

А. Л. Семёнов

МАТЕМАТИКА ТЕКСТОВ



Научно-редакционный совет серии:

*В. В. Прасолов, А. Б. Сосинский,
В. М. Тихомиров (гл. ред.), И. В. Яценко.*

Серия основана в 1999 году.

Библиотека
«Математическое просвещение»
Выпуск 22

А. Л. Семёнов

МАТЕМАТИКА ТЕКСТОВ

Издательство Московского центра
непрерывного математического образования
Москва • 2002

УДК 510.5/.6
ББК 22.12
С30

Аннотация

В брошюре рассматриваются идеи и конструкции, лежащие в основе «математики текстов»; среди примеров её результатов — несчётность множества последовательностей из нулей и единиц, невозможность создать программу, распознающую самоприменимость программ. Обсуждается важное понятие *сложности текста по Колмогорову*, позволяющее отличать случайные тексты от неслучайных.

Текст брошюры представляет собой обработанную запись лекции, прочитанной автором 5 декабря 1999 года для участников III Международного математического турнира старшеклассников «Кубок памяти А. Н. Колмогорова» — школьников 8—11 классов. (Запись Е. Н. Осьмовой, обработки Р. М. Кузнеця.)

Для широкого круга читателей, интересующихся математикой: школьников старших классов, студентов младших курсов, учителей...

*Издание осуществлено при поддержке
Московской городской Думы
и Московского комитета образования.*

ISBN 5-94057-006-2

Семёнов Алексей Львович.

Математика текстов.

(Серия: «Библиотека „Математическое просвещение“»).

М.: МЦНМО, 2002. — 16 с.: ил.

Редактор *М. П. Каленков.*

Техн. редактор *М. Ю. Панов.*

Лицензия ИД № 01335 от 24/III 2000 года. Подписано к печати 4/X 2002 года.
Формат бумаги 60×88 $\frac{1}{16}$. Офсетная бумага № 1. Офсетная печать.
Физ. печ. л. 1,00. Усл. печ. л. 0,98. Уч.-изд. л. 1,15. Тираж ?000 экз. Заказ ????.

Издательство Московского центра непрерывного математического образования.
119002, Москва, Г-2, Бол. Власьевский пер., 11. Тел. 241 05 00.

Отпечатано в ФГУП «Производственно-издательский комбинат ВИНТИ».
140010, г. Люберцы Московской обл., Октябрьский пр-т, 403. Тел. 554 21 86.

Когда-то в детской энциклопедии я прочитал, что математика — это наука о числах и фигурах. «Числа и фигуры» — это название замечательной книжки о разных разделах математики, с которой и вам стоит познакомиться. Однако позже я понял, что в математике есть много такого, что можно только условно назвать числами или фигурами, например топологические пространства и абстрактные алгебры. Ещё в восьмом классе я побывал на лекции Андрея Андреевича Маркова для участников олимпиады и узнал, что можно строить всю математику, начиная с представления о *текстах* — последовательностях (цепочках) символов.

Впрочем, название «*математика текстов*» нетрадиционно, обычно употребляется название «*математическая логика и теория алгоритмов*». Вам, наверное, даже более знаком термин «*информатика*»: почти во всех школах есть предмет с таким названием. На уроках информатики в школе проходят разные вещи: так, если в школе есть компьютерный класс, то детей учат работать с компьютером. Но я употребил слово «информатика» в другом смысле. Наука, которую можно назвать также «*математическая информатика*» или «*теоретическая информатика*», в школьном курсе информатики изучается, но совсем немного. (Когда-то я участвовал в написании первого в нашей стране учебника по информатике для всех школьников 10 и 11 классов, и мы включили в него немного математики. Так происходит и сейчас. Но часто математику там трудно увидеть, и требуется некоторое усилие, чтобы понять, где же математическая суть в этой информатике.) Но всё это: математика текстов, математическая логика и теория алгоритмов, математическая информатика — относится к одной и той же области математики.

Математика текстов, в частности математика *математических текстов* (для этой математики ещё придумали название «*метаматематика*»), появилась раньше, чем компьютеры, она существовала ещё в прошлом веке. Необходимость и полезность этой науки хорошо осознавал один из величайших математиков конца XIX — начала XX века Давид Гильберт. Но именно с возникновением компьютеров математика текстов стала особенно важной областью математики.

Метаматематика изучает, в частности, устройство доказательств и аксиоматических систем в различных разделах математики, таких как арифметика, геометрия, теория множеств: что в этих системах можно доказать, чего доказать невозможно и т. д. Но в этой брошюре речь пойдёт о более понятных и наглядных вещах.

Начнём с нескольких примеров.

ЛОГИЧЕСКИЕ ПАРАДОКСЫ

Возьмём, например, такой простой текст:

Утверждение, записанное в последней и предпоследней строках на третьей странице этой брошюры, ложно.

Многие из вас, наверное, встречали в тех или иных книжках похожий текст. Напомню, что проблема (кстати вполне математическая) состоит в том, чтобы понять, истинно это утверждение или ложно. Подумайте над этим.

Вот ещё один пример.

