

УДК 519.767.6:577

АНАЛИЗ ГЕНЕТИЧЕСКИХ ТЕКСТОВ. I. 1-ГРАММНЫЕ ХАРАКТЕРИСТИКИ

В.Д.Гусев, В.А.Куличков, Т.Н.Титкова

Назовем текстом конечную упорядоченную символьную последовательность, составленную из элементов конечного алфавита. Пусть n — длина текста (в символах), $n = |A|$ — число элементов (мощность) алфавита A . Подпоследовательность текста, состоящую из 1 расположенных подряд символов, будем называть 1-граммой ($1 = 1, 2, \dots, n$), а анализ текста, основанный на подсчете частот встречаемости в нем различных 1-грамм, — 1-граммным анализом.

Под генетическим текстом будем понимать представление молекул нерегулярных полимеров (ДНК, РНК, белков) в виде упорядоченной последовательности мономеров (нуклеотидов или аминокислот). Важность исследования данного класса молекул обусловлена тем, что они определяют протекание наиболее существенных генетических процессов в организмах. Алфавит нуклеотидов состоит из 4 символов (А, Т, Г, Ц для ДНК и А, У, Г, Ц для РНК), алфавит аминокислот — из 20 символов.

В последние годы благодаря совершенствованию техники определения нуклеотидного состава молекул были расшифрованы и опубликованы первые полные нуклеотидные последовательности (геномы) простейших микроорганизмов (*X174*, *G4*, *SV40*, *PBB322*, *FD*, *M32* и ряда других). Длины этих геномов составляют порядка ($10^3 - 10^4$) символов. Весьма актуальной задачей становится создание и пополнение банка данных по микробиологическим объектам и пакетов прикладных программ по исследованию различных свойств этих объектов.

Цель данной работы — исследование 1-граммных характеристик текстов, кодирующих наследственную информацию простейших микроорганизмов. Указан ряд содержательных задач, для решения которых представляется перспективным использовать 1-граммное описание текста. Ряд примеров иллюстрирует возможность выявления по 1-грам-

мым характеристикам функционально значимых фрагментов (или особенностей) текста.

Отметим, что наряду с информацией о частотах встречаемости различных 1-грамм в тексте очень важна и информация о местах вхождения каждой конкретной 1-граммы в текст. Эта информация зачастую дает возможность отличать случайные совпадения 1-грамм от неслучайных, функционально значимых.

1. Структура геномов вирусных частиц [1]. Из шести исследованных геномов (см. выше) пять первых представляют собой кольцевые молекулы ДНК. Поскольку двуцепочечные молекулы ДНК построены по принципу комплементарности (т.е. любая цепь однозначно получается из другой повсеместной заменой нуклеотидов А,Т,Г,Ц нуклеотидами Т,А,Ц,Г соответственно), анализу подвергалась лишь одна из них. В тех случаях, когда это было возможно, анализировались цепи, соответствующие структуре м-РНК (ΦХ174, G4, FD, MS2). Шестой геном (MS2) описывает одноцепочечную молекулу РНК.

Структура геномов фагов и вирусов определяется основными генетическими процессами, обеспечивающими воспроизводство этих организмов. К таким процессам относятся редупликация (копирование ДНК или РНК), транскрипция (синтез РНК на ДНК) и трансляция (синтез белка на РНК). Части текста, отвечающие отдельным актам редупликации, транскрипции и трансляции, называются соответственно репликаном, скриптом и цистроном. Каждая из этих функциональных единиц имеет свои знаки начала и конца — инициаторы и терминаторы, представляющие собой нуклеотидные последовательности длиной от трех до нескольких десятков символов.

Между перечисленными функциональными единицами существует, как правило, отношение иерархического вложения: геном — это совокупность репликанов, репликон — последовательность скриптов, скриптон, в свою очередь, разбивается на цистроны, каждый из которых уже кодирует определенный белок. Текст, отвечающий цистрону, разбивается при этом на последовательность триплетов (кодонов), кодирующих в соответствии с генетическим кодом определенные аминокислоты. Отдельные цистроны могут располагаться последовательно, не пересекаясь друг с другом, могут частично пересекаться, могут быть даже целиком вложенными один в другой. Различие в синтезируемых ими белках достигается в последнем случае за счет использования отличающихся по фазе триплетных рамок считывания, либо (при совпадающих фазах) за счет несовпадения длин кодирующих последовательностей.

Приведенное выше описание структуры генома весьма схематично, но достаточно для наших целей. Следует отметить лишь, что наряду с указанными функциональными единицами, назначение которых в той или иной степени выяснено, существуют некоторые другие, функциональное назначение которых пока непонятно (нетранскрибируемые участки ДНК, межцисстроинные интервалы, интроны). Исследование различных характеристик этих участков (в том числе и 1-граммных) и сопоставление их с аналогичными характеристиками известных функциональных единиц может дать толчок к выявлению их функционального назначения.

2. Примеры содержательных задач. Укажем несколько задач из области анализа нуклеотидных последовательностей, при решении которых существенную пользу может принести знание 1-граммных характеристик этих последовательностей.

Следует сразу заметить, что 1-граммные характеристики раскодировываемых нуклеотидных текстов (или отдельных участков этих текстов) всегда интересовали исследователей. Однако интерес этот в подавляющем большинстве случаев был односторонним. А именно анализировались лишь участки геномов, соответствующие цистронам (т.е. участки, кодирующие белки), причем анализ последних ограничивался преимущественно подсчетом частот использования различных кодонов (т.е. триграмм) и исследованием распределения нуклеотидов по различным позициям кодонов.

Интерес к характеристикам третьего порядка обусловлен, естественно, триплетной структурой генетического кода. Нет никаких оснований предполагать, что они будут играть основную роль и при исследовании некодирующих частей геномов. Ниже будет показано, что важная информация содержится не только в характеристиках третьего порядка, но и в характеристиках меньшего ($l = 1, 2$) и большего ($l = 4, 5, \dots, l_{\max}$) порядков. Параметр $l_{\max}$ здесь соответствует тому минимальному значению l , начиная с которого в тексте уже отсутствуют повторяющиеся 1-граммы, т.е. для любой 1-граммы текста x частота ее встречаемости $F_1(x) = 1$ при $l_{\max} \leq l \leq N$.

