у них разные. Однако есть и общее: сам факт качественной определенности, и потому — измеримости. Значит, можно в каком-то смысле (то есть, в контексте определенной деятельности) уподоблять положительное отрицательному и рассматривать их как явления однопорядковые. Разумеется, если не зарпортоваться и не отринуть напрочь всякие различия.

Чтобы держать себя в рамках и не допускать излишне формального отношения к делу — есть проверенные формальные приемы. Один из них — то самое разведение дебета и кредита, подчеркнутое различие того, что по жизни различно. Дебет и кредит — противоположности, одно отрицает другое; но это не количественное отрицание, а противопоставление разных качеств. Где-то в другом месте, для высшего руководства, — оно без разницы, и можно одну большую кучу вычестить из другой. Но это вид издалека — а в живой бухгалтерии надо честно записывать в два столбика. Еще раз: это не дикий примитивизм, не замшелая традиция, а точное отражение реалий бытия. У каждого есть отец и мать — это плюс; кого-то поднимала мать-одиночка — и одного родителя не хватает, он с минусом, — но это вовсе не то же самое, что фактически иметь только одного родителя (например, при партеногенезе). То есть,  $2-1$  качественно отличается от простой единицы. А школьной математике все едино...

Точно так же, при подсчете ресурсов: в плюсе у нас все, что когда-либо поступало, — это факт настоящего времени. Однако часть из этого уже израсходована и присутствует только в памяти — это минус прошедшего времени. С другой стороны наши потребности — минус будущего времени, который может существенно перевешивать наличное изобилие. Заметим, что положительные числа характеризуют объективное положение вещей — тогда как отрицательные больше о субъективных переживаниях. По крайней мере, пока мы ограничиваемся одним уровнем иерархии. В этом смысле правы те, кто считает отрицательные числа всего лишь условностью.

Чтобы учесть условности, математик может представить «обобщенное» число  $s$  кортежем из двух компонент  $\langle a, b \rangle$ , где  $a$  и  $b$  — неотрицательные вещественные числа (которые мы предположительно определить умеем). Компоненты  $a$  и  $b$  мы называем, соответственно, положительной и отрицательной частью числа  $s$  — а его мы, для определенности, будем

именовать *аддитивно расщепленным вещественным числом* (или просто *расщеплением*, ради экономии места и читательского внимания).

Было в математической практике нечто подобное? Да сколько угодно! Например, рациональные числа мы определяем как пару целых — с правилами комбинирования компонент, подобранными так, чтобы это напоминало обычную арифметику. Точно так же, комплексные числа мы определяем через действительную и мнимую части — опять же, с подходящей детализацией арифметики. Так почему с отрицательностью не пойти проторенным путем?

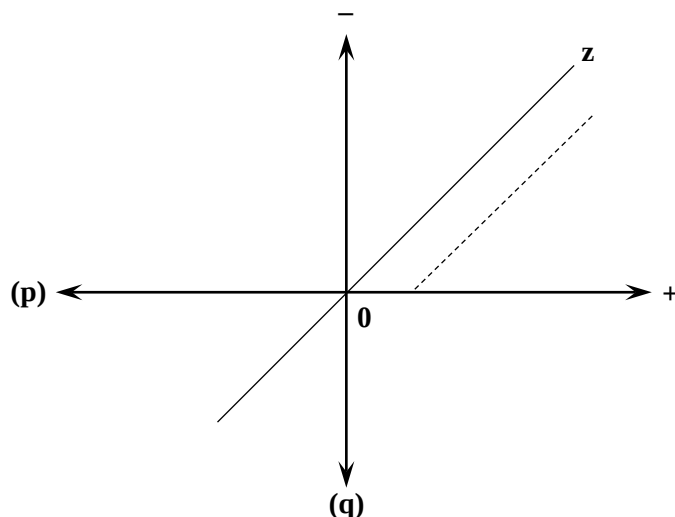
Если трактовать компоненты кортежа как координаты точки на плоскости, базисные вектора этого пространства называются положительной и отрицательной единицей, которые естественно обозначить символами  $(+1)$  и  $(-1)$ ; формально,  $(-1)$  — нерасчленимый знак для некоторой единицы измерения, вообще говоря, никак не связанной с единицей измерения положительной части. Перевод одних единиц в другие требует соответствующего размерного множителя.