Некоторые русские слова и словосочетания описывают числа (например «сто двадцать восемь», «миллион» и т. п.). Рассмотрим такое число:

наименьшее натуральное число, которое нельзя
описать текстом короче 100 символов (1)

(символами мы считаем буквы русского языка, цифры, знаки препинания, пробелы).

С одной стороны, ясно, что существуют числа, которые нельзя описать текстом короче 100 символов — просто потому, что таких текстов конечное число, а чисел бесконечно много. С другой стороны, наименьшее такое число описано текстом (1), в котором меньше 100 символов.

И наконец, ещё один пример.

Однажды учитель математики объявил ученикам, что на следующей неделе он проведёт контрольную работу. При этом накануне контрольной они не будут знать точно, что контрольная завтра. (Уроки математики проходят каждый день с понедельника по пятницу.)

Ясно, что в пятницу контрольной быть не может: тогда в четверг вечером школьники точно будут знать, что контрольная завтра. Но тогда в среду вечером они точно будут знать, что контрольная в четверг, ведь в пятницу её быть не может. Значит, и в четверг не может быть контрольной... Словом, ученики пришли к выводу, что контрольной вообще не будет, и успокоились. А контрольная произошла в среду.

Не пытаясь сейчас до конца выяснить причину странности, парадоксальности приведённых примеров, заметим только, что парадокс возникал при попытке применить в рассуждении конструкцию

«с одной стороны ..., с другой стороны ...»,

и при этом получались взаимно противоречащие друг другу утверждения.

Оказывается, эти конструкции могут быть применены к важным математическим построениям, уже не парадоксальным, хотя часто и демонстрирующим невозможность чего-то.

ДИАГОНАЛЬНЫЙ МЕТОД КАНТОРА

Рассмотрим квадратную таблицу $n \times n$ клеток, в которой некоторые клетки заштрихованы, а остальные клетки — белые. Пусть нам нужно построить такую строку, которой нет в этой таблице. Давайте последовательно выпишем в строку клетки, которые в таблице стоят по диагонали (рис. 1); полученную строку обозначим через d . Строка d может присутствовать в нашей таблице, а может и не присутствовать. Поменяем в этой строке цвета, т. е. сделаем белые клетки чёрными, а чёрные — белыми; переделанную таким образом строку обозначим через $\bar{d}$. Вот строки $\bar{d}$ в нашей таблице точно нет: от первой строки она отличается первой клеткой, от второй строки — второй клеткой, ..., от n -й строки — n -й клеткой.

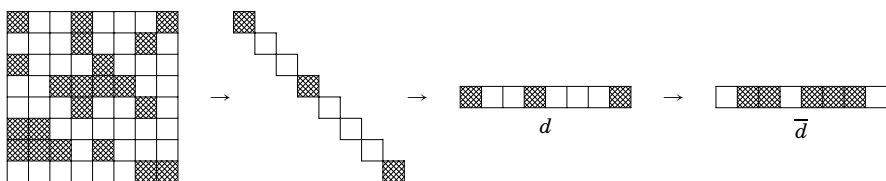


Рис. 1

Рассмотрим теперь бесконечную таблицу. Как и раньше, мы можем построить строку d из клеток, стоящих на диагонали (только теперь она получится бесконечной), и получить из неё строку $\bar{d}$ (рис. 2). Так же, как в конечном случае, доказывается, что бесконечной строки, совпадающей с $\bar{d}$, в таблице нет.

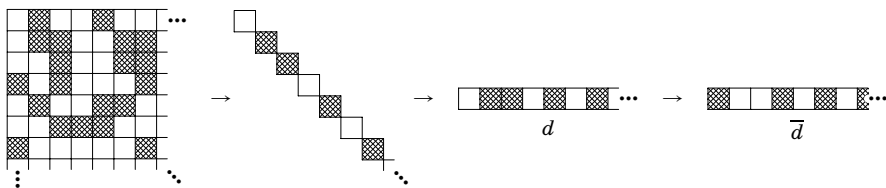


Рис. 2

Что же в нашей конструкции замечательного? Да, мы научились строить бесконечную строку, которой нет в нашей таблице. (Такое построение и называется *диагональным*.) Какие из этого можно извлечь следствия? Самое знаменитое из них получится, если мы предположим, что в нашей таблице имеются вообще все бесконечные строки. В этом случае наше построение сразу приводит к ложности такого предположения.

Мы установили следующий замечательный

Факт. Невозможно все бесконечные последовательности из нулей и единиц выписать в таблицу.

Иными словами, все такие последовательности нельзя пересчитать, перенумеровать, решить, какая последовательность будет первой строкой в нашей таблице, какая — второй и т. д. Говоря научно, невозможно установить взаимно однозначное соответствие между натуральными числами и всеми последовательностями из нулей и единиц.

Эту знаменитую теорему впервые доказал в XIX веке Георг Кантор, основатель теории множеств.

Контрольный вопрос. Можно ли выписать в бесконечную таблицу все строчки, в которых лишь конечное число чёрных клеток, а остальные — белые?