2.1. Задача разметки текста сводится к построению алгоритмов автоматического выделения функциональных единиц геномов (см. п. I) и соответствующих им знаков пунктуации (инициаторов и терминаторов). Эта задача разбивается на ряд отдельных подзадач, многие из которых в формальной постановке сводятся к задаче распознавания.

Укажем для примера на задачу автоматического выделения промоторов - знаков пунктуации, ответственных за начало процесса транскрипции. Каждый промотор представляет собой последовательность длиной до 50 символов. В настоящее время расшифровано уже несколько десятков промоторов, т.е. установлена их первичная структура и доказано (молекулярно-биологическим анализом), что соответствующие последовательности нуклеотидов выполняют функции инициаторов транскрипции. Совокупность этих последовательностей может рассматриваться как обучающая выборка по промоторам в задаче классификации "промотор-не промотор". Для элементов этой выборки существует реальное опознающее устройство - РНК-полимераза. Обучающая выборка по "не промоторам" может быть получена из текстов уже расшифрованных геномов с исключенными промоторными зонами.

Целью обучения является построение решающего правила, моделирующего принцип действия РНК-полимеразы, т.е. позволяющего отнести произвольную последовательность к одному из двух указанных классов. Результат классификации может быть подтвержден (либо опровергнут) весьма трудоемким молекулярно-биологическим анализом. Удачное решение задачи классификации формальными методами (т.е. прямо по виду символьной последовательности) может существенно упростить и ускорить процедуру выявления генетической структуры генома.

В основу построения решающего правила кладется различие в характеристиках элементов первого и второго образов, выявляемое при анализе обучающих выборок. Нами исследовалась возможность использования для этой цели 1-граммных характеристик. А именно была выдвинута и экспериментально подтверждена гипотеза о том, что различие между промоторами и "не промоторами" наиболее ярко проявляется на уровне длинных 1-грамм ($l \geq 7$). Было показано, что существуют достаточно длинные 1-граммы, типичные только для промоторов, т.е. 1-граммы, каждая из которых встречается как минимум в двух промоторах, но не встречается у "не промоторов".

Примерами таких 1-грамм являются: TTTTGAЦА (в промоторах λ -PL, E.COLI-TRP, Φ X174D, S.TUPH-TRP), АЦАЦТТТ (в промоторах E.COLI-LAC, E.COLI-GAL, TRNK-TRP, E.COLIK12-ARAC), ГЦГТГТАТА (в промоторах λ -PL, λ -PR, λ -PRM) и ряд других. Анализ позиций, занимаемых этими 1-граммами внутри промоторов, подтверждает, что совпадение 1-грамм у разных промоторов носит случайный характер. Каждой из указанных 1-грамм внутри разных промоторов соот-

ответствуют одинаковые или очень близкие позиции. Таким образом, сочетание 1-граммного и позиционного анализа позволяет в данном случае выявить функционально важные зоны внутри промоторов и установить возможные варианты нуклеотидных последовательностей для этих зон.

Выделенное множество характерных (в упомянутом выше смысле) 1-грамм естественно назвать "покрытием" обучающей выборки промоторов в случае, если каждый промотор из этой выборки содержит хотя бы одну характерную 1-грамму. В качестве примера укажем, что для обучающей выборки из 24 промоторов удалось получить покрытие из 10 характерных 1-грамм ($6 \leq l \leq 10$). Элементы покрытия в сочетании с информацией о наиболее вероятных зонах их расположения внутри промоторов представляют основу для выработки различных стратегий распознавания.

2.2. Задача вычисления эволюционного расстояния между геномами является составным элементом более общей задачи построения филогенетических деревьев для различных семейств геномов. Можно выделить следующие элементарные преобразования, посредством которых осуществляется эволюция геномов:

- 1) замена одного основания другим;
- 2) вставка цепочки нуклеотидов в исходную последовательность;
- 3) удаление цепочки нуклеотидов из исходной последовательности;
- 4) инверсия (замена цепочки нуклеотидов в исходной последовательности элементами комплементарной цепочки, прочитанными в обратном порядке). Например, цепочка ААТЦГТТЦ исходной последовательности будет заменена инвертированной ГАЦЦГАТТ, которая получается из комплементарной к исходной (ТТАГЦЦАГ) прочтением ее справа налево;
- 5) дупликация (повторение цепочки символов);
- 6) транспозиция (перенос цепочки символов из одного места последовательности в другое).

Естественно определить меру различия между двумя геномами как минимальное число осмысленных с биологической точки зрения преобразований типа 1-6, переводящих один геном в другой. Вычисление этой меры представляет нетривиальную комбинаторную задачу, пока нерешенную в общем виде. Для частного случая, когда допустимыми являются лишь преобразования 1-3, причем два последних при -

менным только к цепочкам длины 1, расстояние вычисляется с помощью метода динамического программирования. При этом существует однозначная связь между расстоянием и длиной максимально длинной общей подпоследовательности двух последовательностей (геномов) [2].

Отметим, что преобразования 2-6 несут "групповой" характер, т.е. применимы к цепочкам, состоящим из многих (иногда до нескольких десятков) символов. Для вычисления расстояния необходимо уметь отыскивать в текстах участки, соответствующие вставкам (устранениям), инверсиям, дупликациям и транспозициям. Основу для этого дает вычисление и сопоставление полных частотных спектров обоих текстов.

Так, например, различие в частотах встречаемости какой-либо достаточно длинной*) 1-граммы в двух родственных текстах, как правило, свидетельствует о том, что в процессе эволюции имела место дупликация данной 1-граммы. Сопоставление результатов частотного и позиционного анализов зачастую позволяет выявить дупликацию даже при наличии одного текста. Об этом сигнализирует близость (а при многократной дупликации - периодичность) расположения повторяющихся 1-грамм.