По аналогии с комплексными числами (и линейной алгеброй), допустимо ввести условно-алгебраическую нотацию  $s = a + (-1)b$ . Базисный вектор  $(+1)$  здесь по традиции опущен (но всегда подразумевается). Минимальный набор полезных свойств вещественных расщеплений можно проиллюстрировать приведенной ниже табличкой, где также показаны аналогичные правила для комплексных чисел.

| <i>расщепления</i>                                     | <i>комплексные числа</i>                                   |
|--------------------------------------------------------|------------------------------------------------------------|
| $(-1)^2 = 1$                                           | $i^2 = (-1)$                                               |
| $1 / (-1) = (-1)$                                      | $1 / i = (-1)i$                                            |
| $s = a + (-1)b$                                        | $c = a + ib$                                               |
| $a = \text{Pos } s$                                    | $a = \text{Re } c$                                         |
| $b = \text{Neg } s$                                    | $b = \text{Re } c$                                         |
| $s = \text{Pos } s + (-1) \text{Neg } s$               | $c = \text{Re } c + i \text{Im } c$                        |
| $s_1 + s_2 = (a_1 + a_2) + m(b_1 + b_2)$               | $c_1 + c_2 = (a_1 + a_2) + i(b_1 + b_2)$                   |
| $s_1 s_2 = (a_1 a_2 + b_1 b_2) + m(a_1 b_2 + b_1 a_2)$ | $c_1 c_2 = (a_1 a_2 + (-1)b_1 b_2) + i(a_1 b_2 + b_1 a_2)$ |
| $\text{Pos } (-1)s = \text{Neg } s$                    | $\text{Re } ic = (-1) \text{Im } c$                        |
| $\text{Neg } (-1)s = \text{Pos } s$                    | $\text{Im } ic = \text{Re } c$                             |
| $(-1)s = b + (-1)a$                                    | $ic = (-1)b + ia$                                          |

Разумеется, с формальной точки зрения, эти равенства не будут независимыми — но нас интересует суть дела, а не формальности. Что из чего выводить — зависит от принятой стратегии арифметизации; например, можно исходить из алгебры операций над расщеплениями и обойтись без отсылок к их внутреннему строению (представление кортежами): положительная и отрицательная части расщепления появятся тогда в теории как функционалы (отображения из пространства расщеплений в пространство вещественных чисел), или как отображения одних расщеплений в другие (проекторы). Традиционно, мы идем от вещественных чисел к комплексным; при этом отрицательная единица  $(-1)$  естественно возникает в равенствах комплексной арифметики. Однако ничто не мешает, наоборот, *определить*  $(-1)$  как  $i^2$  — и тогда арифметику расщеплений придется выводить из комплексной плоскости.

Поскольку на пространстве неотрицательных вещественных чисел уже задан порядок, можно разделить расщепления на два класса: при  $a > b$  ( $\text{Pos } s > \text{Neg } s$ ), расщепление  $s$  называется положительным; при  $a < b$  ( $\text{Pos } s < \text{Neg } s$ ), расщепление  $s$  называется отрицательным. Это определение никоим образом не предполагает прямого комбинирования положительной и отрицательной частей; геометрически, речь идет о разбиении пространства расщеплений (первый квадрант на обычной евклидовой плоскости; на рисунке — область  $(+|\mathbf{0}| -)$ ) на два подпространства (ниже или выше главной диагонали  $(\mathbf{0}|\mathbf{z})$ ). Разумеется, возможны любые другие подразделения, в зависимости от практических нужд.



В алгебраической нотации отрицательная единица  $(-1)$  приобретает смысл оператора, (который мы называем *отрицанием*), однозначно переводящего одни расщепления в другие по определенному правилу: положительная и отрицательная части расщепления меняются местами; графически, это отражение относительно главной диагонали. В частности, точки оси  $(0|+)$  этот оператор переводит в точки оси  $(0|-)$ . Точно так же, положительная единица как оператор просто переводит каждую точку в ту же самую. В компонентном представлении, положительная единица представлена кортежем  $\langle 1, 0 \rangle$ , а отрицательная единица — кортежем  $\langle 0, 1 \rangle$ . Однако единица как оператор — не то же самое, что единичный вектор; отрицание как действие и как результат — вещи разные.