ПРОГРАММЫ

Некоторые тексты, в частности, тексты на русском языке, являются *программами*, или *алгоритмами*, — предписаниями, которые можно выполнить. То, к чему применяется программа, тоже текст. Например, алгоритм сложения двух натуральных чисел «столбиком» — это следующий текст на русском языке:

1. Запиши два числа друг под другом «столбиком».
2. Если числа различаются по длине, дополни короткое слева нулями, чтобы их длины сравнялись.
3. Запомни 0 в уме.
4. Рассмотрим самую правую пару цифр (по вертикали), которая ещё не была рассмотрена. Сложи эти две цифры и добавь то, что в уме. Если сумма меньше 10, напиши её под взятой парой, запомни в уме 0; если сумма больше или равна 10, запиши под взятой парой последнюю цифру суммы и запомни в уме 1.
5. Если все цифры слагаемых рассмотрены и в уме 0, то результат уже получен в нижней строке; если все цифры слагаемых рассмотрены и в уме 1, допиши справа 1 к нижней строке, и это — результат; если не все цифры слагаемых рассмотрены, перейди к пункту 4.

Применяется эта программа к тексту, который представляет собой два числа, записанные в десятичной системе счисления

Двумя чертами слева выделен материал, над которым читателю рекомендуется тщательно подумать самостоятельно, прежде чем читать раздел «Решения и комментарии» (стр. 12).

и разделённые запятой и пробелом, например к такому:

4850493897329785961, 29785565428497

Бывают программы, распознающие некоторое свойство. Такие программы на любом тексте*) дают ответ либо «да», либо «нет». Например, можно себе представить программу, которая определяет, просто ли данное число.

Ещё один пример программы:

1. К данному тексту припиши в конце букву «А».

Это очень простая программа, она заканчивает работу на любом тексте всего за один шаг. Бывают программы, которые не заканчивают работу, как, например, такая:

1. К данному тексту припиши в конце букву «А».
2. Перейди к пункту 1.

Она припишет к тексту букву «А» в конце, потом ещё раз припишет букву «А», потом ещё раз припишет букву «А» и т. д.

Встречаются и более сложные случаи, когда программа на одних текстах заканчивает работу, а на других — нет. Если программа Π на тексте T не заканчивает работу, будем это обозначать так: $\Pi(T) = \infty$. Если же программа Π заканчивает работу на тексте T будем писать: $\Pi(T) = !$.

Самоприменимые программы

Поскольку каждая программа представляет собой текст, то мы можем запустить её на самой себе, на собственном тексте. Программа Π называется *самоприменимой*, если $\Pi(\Pi) = !$. А можно ли по программе Π узнать, является ли она самоприменимой?

Если просто запустить программу Π на тексте Π , то со временем может выясниться, что она закончила работу. Но ни в какой конечный момент времени, если программа всё ещё работает, мы не можем быть уверены, что она никогда не закончит работать. Ведь может случиться так, что программа работает несколько суток, решая какую-нибудь сложную задачу, и это не означает, что она собирается работать бесконечно долго; может быть, она закончит работу через месяц, или через миллион лет.

Но может быть, есть какой-то другой способ читать текст программы понять, что она не кончает работать на самой себе?

Допустим, что существует программа S , которая, получая текст программы Π , распознаёт, является ли Π самоприменимой, и даёт

*) Математики говорят о работе «программы на тексте» вместо более длинного выражения «программы в применении к тексту».

ответ либо «да», либо «нет»^{*)}). Тогда мы можем построить программу $\tilde{S}$, которая, получив на входе программу Π , запускает программу S и, если $S(\Pi)$ = «да», работает бесконечно долго (выполняет какую-нибудь бесконечную процедуру, например, бесконечное число раз умножает 0 на 0), а если $S(\Pi)$ = «нет», тоже говорит «нет», т. е. заканчивает работу. Таким образом, программа $\tilde{S}$ не оканчивает работу на самоприменимых программах и заканчивает — на несамоприменимых.

Заметим, что конструкция здесь тоже «диагональная», как и в случае таблицы с закрашенными клетками. Только сейчас мы рассматриваем таблицу, в которой и по горизонтали, и по вертикали выписаны все программы (точнее все конечные последовательности символов того языка, на котором мы решили записывать программы (рис. 3), при этом, если текст не является программой, мы можем считать, что он просто ничего не делает с исходными данными). В клетки таблицы мы записываем результат работы программы (стоящей в заголовке столбца) на исходном тексте (стоящем в заголовке строки).

	A	Б	АА	АБ	БА	ББ	В	ВАБ	...
A	АБ	«да»	A	)	p?	∞	A	A	
Б	!,7	«да»	Б	Г	ё	∞	Б	Б	
АА	у,	«нет»	АА	Ж	»	∞	АА	АА	
АБ	3,	∞	АБ	,	—«	∞	АБ	АБ	
БА	ьы	«да»	БА	∞	∞	∞	БА	БА	
ББ	1 1	«нет»	ББ	2ы3	фщ	∞	ББ	ББ	
В	й	«нет»	В	(1(	ч)!	∞	В	В	
ВАБ	т	«нет»	ВАБ	2Ц1	600:	∞	ВАБ	ВАБ	
:									

Рис. 3

Если программа S существует, то существует и программа $\tilde{S}$. Предположим, что $\tilde{S}(\tilde{S}) = !$, т. е. $\tilde{S}$ кончает работу. Но тогда $S(\tilde{S})$ = «нет» по определению $\tilde{S}$, следовательно, $\tilde{S}$ несамоприменима. И наоборот, если $\tilde{S}(\tilde{S}) = \infty$, то $S(\tilde{S})$ = «да» и $\tilde{S}$ самоприменима. Противоречие.