Для выявления инверсий необходимо сравнить частотные спектры двух текстов: любого из исходных и инвертированного оставшегося. Наличие достаточно длинных (например, по отношению к схеме независимого порождения) совпадающих 1-грамм в обоих текстах говорит о возможности инверсий. Если в результате позиционного анализа выявляется, что номера позиций, занимаемых интересующей нас 1-граммой, в обоих исходных текстах совпадают или близки друг к другу, то предположение о наличии инверсии подтверждается, в противном случае возможна транспозиция с инверсией, либо случайное совпадение.

Соображения, аналогичные вышеизложенным, могут быть положены и в основу отыскания транспозиций. О наличии вставок (устранений) сигнализирует различие в длинах анализируемых текстов. Для локализации вставок целесообразно осуществить предварительную синхронизацию обоих текстов по выявленным характерным участкам, в качестве которых могут выступать знаки пунктуации, совпадающие 1-грам-

*) При малых значениях l частоты одинаковых 1-грамм в родственных геномах будут отличаться почти всегда, поскольку степень искажения частотных характеристик под влиянием любого из преобразований 1-6 обычно тем больше, чем меньше l .

мы, инверсии, дубликации. Различие длин последовательностей, заключенных между идентичными знаками синхронизации в обоих геномах, свидетельствует о наличии вставки (устранения).

2.3. Восстановление текста по фрагментам. Определение первичной структуры длинных последовательностей (например, целых геномов) в настоящее время технически трудно осуществимо. С помощью существующих методик удается восстанавливать первичную структуру лишь относительно коротких участков (~20-300 нуклеотидов). Поэтому для определения первичной структуры всего генома его предварительно разрезают на более короткие участки (фрагменты) специальными ферментами-рестриктазами, определяют первичную структуру каждого участка, а затем, соединяя нужным образом расшифрованные куски, восстанавливают последовательность, соответствующую геному в целом.

Определение порядка следования фрагментов является нетривиальной задачей. Для восстановления текста используют наборы фрагментов, соответствующие разным рестриктазам. Поскольку каждая рестриктаза разрезает геном лишь в характерных для нее участках (например, рестриктаза HaeIII реагирует только на сочетание ГГЦЦ, разрывая связи между Г и Ц), фрагменты, получаемые от разных рестриктаз, оказываются перекрывающимися. Это свойство и используется для упорядочения фрагментов.

С формальной точки зрения задача восстановления текста по набору фрагментов, полностью его покрывающих, сводится к исследованию вопроса о существовании единственного (с учетом перекрываемости) варианта упорядочения фрагментов, полученных при фиксации определенного набора стратегий расчленения текста (набора рестриктаз). Предложено несколько способов сведения данной задачи к задаче отыскания эйлеровых цепей на графе. Структура соответствующего графа определяется типом текста и используемым набором рестриктаз.

Поскольку с каждой рестриктазой ассоциируется определенная 1-грамма, то количество таких 1-грамм в тексте и их расположение будут определять соответствующий набор фрагментов. Зная 1-граммные и позиционные характеристики родственных текстов, можно целенаправленно формировать набор рестриктаз, так чтобы добавление каждой новой рестриктазы в максимальной степени уменьшало неоднозначность, возникающую при восстановлении текста по уже имеющемуся набору фрагментов.

3. Частотные характеристики генетических текстов. В табл.1 приведены частотные характеристики текстов $\Phi X174, G4, FD, SV40, PBR322$ и MS2 для значений $l = 1, 2, 3$. Все l -граммы упорядочены по убыванию частоты встречаемости их в тексте. Порядковый номер r каждой l -граммы будем называть ее рангом. Параметр r изменяется в диапазоне от 1 до M_l , где M_l - количество различных l -грамм в тексте.

Поскольку с увеличением l объем таблиц возрастает очень быстро (при малых l , как правило, $M_l = n^l$; при $l \rightarrow l_{\max}$ значение M_l ограничено величиной $M - l + 1$), для значений $l > 3$ выборочно приведены лишь интегральные частотные характеристики E_l^k и S_l (табл.2). Здесь E_l^k - количество различных l -грамм текста, каждая из которых встречается ровно k раз; $S_l = \sum_{i=1}^n k(i) \cdot E_l^{k(i)}$ -

текущая сумма числа l -грамм, соответствующих n первым частотным градациям, упорядоченным по убыванию.

Наиболее длинные повторяющиеся l -граммы, встретившиеся в каждом из текстов, приведены в табл. 3.

Табл. 4 содержит энтропийные характеристики текстов [4] (значения $\hat{H}_l$ и $\hat{H}_l$ для $l = 2, \dots, 6$). Ввиду ограниченности длины текстов энтропийные оценки с увеличением l становятся статистически недостоверными. Формально это выражается в том, что с ростом l величины $\hat{H}_l \rightarrow 0$, а $\hat{H}_l \rightarrow 1$. Ввиду отсутствия эффективных подходов к определению порогового значения $l(N)$, определяющего статистически достоверные результаты от недостоверных, мы можем воспользоваться для этой цели лишь некоторыми косвенными соображениями. Не обсуждая их детально, укажем, что в данном случае можно, по-видимому, в качестве оценки для $l(N)$ взять значение $\hat{l}(N) = 4$.

В табл.6 приведены величины коэффициентов ранговой корреляции Спирмена между каждой парой текстов для значений $l = 2, 3, 4$. При больших значениях l уже не все из 4^l возможных комбинаций нуклеотидов присутствуют в тексте, поэтому использовать данную меру сходства без соответствующей модификации нецелесообразно. Нуклеотиды Т (в ДНК-содержащих геномах) и У (в РНК-содержащих) при вычислении ранговых мер сходства отождествлялись.

4. Сравнительный анализ l -граммных характеристик генетических текстов. Ниже обсуждаются наиболее интересные, на наш взгляд, закономерности, выявленные при l -граммном анализе текстов. Далеко

не для всех из них удастся предложить соответствующую интерпретацию. И даже там, где это делается, интерпретация, как правило, носит характер гипотезы. Укажем вначале на два интересных свойства, вытекающих из анализа характеристик первого и второго порядков.