Переход к полярной системе координат от компонентного представления формально одинаков для расщеплений и комплексных чисел:

$$s = r(\cos \varphi + (-1)\sin \varphi)$$

Такая запись позволяет обсуждать тонкие особенности изучаемых структур — однако оправдана она лишь там, где это отвечает набору наложенных на систему симметрий и выбору метрики. А традиционные симметрии комплексной плоскости сильно отличаются от обычной симметрии расщеплений. Это ясно уже из того факта, что определены расщепления только в одном квадранте плоскости, и для расширения на всю плоскость потребовалось бы (как показано на рисунке) ввести еще две фундаментальные единицы (например,  $p$  и  $q$ ) и расщепление общего вида записывать как  $a + (-1)b + px + qy$ , — нечто вроде кватерниона. Дополнительные симметрии «склеивают» пространство специфическим образом, отождествляя его точки. В случае обычных аддитивных расщеплений предполагается фундаментальная симметрия

$$\begin{aligned} \langle a, b \rangle + \langle z, z \rangle &= \langle a + z, b + z \rangle = \langle a, b \rangle \\ (a + z) + (-1)(b + z) &= a + (-1)b \\ s + (-1)s &= 0 \end{aligned}$$

где  $z$  — положительное вещественное число. Графически, это отождествляет все точки главной диагонали — и любая прямая, параллельная главной диагонали, формально склеивается в одну точку (задает класс эквивалентности). Аналог из комплексной математики

$$c + i^2 c = 0$$

квадратичен по мнимой единице — что логично порождает квадратичную метрику. По аналогии с ортогональностью векторов на плоскости, когда их скалярное произведение равно нулю, можно трактовать равенство  $1 + (-1) = 0$  как своего рода аддитивную ортогональность положительной и отрицательной единиц.

Симметрии (наложенные на систему связи), вообще говоря, уменьшают количество степеней свободы, меняют эффективную размерность пространства. В случае расщеплений мы формально сводим двумерное пространство к одномерному (или даже четыре измерения склеиваются в одно, как для «кватернионных» расщеплений) — это существенно нелинейная операция, нечто вроде проекции. Существуют и другие проекции: например, выделение положительной и отрицательной (вещественной или мнимой) частей расщепления (комплексного числа), вычисление нормы (квадратичной — или по линейной комбинации компонент), вычисление фазы, и т. д. На практике мы всегда видим систему в одной из проекций, и существование других — надо восстанавливать умозрительно. В релятивистской физике аналогичную роль играет калибровка. Необходимость проекций связана с непосредственно не наблюдаемым внутренним движением. Однако симметрии этого «виртуального» пространства определяются тем, что мы реально умеем делать с системой и как регистрируем результат — что считаем готовым продуктом.

Эффективная одномерность пространства расщеплений с глобальной симметрией не должна вводить нас в заблуждение: эквивалентность не есть тождество. Совокупность прямых, параллельных главной диагонали — это нечто иное, чем совокупность точек на прямой, а сдвиги прямых вверх или вниз от главной диагонали (области с  $a < b$  и  $a > b$ ) все еще качественно различны.

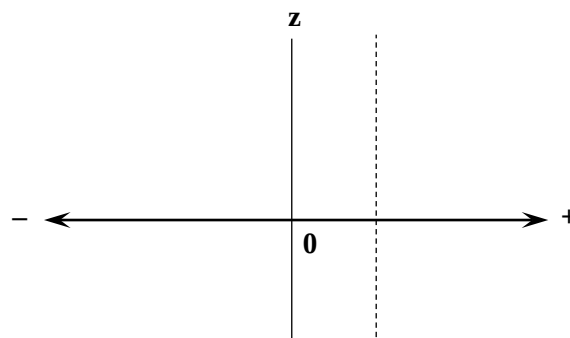
В условиях глобальной симметрии, для любого расщепления  $s$  существует единственное эквивалентное ему «редуцированное» расщепление  $R(s)$ , такое, что либо  $\text{Neg } R(s) = 0$ , либо  $\text{Pos } R(s) = 0$  (положительная и отрицательная редукции). Можно ввести линейный порядок на пространстве расщеплений, если положить, что все положительные расщепления больше отрицательных; при этом положительные расщепления упорядочены по возрастанию  $\text{Pos } R(s)$ , а отрицательные расщепления — по убыванию  $\text{Neg } R(s)$ :

$$\text{Neg } R(s) = 0 \ \& \ \text{Pos } R(s') = 0 \Rightarrow s > s'$$

$$\text{Pos } R(s) > \text{Pos } R(s') \Rightarrow s > s'$$

$$\text{Neg } R(s) > \text{Neg } R(s') \Rightarrow s < s'$$

Следовательно, мы можем формально «продолжить» вещественную полупрямую в область отрицательных значений, эффективно переходя от двух независимых измерений к одному; графически это выглядит как развертывание первого квадранта в верхнюю полуплоскость, а редукция представляется проекцией на горизонталь ( $-|+$ ):