^{*)} Можно рассмотреть и более сложную программу, которая получает два текста: программу Π и произвольный текст X — и определяет, заканчивает ли Π работу на X . Но о программах, которые применяются к парам текстов — чуть позже.

Итак, мы установили следующий

Факт. Не существует программы, которая распознавала бы самоприменимые программы.

Это один из фундаментальных фактов теории алгоритмов, а конструкция, на которой основано его доказательство, — одна из важнейших в этой теории.

Конечно, вопрос о самоприменимости программ кажется довольно экзотическим: не так уж часто мы применяем программу к самой себе. А вот применение программы к другой программе — обычное дело: операционная система всякого реального компьютера (а она, конечно, является программой) выполняет другие программы, заданные человеком, т. е. применяется к ним.

Заметим, что и примеры, приводимые ранее, когда не удавалось определить истинность простого на вид утверждения, выглядят «далёкими от жизни», и неясно, стоит ли им придавать серьёзное значение. Попробуем объяснить, почему дело обстоит серьёзно. Первая причина состоит в том, что один контрпример — в математике уже опровержение общей теории или метода. Вторая причина в том, что как-то отделить «диагональные» случаи от «недиагональных» не удаётся, и это может быть доказано после соответствующего уточнения понятий. Однако во многих случаях и математики спрашивают друг друга: «А есть ли „естественный“ недиагональный пример?».

Универсальная программа

Иногда возникает необходимость в программах, которые применяются к парам текстов. Каким образом можно из пары текстов сделать один? Записать два текста рядом, один за другим? К сожалению, после этого уже невозможно будет определить, где кончается первый текст и начинается второй. Чередовать буквы первого и второго текста? Но, если тексты разной длины, снова не удастся разделить полученный текст на два исходных*).

Попробуем теперь применить более сложный приём. Запишем первый текст, удваивая каждую его букву. Например, текст

Универсальная_программа

(знак _ обозначает пробел) запишется так:

УУннииввееррссалльньнаая_пппрооггррааммммаа

После первого текста запишем *разделитель* «аб». Потом запишем второй текст без изменений. Допустим, что второй текст выглядит так:

Сложность_текста

*) Вспомним, когда в алгоритме сложения «столбиком» мы из пары чисел делали один текст для работы программы на нём, мы использовали разделитель «,» — запятую и пробел, но это было возможно, поскольку в записи самих чисел эти два символа встречаться не могли.

Тогда окончательный результат будет таким:

УУнниивееррссаалпыннаая_л_пппрооггррааммммаабСложность_текста

Покажем, что по такой записи можно однозначно восстановить оба текста. Будем читать из полученной строки по два символа. Если прочитанные символы совпадают, то они представляют один символ первого текста (рис. 4). Если в некоторый момент мы натолкнёмся на два различных символа, то это означает, что мы дошли до разделителя «аб», который при восстановлении текстов нужно выбросить. После разделителя прочитаем по одному символу второй текст.

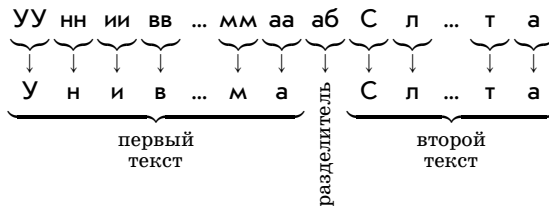


Рис. 4

Оказывается, существуют способы кодировать пару текстов одним текстом, причём делать это экономно, добавляя не слишком много новых букв, так чтобы длина кода была довольно близка к сумме длин исходных текстов.

Задача. Насколько длина кода пары текстов может быть близкой к сумме их длин?

Будем считать, что мы уже научились из двух текстов делать один. После этого можно построить программу U , которая применяется к паре текстов P и X , а результатом её работы будет $P(X)$. Это тоже один из важнейших фактов теории алгоритмов.

Факт. Существует *универсальная* программа, которая, получая на вход программу P и текст X , применяет P к X .

Это утверждение не кажется очень сложным, хотя полные формальные доказательства его довольно громоздки.

СЛОЖНОСТЬ ТЕКСТА. СЛУЧАЙНЫЕ И НЕСЛУЧАЙНЫЕ ПОСЛЕДОВАТЕЛЬНОСТИ

Не так давно, в середине 60-х годов XX века, было выяснено, что математика текстов может быть использована для прояснения наших представлений о природе случайности. Начнём с обсуждения того, что нам самим кажется случайным, а что неслучайным.

Предположим, что мы бросаем монету и записываем результаты в виде последовательности из нулей и единиц: если выпадает орёл, пишем ноль, а если решка — единицу. Может ли получиться такая

последовательность:

0110110010111010110100100000011110100111? (2)

Наверное, может. А такая:

$$000000000000000000000000000000000000000? \quad (3)$$

Вряд ли.