4.1. Классификация по Т, А- или Ц, Г-преобладанию. Сопоставление характеристик первого порядка показывает, что все тексты можно разделить на две группы в зависимости от ранжировки нуклеотидов по частоте встречаемости. В первую группу входят тексты ФХ174, G4, FD и SV40, в которых преобладают нуклеотиды Т и А (Т, А - богатые тексты), во вторую - PBR322 и MS2 (Ц, Г - богатые). Промежуточные варианты, когда Т и А перемежаются с Ц и Г при частотном упорядочении, отсутствуют.

Данная классификация, по-видимому, характеризует вторичную структуру молекул РНК, кодируемых этими геномами, т.е. их пространственную конфигурацию. Известно, что бактериофаг MS2 в отличие от объектов первой группы имеет сложную вторичную структуру. Ее отличительным признаком является наличие большого количества длинных "шпильек" - участков самокомплементарности одноцепочечной молекулы MS2. В [3] показано, что шпильки формируются преимущественно из кодонов, содержащих нуклеотиды Ц и Г в третьей позиции. Поскольку комплементарная пара Ц $\equiv$ Г обладает тремя водородными связями, а пара А = У лишь двумя, преимущественное использование Ц и Г в третьей позиции (а только здесь и допустимы вариации) должно обеспечить большую устойчивость шпильек.

Возможность подтверждения либо опровержения предложенной интерпретации содержится в анализе пока еще не установленной вторичной структуры РНК, транскрибируемых с генома плазмиды PBR322. Он существенно отличается от генома MS2 тем, что является ДНК-содержащим и двухцепочечным, аналогично объектам из первой группы. Если его РНК обладают слабо выраженной вторичной структурой, то гипотеза о связи ранжировки нуклеотидов в частотной характеристике первого порядка со сложностью и стабильностью вторичной структуры соответствующих молекул отпадает. В противном случае эта гипотеза получает сильное подтверждение.

4.2. ЦГ - эффект у SV40. Анализ частотных характеристик второго порядка (табл. I) обнаруживает очень сильный аномальный эффект в частоте встречаемости биграммы ЦГ у вируса SV40. Эта биграмма встречается в тексте всего лишь 27 раз ($\tau = 16$). Для

Т а б л и ц а 1

Частотные характеристики генерических токов для значений
1 = 1, 2, 3

l	x	№ 174		G4		FD		SV40		PBR322		№ 2	
		1-группа	$P_1(x)$	1-группа	$P_1(x)$	1-группа	$P_1(x)$	1-группа	$P_1(x)$	1-группа	$P_1(x)$	1-группа	$P_1(x)$
1	1	Г	1684	А	1519	Г	2210	Г	1582	Ц	1208	Ц	932
	2	А	1291	Г	1510	А	1577	А	1518	Г	1134	Г	928
	3	Г	1254	Ц	1446	Г	1321	Ц	1094	Г	1036	У	875
	4	Ц	1157	Г	1102	Ц	1299	Г	1033	А	984	А	834
2	1	ГГ	572	АА	554	ГГ	815	ГГ	575	ГЦ	378	ЦГ	266
	2	ГГ	480	ЦГ	503	АА	544	АА	533	ЦГ	329	УЦ	263
	3	ЦГ	404	ГЦ	443	АГ	515	ЦА	422	ГЦ	306	ГГ	250
	4	АА	395	АГ	392	ЦГ	483	ЦГ	398	ЦА	300	ЦУ	242
	5	АГ	383	ГГ	391	ГГ	481	ГГ	389	ЦЦ	292	ГЦ	231
	6	ГЦ	347	АЦ	386	ТА	469	АГ	352	ГГ	287	ГГ	230
	7	ГА	327	ЦА	352	ГЦ	445	АГ	345	ГГ	285	ЦЦ	223
	8	ГГ	325	ГГ	342	ГГ	397	ТА	326	ГГ	281	АА	222
	9	ГЦ	320	ТА	334	ГГ	322	ГЦ	292	ГГ	259	ГА	217
	10	ТА	312	ГЦ	313	ГЦ	316	АЦ	288	АА	258	АЦ	215
	11	ЦГ	267	ЦЦ	304	ГА	286	ГГ	272	АГ	254	УГ	211
	12	АЦ	260	ЦГ	287	ЦА	278	ГЦ	266	АГ	239	УГ	207
	13	ЦА	257	ГА	279	ЦГ	274	ГГ	257	ГА	236	АГ	204
	14	ГГ	255	ГГ	272	АЦ	274	ЦЦ	247	ГГ	235	ЦА	201

15	16	AT	252	IT	258	III	264	TA	237	AI	232	YA	194
16		III	229	AT	186	AT	244	IT	27	TA	190	AY	192
1		TTT	178	AAA	197	TTT	273	TTT	238	TTU	112	YUT	89
2		TUT	149	UAA	157	ATT	223	AAA	211	TUT	105	UTU	80
3		TTT	142	TUT	152	AAA	201	UTT	134	TTU	101	UTY	75
4		ATT	140	UTU	144	TAA	188	UAA	132	TUT	99	UTY	74
5		ITT	140	AAU	135	TTU	186	AUA	129	UTT	95	UTA	73
6		TAA	139	AAT	135	AAT	178	ATT	129	UAT	93	UTY	70
7		AAA	133	UTT	132	TTT	177	TAA	122	TUA	91	YTT	69
8		UTT	132	TUT	129	TTT	168	UAT	112	ATU	90	YUU	68
9		TTU	132	AUT	126	TTT	164	UTT	111	TUA	90	YUU	68
10		TTU	130	TUA	119	TAA	164	AAT	111	UTT	89	ITT	64
11		TTT	130	TAT	119	TAT	158	TTU	109	UTT	88	UTT	63
12		UTT	127	ATT	118	TAA	154	TAA	108	AAA	88	TTU	63
13		ATT	127	ATU	118	UTT	142	TTT	106	TUU	87	UTY	63
14		TAA	120	UTT	116	ITT	136	AAT	106	TUU	84	YAU	63
15		TAT	115	TTU	116	UTT	129	AAU	105	UAT	84	UAA	62
16		TUT	104	UTA	111	TUT	122	UAT	105	UTU	80	UTY	62
17		AAT	103	TTU	107	UTU	121	TUU	104	TUT	80	AUT	62
18		TAT	101	TTT	106	TUT	119	TTT	103	ATU	78	UTY	61
19		TAA	101	TAA	103	TTT	118	TUA	102	UTU	78	AUU	61
20		TTT	100	TAA	101	TUA	113	TUT	102	AUU	76	UTU	61
21		AAT	92	TAA	99	ATA	111	AUT	102	TTU	76	UTT	60