Тем не менее, положительная и отрицательная ветви все еще остаются разными пространствами: да, одно однозначно отображается в другое с сохранением алгебраической структуры внутри каждого из пространств, — но заранее не очевидно, как определить алгебраические операции с элементами разных пространств, что позволило бы трактовать эти элементы единообразно. Легко видеть, что

$$R(s_1 + s_2) = R(R(s_1) + R(s_2))$$

$$R(s_1 s_2) = R(R(s_1) R(s_2)) = R(s_1) R(s_2)$$

Умножение остается умножением и в редукциях — однако сложение отрицательных редукций с положительными выводит нас в нередуцированную область — и приходится прибегнуть к еще одной проекции. Таким образом, построение единого пространства вещественных чисел требует достаточно сильных допущений — которые не всегда оправданы приложениями.

Слишком прямолинейная трактовка условно-алгебраической нотации расщеплений (или комплексных чисел) может легко ввести в заблуждение: качественно разные объекты трактуют на единых основаниях. Но запись  $s = a + (-1)b$  — это лишь другая графика для  $s = \langle a, b \rangle$ ; здесь символы  $a$  и  $b$  ссылаются на вещественные числа,  $(-1)$  указывает на позицию в кортеже, а знак суммы говорит о совместном рассмотрении компонент. Чтобы превратить это в собственно алгебраическое выражение, следовало бы писать нечто вроде

$$\langle a, 0 \rangle + \langle 0, 1 \rangle \cdot \langle b, 0 \rangle = \langle a, b \rangle$$

В частности, в правой части связи  $s + (-1)s = 0$ , понимаемой алгебраически, подразумевается вовсе не вещественный нуль, а кортеж  $\langle 0, 0 \rangle$ , так что

$$\langle k, k \rangle = \langle 0, 0 \rangle \neq 0$$

Расположение положительной и отрицательной полуосей на одной прямой ничего в этом плане не меняет, а лишь скрадывает (аддитивную) ортогональность. Совершенно точно так же, на пространстве рациональных чисел (представленных кортежами из двух целых чисел) действует симметрия

$$\langle k, k \rangle = \langle 1, 1 \rangle$$

поскольку мы можем сокращать числитель и знаменатель на общий множитель. Нечто подобное возникает в полярных координатах для комплексных чисел:  $\langle x, y \rangle \rightarrow \langle r, \varphi \rangle$ . При отсутствии ветвления, действует симметрия

$$\begin{aligned} \langle r, \varphi \rangle &= \langle r, \varphi + 2\pi k \rangle \\ \langle 0, \varepsilon \rangle &= \langle 0, 0 \rangle \neq 0 \end{aligned}$$

для любого целого  $k$  и любого угла  $\varepsilon$ . Но в общем случае для каких-то приложений эти симметрии могут не соблюдаться: сдвиг фазы на достаточно большие углы выводит на другие листы, а приближение к нулю с разных направлений может давать разные пределы.

Сухой остаток: во всякой теории нуль (как отметка на шкале) и ноль (как количество) — не отсутствие предмета как такового, а отсутствие определенного качества. То есть, нет нуля вообще — а есть нуль по отношению к некоторой единице, качественно определенный нуль. Именно поэтому мы можем (в каких-то допущениях) доопределить операции с объектами одного типа, включая в элементную базу того же типа нули. Однако в строгой математике, ноль вовсе не число (единичный объект); это, скорее, способ порождения объектов: шаблон, форма, схема, образец, тип, — то есть, класс объектов, характеристика самой возможности их построения. Симметрия  $\langle k, k \rangle = \langle 0, 0 \rangle$  может быть интерпретирована не только как эквивалентность точек диагонали, сведение их к нулевому кортежу, — но и в обратном смысле, как порождение виртуальных пар из нуля, который тогда оказывается подобен физическому вакууму (а это вовсе не пустота!). Если мы выдрали из вакуума частицу — в нем возникает дырка того же типа и размера. Как и в физике, нулевой уровень часто оказывается подвижным — но такая симметрия не меняет поведения системы в отношении прочих симметрий.

Точно так же, бесконечность — не число (хотя бы и кардинальное, или ординал), а та самая деятельность, в которой мы различаем противоположности — и переводим их друг в друга. Ноль связан с бесконечностью, как возможность и действительность; это одно и того же с разных точек зрения. Одно дело — компьютерная программа, другое — работа по программе; однако языки программирования исходят из возможностей железа, а железо понемногу адаптируется под идиоматику языков.