Почему, глядя на эти последовательности, мы сразу понимаем: так бывает, а вот так — не бывает? Первый приходящий в голову ответ: потому что первая последовательность «вероятна», а вторая — «невероятна». Однако обе последовательности имеют одинаковую вероятность! (Равную 2^{-n} , где n — длина последовательности.) Чем случайные последовательности отличаются от неслучайных? Попробуйте прояснить этот вопрос привела к такому важному понятию как *сложность конечной последовательности*, или *сложность текста*.

Сейчас мы временно отложим вопрос о случайности и поговорим о сложности текстов.

Бывает, что короткий текст описывает длинный. Например, текст

миллион слов «да»,

состоящий из 17 символов, является описанием текста длиной два миллиона символов.

Пусть Π — какая-нибудь программа. *Сложностью* текста S при способе описания Π называется наименьшая длина текста T , такого что $\Pi(T)=S$. Текст T при этом называется *описанием* текста S .

Иными словами, мы берём текст S , рассматриваем все его описания, из которых программа Π может получить S , и среди них ищем самое короткое. Например,

МИЛЛИОН СЛОВ «ДА» И ПЯТСОТ ТЫСЯЧ СЛОВ «ДАДА»

— два описания одного текста, но первое описание короче.

Заметим, что для последовательности (3) мы сразу можем указать короткое описание, в отличие от последовательности (2). Естественно ожидать, что у большинства последовательностей нет коротких описаний. Такие последовательности и называются *случайными*.

Это понятие случайности (его в 1960-х годах ввёл Андрей Николаевич Колмогоров) оказывается очень интересным и продуктивным.

Докажите, что существует способ, который даёт почти самые короткие описания для всех последовательностей. Более подробно, существует такой способ K , что сложность любого текста S при способе описания K меньше, чем его сложность при любом другом способе описания Π плюс некоторая постоянная, зависящая только от Π .

(Эту теорему доказал А. Н. Колмогоров.)

Построение этого способа описания K использует конструкцию универсальной программы, о которой говорилось ранее. Попробуйте сами построить этот способ. У вас получится!

РЕШЕНИЯ И КОММЕНТАРИИ

К стр. 6.

О т в е т: можно. Будем выписывать строки в следующем порядке. Сначала пустую строку и строку с одной чёрной клеткой на самом левом, первом месте. Затем выпишем все строки, у которых не закрашены все клетки, кроме левых двух: таких строк осталось 2. Потом все строки, у которых закрашены только некоторые из левых трёх клеток (их осталось 4). Вообще, на n -м шаге ($n > 1$) выписываются все строки, в ко-

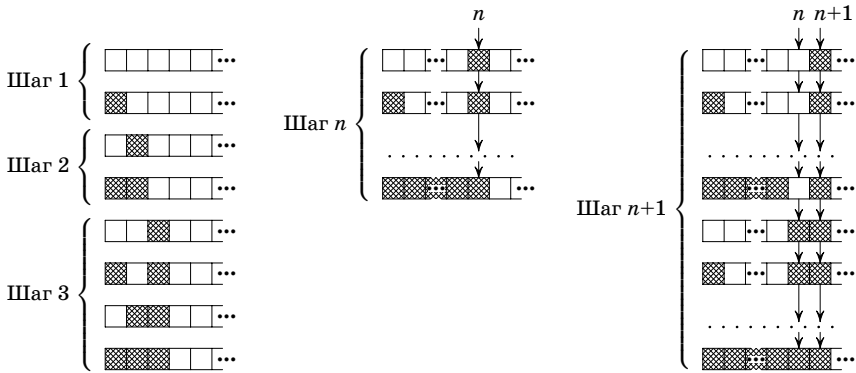


Рис. 5

торых закрашенные клетки все находятся на первых n местах, в количестве 2^{n-1} строк (рис. 5). Нетрудно убедиться, что таким образом будут выписаны все возможные строки с конечным числом чёрных клеток.

К стр. 10.

Обозначим через $N(k)$ текст, кодирующий число k . Чтобы получить такой текст, занумеруем символы нашего алфавита числами от 0 до $n-1$, где n — количество символов в нашем алфавите. Теперь просто запишем число k в системе счисления с основанием n и заменим все цифры в этой записи на соответствующие символы нашего алфавита. Легко понять, что по такому тексту однозначно восстанавливается число k . Как известно, длина числа k в системе счисления с основанием n равна $\lceil \log_n(k+1) \rceil$ (*). Если мы не будем писать в начале числа незначащие нули, то длина текста, кодирующего число k , равна $\lceil \log_n(k+1) \rceil$. При $k=0$ будем считать, что $N(k)$ — это пустой текст, ведь мы выкидываем из числа все нули слева. Но это только вопрос договорённости.

*) $\lceil x \rceil$ — это наименьшее целое число, не меньшее x .

Пусть $|X|$ — длина текста X . Если бы нам удалось передать программе эту длину вместе с текстом XY , то мы легко смогли бы восстановить текст X , просто прочитав нужное количество символов сначала. Таким образом, вместо того, чтобы закодировать пару текстов X и Y , мы можем закодировать пару текстов $N(|X|)$ и XY . На первый взгляд такая замена никак не упрощает задачу, но это не совсем так по двум причинам.