Продолжение таблицы 1

1	2	3434		64		7D		8740		PBE322		MS2	
		1-группа	$P_1(x)$	1-группа	$P_1(x)$	1-группа	$P_1(x)$	1-группа	$P_1(x)$	1-группа	$P_1(x)$	1-группа	$P_1(x)$
	22	КТН	92	ТАЦ	98	ТАГ	110	ЦПГ	102	ТТТ	75	АУЦ	59
	23	ЦПГ	92	ТТА	98	ЦПГ	109	АУЦ	97	АТГ	73	ААА	58
	24	ТТА	91	АУА	96	АТГ	108	ЦПА	97	ЦПА	71	ЦУУ	57
	25	ААА	91	ЦПТ	96	ЦПА	103	ГТТ	97	ГТТ	70	ТАГ	57
	26	АПГ	89	ЦПТ	96	ЦПА	91	АТГ	95	ТТТ	69	УПГ	56
	27	ТАЦ	82	ЦПГ	92	АПГ	91	ЦПА	94	ЦПА	68	АЩ	56
	28	ТТТ	81	ЦПЦ	92	ТЦН	90	ГПА	94	ЦПТ	67	ГУЦ	56
	29	ЦПА	78	ТАЦ	90	ТАЦ	89	ТЦГ	92	ГТА	67	АГА	55
	30	ТАГ	74	АТГ	90	ТЦН	87	ТТТ	92	ЦПГ	67	АУЦ	55
	31	АТГ	74	АЩ	89	ЦПТ	86	ТТА	90	ТЦЦ	66	ГЦГ	55
	32	АГА	73	ААГ	86	ААН	84	ТАГ	87	ГАА	66	ТАЦ	54
	33	ТАЦ	72	ЦАТ	83	ГУЦ	83	АТГ	87	ЦЩ	66	ААГ	54
	34	ГУЦ	72	ТЦН	80	ААГ	81	АТГ	83	ТТА	65	УУУ	54
	35	ГЩ	71	ЦПГ	77	ГТА	79	ГАА	82	ААЦ	64	ГАА	54
	36	ЩПГ	70	АПГ	75	ЦПЦ	78	ЦПЦ	80	ТТТ	64	ЦПА	54
	37	ГАА	70	ГТТ	74	ТЦГ	78	ГТА	79	ТАТ	63	ГПА	54
	38	КПН	67	ГУЦ	74	АТГ	77	АТА	74	ГТТ	63	УПЦ	53
	39	ГЦГ	67	ГТА	73	ГТТ	76	ТТТ	73	ТТГ	62	АУГ	53
	40	ТЦГ	67	ГЩ	72	ГАА	76	ЦАЦ	73	ТАТ	61	АТГ	53
	41	ААН	66	ТТТ	71	АТН	73	ЦПА	73	АТГ	61	ТАУ	52

[illegible]

сравнения укажем, что частоты встречаемости биграмм с рангом 16 во всех остальных текстах в 7-9 раз выше.

Отметим, что исследователи, занимавшиеся изучением кодирующих частей генома SV40 и родственных вирусов, уже обращали внимание на чрезвычайно малое количество кодонов, содержащих биграмму ЦГ в позициях 1-2 или 2-3. Наш анализ показывает, что данный эффект возникает уже на уровне биграмм, т.е. аномалии в использовании кодонов - его следствие, а не причина. Эффект наблюдается для генома в целом, как для кодирующих частей (при любых фазах считывания), так и для не кодирующих. Выявление данного и ему подобных эффектов не требует предварительной разметки текстов и может быть автоматизировано с помощью достаточно развитой в настоящее время техники обнаружения аномальных наблюдений.

Мы не можем пока указать, какую функциональную нагрузку несет данный эффект. В [6] предпринята попытка объяснения механизма эффекта тем, что в геномах вирусов эукариотов существует механизм перехода ЦГ в ТГ (5-метил-цитозин часто мутирует в Т, причем Ц метилируется именно в составе биграммы ЦГ). Этот процесс должен был бы приводить к соответствующему избытку ТГ. Проверка данной гипотезы на схеме независимого порождения (см. п.4.5) показала, однако, что некоторый избыток ТГ имеет место, но он далеко не компенсирует недостатка ЦГ.

Отметим, что из шести анализировавшихся микроорганизмов SV40 является единственным, который паразитирует на эукариотах, остальные - на прокариотах. Если данный эффект подтвердится и на других вирусах, паразитирующих на эукариотах, он может быть положен в основу алгоритма классификации текстов по признаку: эукариотический-прокариотический.

4.3. Наиболее длинные повторения. Значения параметра $l_{\max}$ составляют: для текста SV40 - 59, FD - 33, G4 и FBR322 - 14, $\Phi X174$ - 13, MS2 - 11. При $l = l_{\max} - 1$ в каждом из текстов уже появляются повторяющиеся 1-граммы (см. табл.3). Анализ их расположения в текстах проливает некоторый свет на их происхождение.