Заметим, что особая роль нуля связана с логикой теории: например, мы можем говорить, что некоторое суждение является истинным (+1) или ложным (−1) — но и то, и другое утверждает его логическую определенность, возможность проверки, — тогда как нулевая оценка означает, скорее, некорректность постановки вопроса, неприменимость категорий истинности или ложности в данном контексте (хотя в других отношениях то же суждение вполне может претендовать на положительную или отрицательную оценку: например, быть правильным или неправильным).

В теории расщеплений нулевые компоненты означают просто отсутствие каких-либо операций по «дебету» или «кредиту»; с учетом «накладных расходов», кортеж  $\langle 0, 0 \rangle$  — совсем не то же самое, что виртуальный обмен  $\langle k, k \rangle$  (сколько зачислили — столько же списать). Заметим, что операция обмена может быть нетривиальной, если единицы измерения положительной и отрицательной компонент различны: например, в случае обмена валюты. Однако в классическом учете обмены представляют не расщеплениями  $\langle a, b \rangle$ , а парой редукций  $\langle a, 0 \rangle$  и  $\langle 0, b \rangle$ , отнесенных к разным счетам. Это различие подобно тому, как в физике мы различаем квантовые переходы и каскады; в виртуальных мирах возможна интерференция процессов, приводящая к вполне наблюдаемым («макроскопическим») эффектам, — но точно так же устроена и финансовая спекуляция! На более высоком уровне иерархии, в сводных отчетах, оказывается возможным формальное сложение движений по разным счетам:

$$\langle a, 0 \rangle + \langle 0, b \rangle = \langle a, b \rangle$$

Следующий уровень — предполагает положительную или отрицательную редукцию (или нулевой баланс).

Глубокий математический (и философский) смысл расщеплений состоит в том, что один и тот же объект возможно получать различными способами, которые в каких-то отношениях будут эквивалентными — но в другом контексте возникают серьезные различия. Здесь мы говорили об аддитивных расщеплениях вещественных чисел: каждое число виртуально представляется разностью двух других. Точно так же устроены мультипликативные расщепления, когда вещественное число представлено отношением двух чисел; легко видеть, что это расщепление становится аддитивным при логарифмировании. Можно обсуждать расщепления в суммы и произведения (в том числе бесконечные). Пример — разложение чисел на простые множители, ядро и смысл классической теории чисел. Еще примеры: вращение в разных направлениях, инверсия относительно окружности, различение внешней и спинорной размерности, расслоения многообразий... Наконец, из той же серии — противоположность объекта и субъекта в контексте определенной деятельности. Поскольку же любая деятельность регулярно воспроизводит свой продукт, расщепления можно понимать как циклы, замкнутые пути; соответственно, возможны многокомпонентные и многоуровневые расщепления. Любой математический объект в общем случае представляется иерархией всевозможных расщеплений — способов порождения определенности. Развертывание этой иерархии в одну из допустимых иерархических структур называется математической теорией. Способ развертывания — их практики, от насущных потребностей. Математика (как и любая наука) нужна людям для того, чтобы вооружать их формальными приемами, типовыми схемами деятельности, освобождая разум от рутинных операций ради поиска неизведанного.

декабрь 1983

### Об ориентированных кривых

Вероятно, в детстве все играли с лентой Мебиуса. Вещица действительно забавная — и стоит за ней математическая теория не совсем формального свойства. Бутылка Клейна попроще: это всего лишь двумерная поверхность в четырех измерениях, совершенно гладкая и



беспроблемная. А тут — поверхность с краем. Определить край многообразия в его в собственных терминах совершенно невозможно, надо выходить вовне, изучать вместе с чем-то еще — и многое зависит от выбора объемлющего пространства и способа вложения. По отношению к одному может получиться совсем не то, что вместе с другим. Поэтому рассуждения о пространствах с особыми точками выглядят не очень убедительно: остается ощущение подгонки под заданные параметры, теории *ad hoc*. Что, впрочем, не мешает науке оставаться занимательной и полезной.

Можно ли избавиться от разрывов и перейти к чему-то более удобному? Традиционно топологи делают это при помощи склейки краев — но у ленты Мебиуса только один край, и не очень понятно, что к чему клеить. Еще один стандартный прием — стягивание. Например, можно стянуть край в (выколотую) точку; такая конструкция остается сингулярной — но в каком-то смысле более проста.