Во-первых, если мы просто удвоим все символы текста $N(|X|)$, припишем к ним справа два разных символа из нашего алфавита, а потом припишем ещё и текст XY , мы сможем, как это делали раньше, восстановить тексты $N(|X|)$ и XY , а значит, и тексты X и Y . Таким образом, мы закодировали пару текстов X и Y текстом длиной $2\lceil\log_n(|X|+1)\rceil + 2 + |X| + |Y|$, что почти всегда меньше, чем $2|X| + 2 + |Y|$ символов, которые получаются при использовании старого способа. Ниже мы ещё улучшим этот новый способ кодирования пары текстов.

Во-вторых, длина текста $N(|X|)$ сильно зависит от длины текста X . Как это можно использовать? Можно придумать ещё один способ кодирования. Пусть $l = |X| + |Y|$ ($= |XY|$). Тогда $|X|$ может принимать всего $l + 1$ значение — от 0 до l . Все эти числа кодируются различными текстами, длина максимального из которых равна $\lceil\log_n(l+1)\rceil$. На самом деле можно считать, что все тексты имеют такую длину, ведь если число записывается менее чем $\lceil\log_n(l+1)\rceil$ символами, то мы всегда можем приписать слева несколько незначащих нулей. Таким образом, нам надо закодировать два текста: один длины l , а другой — $\lceil\log_n(l+1)\rceil$. Казалось бы, что и это никак не упрощает задачу. На самом деле нам даже не надо ничего кодировать — нужно просто записать сначала текст длины $\lceil\log_n(l+1)\rceil$, а потом приписать к нему справа текст длины l .

Как же мы сможем понять, где кончается первый текст и начинается второй? Посмотрим, какой информацией мы располагаем. У нас есть текст, и мы хотим найти, где в нём закодирован первый текст, а где — второй. В самом тексте мы ничего найти не можем. Но дополнительной информацией нам будет служить его длина. В самом деле, наш текст имеет длину $\lceil\log_n(l+1)\rceil + l$, а зная это число можно однозначно восстановить l , так как при увеличении l на единицу число $\lceil\log_n(l+1)\rceil + l$ тоже возрастает (иногда на единицу, а иногда на двойку). Значит, зная $\lceil\log_n(l+1)\rceil + l$, можно найти число l , т. е. разделить наш текст на два, которые он кодировал. Теперь мы знаем $|X|$ и XY , а значит, и сами тексты X и Y . Для любой пары текстов X и Y суммарной длины l этот алгоритм кодирует их, используя только $\lceil\log_n(l+1)\rceil$ «лишних» символов. Если, например $n = 3$, а $|X| = |Y| = 1\,000\,000$, то наш алгоритм использует лишь 14 «лишних» символов!

Оказывается, приведённый выше алгоритм в некотором смысле оптимален. Имеется в виду, что любой другой способ кодирования двух текстов, так чтобы их после этого можно было однозначно восстановить, не может быть всегда лучше нашего. Иными словами, для любого числа l найдутся такие два текста X и Y суммарной длины не больше l , что этот другой алгоритм закодирует их, используя не менее $\lceil\log_n(l+1)\rceil$ «лишних» символов.

Пусть какой-то алгоритм умеет кодировать все пары текстов X и Y суммарной длины не больше l , используя не более k «лишних» символов. Подсчитаем, сколько всего существует таких пар текстов. Пусть $|X| + |Y| = m$, где $0 \leq m \leq l$. Посмотрим внимательно на текст XY . Из каждого такого текста длины m получается $m + 1$ различных пар текстов X и Y — X может быть первыми m символами этого текста, первыми $m - 1$ символами, ..., первым 1 символом, первыми 0 символами. При этом из разных текстов длины m будут получаться разные пары текстов X и Y . Легко понять, что каждая пара текстов X и Y получится ровно из одного текста длины m , а именно, из текста XY . Всего существует n^m текстов длины m , значит, пар текстов X и Y суммарной длины m ровно $(m + 1)n^m$.

Таким образом, количество пар текстов X и Y суммарной длины не больше l равно

$$S_l = \sum_{i=0}^l (i+1)n^i.$$

Вычислим S_l . С одной стороны, $S_{l+1} = S_l + (l+2)n^{l+1}$, а с другой,

$$\begin{aligned} S_{l+1} &= 1 + \sum_{i=1}^{l+1} (i+1)n^i = 1 + \sum_{i=0}^l ((i+1)+1)n^{i+1} = 1 + n \sum_{i=0}^l ((i+1)n^i + n^i) = \\ &= 1 + nS_l + n \sum_{i=0}^l n^i = 1 + nS_l + \frac{n(n^{l+1}-1)}{n-1}. \end{aligned}$$

Отсюда находим, что

$$S_l = \frac{(l+2)n^{l+1}}{n-1} - \frac{n^{l+2}-1}{(n-1)^2}.$$

Каждой паре текстов X и Y должен соответствовать свой текст не длиннее $l+k$ символов. Всего таких текстов

$$T_{l+k} = 1 + n + n^2 + \dots + n^{l+k} = \frac{n^{l+k+1}-1}{n-1}.$$

Из-за того, что можно провести однозначное декодирование, $T_{l+k} \geq S_l$, или, что то же самое, $(n-1)T_{l+k} \geq (n-1)S_l$. Имеем:

$$n^{l+k+1} - 1 \geq (l+2)n^{l+1} - \frac{n^{l+2}-1}{n-1},$$

т. е.

$$n^k \geq (l+2) - \frac{n-n^{-l}}{n-1} = l+2-1 - \frac{1-n^{-l}}{n-1} > l.$$

А поскольку число n^k — натуральное, выполняется неравенство $n^k \geq l+1$, откуда следует, что

$$k \geq \log_n(l+1) \quad \text{или} \quad k \geq \lceil \log_n(l+1) \rceil.$$

Что и требовалось.