Самый длинный повтор ($l = 58$) наблюдается в тексте SV40 (1-грамма #14 из табл.3). Эта 1-грамма начинается и заканчивается триграммой ТТГ. В тексте обе 58-граммы расположены подряд так, что последние три символа первой 1-граммы накладываются на первые три символа второй. Таким образом, наличие данного повтора можно

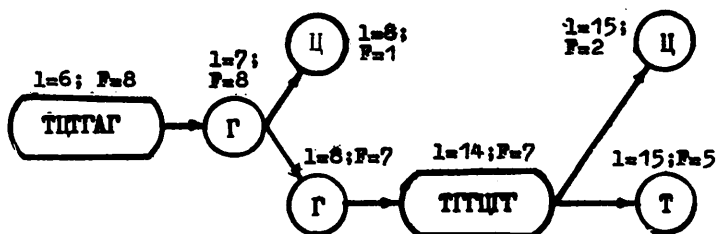
Т а б л и ц а 3

Наиболее длинные повторяющиеся 1-граммы
генетических текстов

Текст	1	F ₁	1-грамма	№ 1-грамм
ΦX174	I2	2	ЦТЦААГТАЦТГ	1
	I2	2	ЦТЦТГЦЦЦТТТ	2
	II	2	ТТЦТГТТГАТТ	3
	II	2	ЦТЦТТТГААГТ	4
G4	I3	2	ГТГТЦЦЦГАТТЦ	5
	II	2	ГЦГТТТАТГТЦ	6
	II	2	ЦГАТЦТЦГГЦ	7
	II	2	АЦЦЦЦТТАААГ	8
FD	32	2	ГГЦЦЦТГАГГТГГГЦГГГТЦТГАГГТГГГЦТ	9
	32	2	ТЦТГАГГГТГГГЦГГГТЦТГАГГТГГГЦГГТЦ	10
	20	2	ГЦГТТТЦТГАГГТГГГЦТТ	11
	I3	2	ААГТТААТЦААА	12
	I2	2	ЦТЦААТТТЦТТ	13
SV40	58	2	ТГГТГГЦГГАНТААТТГАГАТТНАТГЦТТГЦ АТАНТТГГГЦЦТГГГГТТГАГЦЦГТ	14
	2I	2	ТГГГЦГГАГТТАГТГГГЦГТА	15
	I3	2	ГАГТАЦАЦАГАГТ	16
	I3	2	АААААЦЦАГАААГ	17
	II	2	ЦТЦААТТТГТТ	18
PBR322	I3	2	АГЦААААГГЦЦАГ	19
	II	2	ТГГЦГГГЦТТ	20
	II	2	ЦГГТТЦЦЦЦЦ	21
	II	2	АААААГГАТЦТ	22
MS2	10	2	АГАЦГЦЦГЦ	23
	10	2	АУУЦЦЦЦАГ	24
	8	4	АТТГГЦЦ	25
	8	3	ГУУУУАА	26

было бы объяснить дубликацией соответствующей 55-граммы без наложения (но при этом маловероятно появление еще одного ТТГ после такой композиции) либо дубликацией всей 58-граммы с наложением (одним из возможных объяснений механизма наложения может быть "неравный" кроссинговер). Наличие данного повтора уже отмечалось. Функциональное назначение его пока не выяснено, известно лишь, что в этом районе инициируется репликация генома SV40. С интервалом в три символа от данного участка находится еще одна дубликация (дважды повторяется 21-грамма №15 из табл. 3).

Интересный пример дубликации наблюдается в тексте FD. Ядро этой дубликации составляет 14-грамма X = ТЦТАГТТТТЦТ. Она легко выделяется при анализе дерева всевозможных продолжений 6-граммы ТЦТАГ:



Из приведенного рисунка видно, что при изменении l от 8 до 13 дерево не имеет разветвлений, т.е. продолжается формирование какого-то сообщения (с неясной пока семантикой). При $l = 14$ формирование заканчивается, и далее дерево расщепляется. Три 15-граммы, заканчивающиеся на Т, следуют в тексте друг за другом, образуя тройную дубликацию, которой предшествует триграмма ГТЦ(ГТЦТХТХТ). На расстоянии порядка 400 символов от этого участка расположен другой, образованный четырехкратным (с интервалом в 1 символ) повторением базовой 1-граммы и заканчивающийся триграммой ТТЦ(ХЦТХЦТТЦ). Обладая внутренней периодичностью, два этих участка содержат в себе и более длинные (по сравнению с базовой) повторяющиеся последовательности (1-граммы № 9, 10 и 11 из табл. 3).

Интересно отметить, что расположение базовых 1-грамм синхронизовано относительно рамки считывания при трансляции так, что каждая из них кодирует идентичные фрагменты белка. Это указывает на то, что в данном случае, в отличие от предыдущего, семантическую нагрузку дубликаций следует искать на уровне функционирования кодируемого ими полипептида.

Примером дубликаций являются также 1-граммы №5 (самый длинный повтор в G4) и №13 (PBR322). Интересно отметить, что последняя 1-грамма, аналогично 58-грамме из SV40, имеет конец, совпадающий с началом (биграмма АГ). Механизм дубликации вновь улавливает это и осуществляет дубликацию с наложением. Более короткие 1-граммы, например №1 (из ФХ174), №23 (из MS2), встречающиеся каждая по два раза, уже разнесены внутри своих текстов, т.е. наличие таких повторов может трактоваться и как случайный фактор.

Т а б л и ц а 4
Оценки энтропии и избыточности
генетических текстов

Текст	Энтропия H_1 R_1	1				
		2	3	4	5	6
ФХ174	$\hat{H}_1$	1,968	1,932	1,900	1,798	1,453
	$\hat{R}_1$	0,015	0,033	0,049	0,100	0,273
G4	$\hat{H}_1$	1,963	1,937	1,903	1,815	1,479
	$\hat{R}_1$	0,018	0,031	0,048	0,092	0,260
FD	$\hat{H}_1$	1,963	1,928	1,890	1,798	1,493
	$\hat{R}_1$	0,018	0,036	0,055	0,101	0,254
SV40	$\hat{H}_1$	1,906	1,980	1,862	1,810	1,451
	$\hat{R}_1$	0,047	0,055	0,069	0,095	0,274
PBR322	$\hat{H}_1$	1,989	1,974	1,944	1,855	1,389
	$\hat{R}_1$	0,006	0,013	0,028	0,073	0,305
MS2	$\hat{H}_1$	1,997	1,985	1,956	1,823	1,284
	$\hat{R}_1$	0,001	0,007	0,021	0,088	0,357

4.4. Низкая избыточность. С учетом замечания о достоверности энтропийных оценок, сделанного в п.3, избыточность генетических текстов, как это следует из табл.4, очень невелика (не превышает 0,1). Для сравнения укажем, что оценки избыточности современных естественных языков (русского, английского, немецкого) дают значение $R \approx 0,7$.