Но есть и другой вариант. Заметим, что ширина ленты Мебиуса не влияет на ее топологические свойства. Вот и давайте сделаем ленту бесконечно узкой — превратим ее в замкнутую кривую. В итоге получится (одномерное) многообразие без края — и внешние особенности станут его внутренним строением. Ленту Мебиуса мы получаем склейкой концов обычной плоской ленты, с перекручиванием на пол-оборота. Остается выяснить, как перекрутить (пространственную) кривую перед замыканием концов.

Чисто визуально, каждой точке незамкнутой кривой сопоставляется трехмерный вектор ориентации, ортогональный направлению вдоль кривой; при переходе от одной точки к другой ориентация, вообще говоря, поворачивается в пространстве. Если склеить концы кривой так, чтобы ориентация их совпадала — получим обычную пространственную кривую, которую можно проецировать на плоскость так, что внутренняя и внешняя области оказываются отделены друг от друга, и можно говорить о внешней и внутренней нормали. Если же при склейке ориентация меняет знак — возникает аналог ленты Мебиуса; в проекции на плоскость, при движении вдоль кривой, внешняя нормаль после полного оборота скачком становится внутренней — причем это поведение не зависит от того, с какой точки мы начнем. Кажущаяся сингулярность никак не связана с гладкостью многообразия — это артефакт, следствие нелинейности операций проецирования и нормализации векторов.

Формально, в каждой точке пространственной кривой возникает локальная система ориентаций: одно из измерений задает направление вдоль кривой, другое — поперечное направление (для ленты — это движение от одного края к другому), а их векторное произведение определяет положение трехмерной нормали. Эта система ориентаций порождает внутреннее пространство каждой точки, которое следует отличать от внешних, присоединенных пространств (например, с образующими вдоль скорости и радиального ускорения); внешние характеристики кривой (кривизна, кручение и т. д.), вообще говоря, никак не соотносятся с внутренними (положение системы ориентаций).

При смещении вдоль замкнутой кривой менять знак может не только нормаль, но и направление вдоль кривой. В трехмерном представлении это выглядит как сингулярность, касп, попятное движение. Однако при вложении кривой в четырехмерное пространство возможно сохранить гладкость — и объяснить видимые разрывы выбором проекции. Аналогии из жизни хорошо известны. Например, математически гладкое движение маятника визуально выглядит как попятное движение при максимальном отклонении; мы знаем, что движение летящего вдалеке самолета может в проекции на плоскость наблюдения выстраиваться в причудливые кривые, способные навести кого-то на мысль об НЛО. Многие нелинейные эффекты и сингулярности в физике можно объяснить присутствием скрытых измерений, сложностью динамики, — вплоть до того, что наличие светового барьера (невозможность двигаться быстрее света) или сингулярности Шварцшильда могло бы указывать на многомерность физического

пространства, где на самом деле допустимы любые движения — но трехмерная проекция ставит искусственные границы.

Различие внешнего и внутреннего пространства в какой-то мере аналогично различию геометрии и физических полей. Внутренняя ориентация не зависит от преобразования внешних координат (включая зеркальное отражение) — это формальное выражение независимости физической системы от наблюдателя.

Движение вдоль замкнутой кривой и поворот нормали в плоскости, перпендикулярной продольному направлению, можно описывать двумя угловыми параметрами: фаза и ориентация. Если при изменении фазы на  $2\pi$  ориентация меняется на  $2\pi k$ , кривая соответствует простой ленте (перекрученной целое число раз); если ориентация меняется на  $2\pi(k + 1/2)$ , это аналог ленты Мебиуса (число перекручиваний полуцелое). В последнем варианте кривая, по сути дела, расщепляется на два слоя («листа»), которые последовательно пробегает каждый поворот фазы на  $4\pi$ . Можно нормировать этот удлинённый цикл на  $2\pi$  и превратить кривую Мебиуса в обычную кривую. Легко видеть, что это соответствует операции разреза ленты Мебиуса в продольном направлении (в результате чего получается обычная перекрученная лента). Вообще говоря, динамика изменения ориентации на первой и второй половинах удлинённого цикла может быть разной; при отсутствии связей это различие несущественно — сводится к переопределению параметра «времени».

В общем случае, на одном витке вдоль кривой ориентация меняется на  $2\pi q$ , где  $q$  — произвольное вещественное число; для рациональных  $q$  количество слоев будет конечным, а для иррациональных — кривая расслаивается в тороидальную поверхность. Если динамика поворота ориентации зависит от фазы, спектр (плотность заполнения поверхности витками кривой) может быть нетривиальным.