В приведённом алгоритме количество «лишних» символов зависит от суммарной длины текстов X и Y . Это подразумевает некоторое равноправие текстов X и Y . Однако, часто бывает, что тексты имеют различную природу, например, текст X может быть значительно короче, чем Y . Тогда приведённый алгоритм будет не всегда оптимален. Поэтому хотелось бы найти экономный алгоритм, при использовании которого количество «лишних» символов зависело бы только от $|X|$ (*). Два таких алгоритма уже были рассмотрены: они использовали $|X|+2$ и $2\lceil \log_n(|X|+1) \rceil + 2$ «лишних» символов соответственно. Видно, что второй алгоритм экономичнее первого. Это связано с тем, что пару текстов X и Y заменили на пару $N(|X|)$ и XY . Логично было бы сделать такую замену ещё раз, а именно, заменить пару $N(|X|)$ и XY на пару $N(|N(|X|)|)$ и $N(|X|)XY$, по которой тексты $N(|X|)$ и XY , а значит, и тексты X и Y восстанавливаются однозначно. Но $|N(|X|)| > 0$ (иначе $N(|X|)$ — пустой текст, и такая замена не даёт никакого выигрыша). Поэтому ещё лучше вместо текста $N(|N(|X|)|)$ взять текст $N(|N(|X|)|-1)$. Таким образом, из пары текстов $A_0 = N(|X|)$ и $B_0 = XY$ получится

*) Например, такой алгоритм будет использоваться в решении следующего упражнения.

пара $A_1 = N(|N(|X)| - 1)$ и $B_1 = N(|X|)XY$. Количество «лишних» символов при её кодировании равно

$$2\lceil \log_n \lceil \log_n (|X| + 1) \rceil \rceil + 2 + \lceil \log_n (|X| + 1) \rceil.$$

Эту формулу можно упростить, если заметить, что

$$\lceil \log_n \lceil \log_n x \rceil \rceil = \lceil \log_n \log_n x \rceil$$

(докажите это). Количество «лишних» символов получается равным

$$2\lceil \log_n \log_n (|X| + 1) \rceil + 2 + \lceil \log_n (|X| + 1) \rceil,$$

что почти всегда меньше $2\lceil \log_n (|X| + 1) \rceil + 2$.

Если теперь заменить пару A_1 и B_1 на $A_2 = N(|A_1| - 1)$ и $B_2 = A_1 B_1$, то можно получить ещё более экономичный алгоритм. До каких же пор имеет смысл делать такие замены? Ясно, что когда очередной текст A_k будет иметь длину 1, продолжать замены бессмысленно. С другой стороны, если $|A_k| = 1$, то не нужно удваивать символы $|A_k|$ и тратить два символа на разделитель, так как можно просто восстановить текст A_k — его длина один символ. По A_k и B_k можно восстановить A_{k-1} и B_{k-1} , по A_{k-1} и B_{k-1} — A_{k-2} и B_{k-2} , и т. д. до A_0 и B_0 , а по A_0 и B_0 — X и Y . Теперь можно сформулировать и сам алгоритм кодирования. Надо получить текст $A_0 = N(|X|)$, по нему — текст $A_1 = N(|A_0| - 1)$, затем $A_2 = N(|A_1| - 1)$, и вообще получать текст $A_{i+1} = N(|A_i| - 1)$ до тех пор, пока длина очередного A_k не окажется равной единице. После этого запишем текст

$$S = A_k A_{k-1} \dots A_0 0XY,$$

где 0 — символ с номером ноль. Зачем он нужен? Для того, чтобы можно было понять, что начинается непосредственно текст X , а не очередная длина. В самом деле, никакой из текстов A_i не может начинаться с нуля, так как A_i суть числа, у которых первый символ точно не ноль. Следует отметить, что число ноль кодируется пустым текстом. Более того, среди $A_1, \dots, A_k$ не может быть пустых текстов, так как если A_i пуст, то текст A_{i-1} имел бы единичную длину, а значит, $k = i - 1$ и текст A_i вообще не мог использоваться (по определению k). Если A_0 является пустым текстом, то $|X| = 0$ и наш текст S есть не что иное, как $A_0 0XY = 0Y$, а $k = 0$.

При использовании этого алгоритма количество «лишних» символов равно

$$1 + (\lceil \log_n (|X| + 1) \rceil + \lceil \log_n \log_n (|X| + 1) \rceil + \dots + \underbrace{\lceil \log_n \dots \log_n (|X| + 1) \rceil}_{k+1 \text{ раз}}),$$

если $|X| > 0$, и равно 1, если $|X| = 0$. Важно отметить, что этот способ не оптимальный, но даёт очень хорошие результаты. Так, при $n = 3$ и $|X| = 1\,000\,000$ количество «лишних» символов равно 18 вне зависимости от текста $|Y|$.