Отсутствие избыточности в языке означает, что любая комбинация символов в нем является осмысленной. Именно по такому принципу построен генетический код, использующий все 64 возможные комбинации троек нуклеотидов. Анализ некодирующих частей показывает, что и там на уровне $1 = 3$ нет запрещенных комбинаций.

В порядке нарастания избыточности тексты ранжируются следующим образом: MS2, PBR322, G4 и Φ X174, λ D, SV40. Приведенная ранжировка, возможно, отражает многообразие и степень перекрытия различных функциональных единиц в геномах (в некотором смысле сложность организации геномов). У генома MS2, избыточность которого минимальна, гены практически не перекрываются друг с другом. В геномах G4, Φ X174, SV40 имеются зоны перекрытия, иногда весьма обширные. Нахождение различных функциональных единиц (гена на ген, промотора на ген и т.д.) увеличивает количество ограничений, которым должна удовлетворять символьная последовательность, т.е. повышает избыточность языка. Аналогии подобного рода проводил К.Шеннон [4], обсуждая связь между избыточностью языка и возможностью построения кроссвордов.

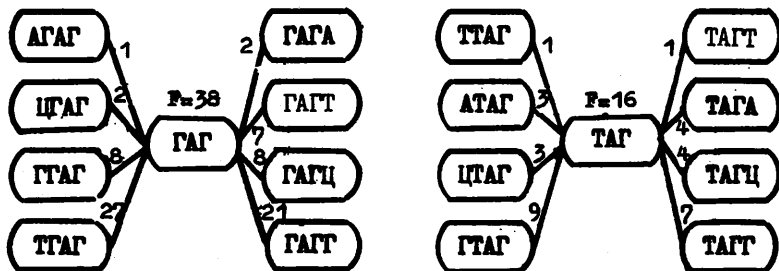
Низкая избыточность находит свое отражение в слабом наклоне частотных характеристик соответствующего порядка. Наклон характеристики 1-го порядка может быть определен как МНК - оценка параметра γ_1 в выражении $F_1(r) = C_1 r^{-\gamma_1}$, используемом для аппроксимации частотной кривой. Значения оценок этого параметра довольно сильно отличаются для текстов двух выделенных групп, но не превышают $0,5 (1 = 2, 3)$. Это говорит о том, что закон Ципфа (близость параметра γ к единице для частотных словарей естественного языка) не выполняется для 1-граммных характеристик генетических текстов. Более того, различия в значениях этого параметра могут быть положены в основу классификации генетических текстов.

4.5. Схема независимого порождения как грубая модель текста. Низкие значения избыточности R для $1 = 2, 3, 4$ наводят на мысль о том, что, по крайней мере, в данном диапазоне значений 1 можно воспользоваться в качестве грубой модели текста схемой независимых испытаний с вероятностями порождения каждого символа α равными $F(\alpha)/N$, где $\alpha \in \{A, T, G, C\}$, а $F(\alpha)$ - частота встречаемости символа α в соответствующем тексте. Основанием для этого служит тот факт, что для последовательности, полученной в соответствии со схемой независимых испытаний, $N_1 = N_2 = \dots = N_l$ для любого l . Анализ табл.4 показывает, что для $1 = 2, 3, 4$ мы делаем не слишком большую натяж-

ку, допуская эту гипотезу. В наилучшем соответствии с ней находятся тексты с наименьшей избыточностью (MS2, PBR322). Это подтверждается и прямым сопоставлением частот 1-грамм, наблюдаемых фактически и вычисляемых по схеме независимых испытаний.

Схема независимого порождения представляет интерес с точки зрения оценки многих интересующих нас параметров, например, таких как $l_{\max}, M_1, E_1^1$, определяющих трудоемкость алгоритмов обработки символьных последовательностей. С биологической же точки зрения несомненный интерес представляет анализ очень сильных отклонений фактических частот от модельных. Зачастую такие отклонения являются функционально значимыми.

Рассмотренный выше ЦГ-эффект у SV40 можно трактовать, с одной стороны, как аномалию в частотной характеристике, с другой стороны, — как сильное отклонение от схемы независимых испытаний ($F_{\text{мод}}(\text{ЦГ}) = 217$, $F_{\text{эксп}}(\text{ЦГ}) = 27$). Пример дубликации в тексте FD, описанный в п.4.3, можно также трактовать как аномальное отклонение от схемы независимых испытаний, поскольку при увеличении 1 от 8 до 14 частота соответствующей 1-граммы не должна была оставаться постоянной, а должна была уменьшаться на каждом шаге за счет разветвления дерева на 4 ветви с вероятностями $P(\alpha)$ ($\alpha \in \{A, T, G, C\}$) для каждой ветви. Другой пример аномального ветвления — лево- и правосторонние расширения 1-грамм ГАГ (из FD) и ТАГ (из G4):



В табл. 5 приведен еще ряд примеров аномальных отклонений экспериментально наблюдаемых частот ($F_{\text{эксп}}$) от вычисленных в соответствии с моделью независимого порождения ($F_{\text{мод}}$). Отметим, что во всех текстах фактически наблюдаемые частоты встречаемости 1-грамм АА, ААА выше, а триграммы ТАГ (терминальный кодон) существенно ниже, чем это следует из схемы независимых испытаний.