Изменение ориентации вдоль замкнутой кривой может быть проиллюстрировано физической моделью, когда в каждой точке мы обнаруживаем перпендикулярный направлению движения диполь; поляризация текущей точки влияет на поляризацию следующей. Конечное число слоев отвечает стоячей волне, а незамкнутость — бегущей волне вдоль кривой. Здесь возможны нетривиальные модельные интерпретации (например, магнитный монополь как кривая типа Мебиуса) и вполне практические приложения.

Ориентированные кривые — это иерархические структуры, в которых каждая точка внешнего пространства развертывается в некоторое внутреннее пространство. Помимо этого движение во внешнем пространстве порождает внешнее расслоение: конфигурационное пространство превращается в многообразие. Когда мы пытаемся расположить такую кривую в плоскости, речь уже не просто о соответствии точек базы точкам плоской кривой, а о проекции внутренних и внешних расслоений. В одних случаях это возможно (например, путем разбиения плоскости на несколько секторов, в каждом из которых размещается один слой) — в других приходится чем-то жертвовать. Аналогичная ситуация, когда речь уже не о кривых, а о более сложных многообразиях (в частности, поверхностях или телах).

Уровни иерархии существенно зависят друг от друга. Например, математическая иллюстрация: представим ориентированную кривую узким эллипсом на плоскости, оси которого поворачиваются при смещении вдоль кривой; если в конце эллипс принимает исходное положение — мы имеем дело с простой кривой.

Ориентированная кривая лишь по видимости одномерна. В принципе, в каждой ее точке возможно развертывание любых иерархических структур, а переход от одной точки к другой означает свертывание текущей структуры и развертывание другой (обращение иерархии). В частности, размерность внутреннего пространства может меняться — при вложении в ту же внешнюю геометрию.

январь 1985

### Виртуальная математика

В квантовой физике мы ссылаемся на внутренние состояния системы (точки конфигурационного пространства), которые не наблюдаются сами по себе, и судим мы о них по внешним проявлениям: нечто наблюдаемое происходит так, *как если бы* эти внутренние состояния существовали и взаимодействовали определенным образом. В принципе, ничто не мешает вычислять наблюдаемые величины каким-то иным способом — в котором прежние представления о виртуальных состояниях никак не участвуют. Это нормально: по жизни, одно и то же можно сделать разными способами, — так почему квантовую систему не построить из разных кирпичиков?

Аналоги этой виртуальности существуют и в математике, и связаны они с выходом за пределы области определения. Фундаментальные математики предпочитают о предметной области не упоминать — как будто математические объекты существуют сами по себе и не относятся вообще ни к чему. Но это иллюзия, самообман. Любая математическая теория предполагает конкретный универсум, поверх которого надстраиваются абстрактные формы. Строгая теория неприменима ни к чему помимо ее предмета, и суждения о чем-то внешнем логически некорректны. Тем не менее, чисто формально, мы можем ввести невозможные в теории (виртуальные) объекты — и построить новую (расширенную) теорию, в которой эти сущности заданы наряду с прочими определениями предмета.

Например, сложение двух натуральных чисел дает натуральное число, большее каждого из слагаемых. Но это означает, что оба слагаемых виртуально присутствуют в сумме, становятся ее внутренними компонентами. Допустим, что наблюдаемыми нашей теории (ее предметной областью) являются только четные числа. Тогда представление четного числа в виде суммы двух нечетных — совершенно аналогично представлению квантовых амплитуд суперпозициями виртуальных состояний. Более того, мы можем пойти дальше — и допустить существование совсем странных отрицательных чисел, способных уменьшать любое число при сложении. Если на практике мы научимся изготавливать вещи, представляющие эти воображаемые сущности, — они переходят в разряд наблюдаемых, и можно смело строить теорию целых чисел вообще.