Приведём ещё алгоритм декодирования текстов X и Y по тексту S , оставляя читателю самому проверить его корректность.

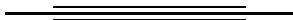
1. Запомни число 0.
2. Прочитай один символ текста S . Если это символ с номером 0, то следующие столько символов из S , какое число в памяти — это текст X , а остальные непрочитанные символы — Y . Иначе припиши к прочитанному символу справа ещё столько символов текста S , какое число находится в памяти. Полученный текст расшифруй как число в системе счисления с основанием n , а полученное число запомни.
3. Если тексты X и Y найдены, то закончи работу. Иначе перейди к пункту 2.

К стр. 11.

На самом деле, идея построения достаточно проста. Мы уже научились строить универсальную программу U , применение которой к программе Π и тексту X даёт $\Pi(X)$:

$$U(\Pi, X) = \Pi(X).$$

Пусть у нас имеется произвольный способ описания Π . Тогда $U(\Pi, X) = \Pi(X) = S$ для любого описания X текста S , т. е. пара текстов (Π, X) является описанием текста S при способе U . Даже не вдаваясь в детали оптимального по длине кодирования пары текстов, при тривиальном, рассмотренном нами способе, для пары текстов: Π длины k и X длины l — длина кода пары равнялась $2k + l + 2$. Так что сложность описания текста S при способе U превосходит сложность при любом другом способе Π не более чем на $2k + 2$, где k — длина текста самого способа Π — это и есть константа, зависящая только от способа Π .



**БИБЛИОТЕКА
«МАТЕМАТИЧЕСКОЕ ПРОСВЕЩЕНИЕ»**

ВЫПУСК 1

В. М. Тихомиров. Великие математики прошлого и их великие теоремы.

ВЫПУСК 2

А. А. Болибрух. Проблемы Гильберта (100 лет спустя).

ВЫПУСК 3

Д. В. Аносов. Взгляд на математику и нечто из неё.

ВЫПУСК 4

В. В. Прасолов. Точки Брокара и изогональное сопряжение.

ВЫПУСК 5

Н. П. Долбилин. Жемчужины теории многогранников.

ВЫПУСК 6

А. Б. Сосинский. Мыльные плёнки и случайные блуждания.

ВЫПУСК 7

И. М. Парамонова. Симметрия в математике.

ВЫПУСК 8

В. В. Острик, М. А. Цфасман. Алгебраическая геометрия и теория чисел: рациональные и эллиптические кривые.

ВЫПУСК 9

Б. П. Гейдман. Площади многоугольников.

ВЫПУСК 10

А. Б. Сосинский. Узлы и косы.

ВЫПУСК 11

Э. Б. Винберг. Симметрия многочленов.

ВЫПУСК 12

В. Г. Сурдин. Динамика звёздных систем.

ВЫПУСК 13

В. О. Бугаенко. Уравнения Пелля.

ВЫПУСК 14

В. И. Арнольд. Цепные дроби.

ВЫПУСК 15

В. М. Тихомиров. Дифференциальное исчисление (теория и приложения).

ВЫПУСК 16

В. А. Скворцов. Примеры метрических пространств.

ВЫПУСК 17

В. Г. Сурдин. Пятая сила.

ВЫПУСК 18

А. В. Жуков. О числе π .

ВЫПУСК 19

А. Г. Мякишев. Элементы геометрии треугольника.

ВЫПУСК 20

И. В. Яценко. Парадоксы теории множеств.

ВЫПУСК 21

И. Х. Сабитов. Объёмы многогранников.

ВЫПУСК 22

А. Л. Семёнов. Математика текстов.



МОСКОВСКИЙ ЦЕНТР НЕПРЕРЫВНОГО МАТЕМАТИЧЕСКОГО ОБРАЗОВАНИЯ

- учреждён Московским комитетом образования, префектурой Центрального административного округа Москвы, Отделением математики РАН, Математическим институтом им. В. А. Стеклова РАН, Московским государственным университетом им. М. В. Ломоносова, Независимым Московским университетом.
- ставит своей целью сохранение и развитие традиций математического образования в Москве, организацию и поддержку различных форм внеклассной работы со школьниками, методическую помощь руководителям кружков и преподавателям классов с углублённым изучением математики.
- является некоммерческой организацией и не стремится к извлечению прибыли. Обучение школьников, студентов, аспирантов и преподавателей средней школы в рамках программ МЦНМО является бесплатным.
- совместно с Московским институтом открытого образования организует курсы повышения квалификации московских учителей математики.
- организует математические и физические кружки, конкурсы, олимпиады и турниры для школьников, участвует в организации классов с углублённым изучением математики.

— осуществляет информационную поддержку большинства московских олимпиад для школьников, информация о них представлена на сервере МЦНМО
<http://www.mccme.ru/olympiads>

ISBN 5-94057-006-2



9 785940 570066

Адрес МЦНМО: 119002, Москва, Г-2,
Бол. Власьевский пер., 11.
Телефоны для справок: 241 05 00, 241 12 37.