Т а б л и ц а 5

АНОМАЛЬНЫЕ ОТКЛОНЕНИЯ ОТ СХЕМЫ НЕЗАВИСИМОГО ПОРОЖДЕНИЯ

Текст	1-граммы	F _{мод}	F _{эксп}	1-граммы	F _{мод}	F _{эксп}	Комментарий
ΦX174	AA TA TT	310 403 392	395 312 480	AAA ATT CTT TTT TAG	74 94 63 68 94	I33 I40 3I 20 24	иниц. терм.
G4	AG AA CT	300 413 391	I86 554 503	AAA AGT ATA TAG TTT	II2 8I II2 8I III	I97 36 66 16 66	 терм.
FD	AA AG	388 325	544 244	AAA CAT TAG AGT GAG	95 II0 II2 II2 67	20I 69 58 53 38	 терм.
SV40	CT TA CA	217 459 318	27 325 424	Все CT-содержащие триплеты			
				AAA ATA CAT CTT TTT TTT	I28 I33 63 65 40 I44	2I2 74 II3 I32 73 238	
PBR322				AAA CCC ATC TAG	50 92 65 61	88 66 90 31	 терм.
MS2	UC UC	228 263	263 263	CAC UCC	57 59	37 89	

Поскольку схема независимых испытаний является весьма грубым приближением к описанию генетических текстов, предпринимается попытка использовать для описания отдельных частей геномов более сложные модели, например, в виде марковских цепей низшего (первого и второго) порядка [5]. Оценки переходных вероятностей получаются при этом на основе 1-граммных характеристик соответствующих частей геномов. Такого рода кусочные представления могут быть использованы для решения задач разметки и генерации текстов.

Т а б л и ц а 6

Коэффициенты ранговой корреляции

Пара текст/текст	1		
	2	3	4
ΦX174/SV40	0,55	0,43	0,37
ΦX174/G4	0,57	0,51	0,58
ΦX174/MS2	-0,08	0,01	
ΦX174/FD	0,86	0,77	0,67
ΦX174/FBR	0,10	0,13	0,21
SV40/G4	0,49	0,41	0,34
SV40/MS2	-0,30	-0,32	-0,15
SV40/FD	0,67	0,55	0,48
SV40/FBR	0,22	-0,04	-0,03
G4/MS2	-0,11	0,14	0,11
G4/FD	0,69	0,56	0,50
G4/FBR	0,24	0,15	0,18
MS2/FD	0,02	0,01	0,10
MS2/FBR	0,62	0,32	0,19
FD/FBR	0,04	-0,02	0,01

4.6. Количественные оценки близости текстов. Простейшими мерами близости текстов являются коэффициенты ранговой корреляции ρ_1 , вычисленные с использованием частотных характеристик 1-го порядка (см. табл. 6, $l = 2, 3, 4$). Значения ρ_1 в целом согласуются с тем разбиением текстов на 2 класса (по ТА- и ЦГ-пробладанию), которое было приведено в п. 4.1, поскольку коэффициенты ранговой корреляции между представителями разных классов близки к нулю или отрицательны. Представляет интерес проследить степень близости между текстами одного класса.

Весьма близки тексты ΦX174 и G4 ($\rho_1 > 0,5$ для всех l), что неудивительно, так как эти тексты кодируют родственные вирусы. Неожиданным оказывается факт близости этих текстов с текстом FD, поскольку внутренняя организация вирусов ΦX174 и G4 сильно отличается от структуры FD. Объяснение может заключаться либо в том, что при различной внутренней организации похожими могут оказаться белки, кодируемые этими текстами, либо в том, что ранговые меры

близости не отражают адекватным образом различия во внутренней организации геномов. Последнее может объясняться тем, что ранговые меры близости не учитывают информации о расположении 1-грамм в тексте. Те же меры, которые это учитывают, весьма сложны для вычисления.

Тексты FBR322 и MS2, будучи весьма близкими по биграммным характеристикам ($\rho_2 = 0.62$), с увеличением n становятся менее похожими ($\rho_n = 0.18$). Так как характеристики четвертого порядка учитывают большую информацию, чем характеристики второго порядка, эти тексты, видимо, не следует считать слишком близкими.

5. Заключительные замечания. Проведенный анализ затронул всего лишь 6 текстов, поэтому многие выводы, сделанные выше, носят предварительный характер и подлежат дальнейшему уточнению по мере накопления материала. Однако уже сейчас можно сказать, что 1-граммный способ описания генетических текстов, сводящийся фактически к поиску всех периодичностей в символьной последовательности, является в сочетании с позиционным анализом адекватным средством для решения многих содержательных генетических задач.

Обобщением 1-граммного способа представления символьных последовательностей является представление, допускающее объединение в один класс не только совпадающих 1-грамм, но и 1-грамм, отличающихся от них в пределах заданной меры различия. Такое описание представляет интерес и с биологической точки зрения, например, в задаче поиска гомологичных участков в родственных геномах.

Л и т е р а т у р а

1. PATHER B.A. Молекулярно-генетические системы управления. - Новосибирск: Наука, 1975. - 286 с.
2. WAGNER R.A., FISHER M.I. The String-to-String Correction Problem. - J. Assoc. Comput. Machinery, 1974, v. 21, N 1, p. 168-173.
3. MASAMI Hasegawa, TERUO Yasunaga, TAKASHI Miyata. Secondary structure of MS2 phage RNA and Bias in code word usage. - Nucl. Acid. Res., 1979, v. 7, N 7, p. 2073-2079.
4. ШЕННОН К. Математическая теория связи. - В кн.: Работы по теории информации и кибернетике. М., 1963, с. 243-332.
5. ERICKSON J.W., ALFMAN J.J. A Search for Patterns in the Nucleotide Sequence of the MS2 Genome. - J. Math. Biology, 1979, v. 7, p. 219-230.
6. BIRD A.P. DNA methylation and the frequency of CpJ in animal DNA. - Nucl. Acid. Res., 1980, v. 8, N 7, p. 1499-1505.

Поступила в ред.-изд. отд.

5 ноября 1980 года