Точно так же, произведение целых положительных чисел — число положительное; все вместе такие числа образуют предметную базу одной из важнейших математических теорий, где операция разложения на целые множители играет фундаментальнейшую роль. Каждое целое положительное число оказывается облаком виртуальных произведений — и существуют «простые» числа, для которых объем этого облака минимален; это своего рода базис нашего универсума. Для каждого числа можно описать его внутреннее строение набором «множеств уровня»  $L_n$ , содержащих все разложения в произведение  $n$  сомножителей (не учитывая единицу, которая по факту лишь задает единицу измерения, физическую размерность); на первом уровне, очевидно, оказывает само исходное число, а у простых чисел подуровни отсутствуют. Количество элементов («меру») множества  $X$  обозначим через  $\mu(X)$ . Определим теперь *индекс* целого положительного числа  $N$  как:

$$\lambda(N) = \sum \frac{\mu(L_n)}{n!}$$

Для простого числа индекс очевидно равняется единице; для всех остальных чисел он больше единицы. Например,  $\lambda(15) = 1.5$ , а  $\lambda(12) = 3.5$ . Порядок сомножителей имеет значение — но одинаковые сомножители мы здесь не различаем (хотя могут быть и другие определения). Факториальный вес введен из сугубо практических соображений: для нас важнее варианты, когда искомый результат получается за минимальное число шагов. С одной стороны, индекс характеризует иерархичность целого числа — а с другой, это показатель «продуктивности»: чем выше индекс, тем больше способов число изготовить на практике. Зная разложение числа на простые множители, легко вычислить его индекс — но выражение не из самых тривиальных.

Поэтому в теории индексов целых положительных чисел возникают эффекты, аналогичные квантовой интерференции, перепутыванию разных внутренних представлений.

Как и в аддитивном примере, возможно введение виртуально отрицательных целых чисел, которые можно считать произведением положительного числа на  $(-1)$  — подобно тому, как положительные числа предполагают общий («размерный») множитель 1; в составе физически наблюдаемых (положительных) чисел тогда будут присутствовать лишь цепочки с четным числом отрицательных сомножителей — или иначе: отрицательные размерности в природе рождаются только парами (как полюса магнита). По этому поводу существует весьма содержательная математика — а теория индексов имеет нетривиальные обобщения.

Переход к рассмотрению целых чисел любого знака сохраняет всю эту сложность — но в новой теории она отходит на второй план, по сравнению с общей алгебраической структурой, в которой предполагается, что отрицательные числа «наблюдаемы» сами по себе (то есть, можно их представить соответствующими вещами).

Аналогично, в теории вещественных чисел при любых вычислениях «физическим» считается именно вещественное значение — и введение мнимой единицы  $i = \sqrt{-1}$  есть явный выход из области определения квадратного корня; однако выражения в комплексных числах вполне допустимы как виртуальные пути, если в конечном итоге мы все равно получим нечто вещественное. Здесь аналогия с квантовой виртуальностью еще нагляднее: к одному и тому же вещественному числу можно приходить по разным траекториям в комплексной плоскости, и каждое вещественное число представляется иерархией комплексных путей (циклов, петель), для которой тоже возможно определить индексы — и развить содержательную математическую теорию. При наличии связей, топология предметной области может быть усложняться, и не факт, что для всякого числа найдется какой-нибудь цикл (хотя бы и нулевой длины). Виртуальность мнимой единицы в рамках вещественной теории — тоже связь, ограничение на множество допустимых путей. С другой стороны, такая теория может рассматривать альтернативные топологии комплексной плоскости, не предполагающие традиционных операций сложения и умножения. Но если мы собираемся перейти к математике комплексных чисел как таковых — правила работы с ними придется явно задать, выбирая одну из возможных (несводимых друг к другу) структур.

Еще один пример — отрицательные и комплексные множества в теории, где каждое множество ассоциировано с виртуальными путями его построения из других.

Точно так же, возможны нетрадиционные логики, в которых классические оценки истинности получаются в результате неклассических рассуждений.

В общем случае, в любой области математики возможны как классические теории, работающие только с объектами, определяемыми в рамках этой теории, — так и «квантовые» расширения, учитывающие разные способы виртуализации.

*декабрь 1982*

**СОДЕРЖАНИЕ**

|                                                             |    |
|-------------------------------------------------------------|----|
| Множества vs. алгебра логики .....                          | 1  |
| Отрицательные множества .....                               | 1  |
| Нечеткие множества и релятивистское сложение скоростей..... | 3  |
| Логические симметрии и комплексная логика .....             | 5  |
| Точки и пределы .....                                       | 8  |
| Иерархическая размерность .....                             | 12 |
| Квантовая теория множеств .....                             | 18 |
| Абстрактные картинки.....                                   | 25 |
| Предметная теория множеств .....                            | 29 |
| Кардинальная иерархичность.....                             | 33 |
| Качество отрицания .....                                    | 34 |
| Об ориентированных кривых .....                             | 40 |
| Виртуальная математика .....                                | 43 |