

1 УДК 621.391 : 519.218.5

2 © 20?? г. В.А. Любецкий, К.Ю. Горбунов, С.А. Пирогов, Г.А. Хазиев, А.И.

3 Агламазова

4 **КЛАСТЕРИЗАЦИЯ ТОЧЕК МНОГОМЕРНОГО ПРОСТРАНСТВА**
5 **НА ОСНОВЕ ИДЕОЛОГИИ SEURAT¹**

6 В статье на математическом уровне обсуждаются современные прикладные
7 задачи дискретной оптимизации, и для одной из них, кластеризации точек в
8 многомерном вещественном пространстве, подробно рассмотрен алгоритм её
9 решения. Характерной особенностью этих задач является большой размер ис-
10 ходных данных, часто задаваемых матрицами с десятками тысяч строк и десят-
11 ками тысяч столбцов. Это приводит к проблемам обработки больших данных
12 в условиях сложного вычислительного алгоритма; при этом нередко требуется
13 быстро получить достаточно точное решение задачи, поэтому вопрос о слож-
14 ности алгоритма имеет принципиальное значение. Решение упомянутой зада-
15 чи кластеризации приводит к трудным математическим проблемам: удалению
16 скрытых параметров; выделению значимых признаков; переходу к оптималь-
17 ным и информационно значимым координатам точек; специфическом представ-
18 лении (картами многообразия) вершин нагруженного графа; выбор функции,
19 зависящей от текущей кластеризации вершин, максимизация которой приводит
20 к искомой кластеризации; для каждого кластера выбор признаков, которые ин-
21 дивидуально характеризуют его; и, наконец, понижения размерности исходных
22 данных.

23 *Ключевые слова:* оптимальное преобразование графов, эволюция строки вдоль
24 дерева, оптимальная дискретная кластеризация.

25 **DOI:** 10.31857/S05552923??, **EDN:** ??

26 **§ 1. Введение**

27 В этой статье на математическом уровне изложены постановки ряда современных
28 (но уже ставших классическими) прикладных задач и для одной из них, кластериза-
29 ции множества точек в многомерном вещественном пространстве, описан алгоритм

¹ Работа выполнена в рамках государственного задания ИППИ РАН, утвержденного Минобрна-
уки России.

30 решения на основе идеологии Seurat. Этот алгоритм представляет собой пример
31 современной обработки данных, направленный на решение сложной вычислитель-
32 ной задачи. Решения поставленных задач широко востребованы и используются в
33 разных прикладных областях, но мы не касаемся здесь собственно прикладных ас-
34 пектов. Хотя эвристические алгоритмы для решения рассмотренной задачи класте-
35 ризации широко применяются, авторам неизвестны цельные математические изло-
36 жения её решения; в какой-то мере алгоритмы могут быть извлечены из прикладных
37 работ, однако, там изложения строятся в контексте соответствующих прикладных
38 областей, что затруднит математически ориентированного читателя.

39 Для таких задач кроме собственно математического исследования требуется най-
40 ти эффективный алгоритм и доказательство того, что алгоритм действительно на-
41 ходит заранее математически описанное решение; кроме того, нужно найти оценку
42 времени работы (сложность) алгоритма; и, в равной мере, найти эффективную ком-
43 пьютерную реализацию этого алгоритма. Поскольку эти задачи возникают в при-
44 кладном аспекте, не всегда осознаётся, что компьютерной программе предшествует
45 математическая постановка и, по меньшей мере, математическое осмысление задачи.
46 Иными словами, случается, что прикладник пишет эвристическую компьютерную
47 программу до математических постановки и исследования задачи.

48 В заключение заметим, что наша научно-методическая статья ориентирована на
49 широкий круг читателей; хотя лучше, если они знакомы с материалом первого курса
50 математического факультета.

51 Наконец, общим местом является замечание, что алгоритмическая и компьютер-
52 ная обработка больших данных – универсальное направление исследований и метод
53 работы буквально во всех областях естественных, инженерных и гуманитарных на-
54 ук. Такая обработка опирается на методы современной математики (от алгоритмов
55 – алгебры – анализа до геометрии) и также опирается на современные приёмы эф-
56 фективного программирования.

57 § 2. Примеры задач

58 Начнём с трёх примеров математических задач, хотя методы решения первых
59 двух из них здесь не обсуждаются (иначе пришлось бы значительно увеличить раз-
60 мер текста, и возникла тематическая разбросанность), но их математическое обсуж-
61 дение доступно по ссылкам.

62 1. Даны ориентированные графы с именами рёбер, без повторения имён или с их по-
63 вторениями. Такие графы иногда называют *структурами*; в сущности, даны две
64 картинки, показанные на рис. 1. Фиксированы шесть *операций*, преобразующие
65 одну структуру в другую; каждой операции сопоставлено строго положительное
66 рациональное (в общем случае, вещественное) число, которое называется *ценой*
67 данной операции. Цена показывает, сколь редко используется данная операция в

68 процессе, который сам описывается как минимум суммарной цены всех исполь-
69 зуемых в нём операций (с их повторениями): высокая частота использования
70 означает низкую цену операции, а низкая частота использования – её высокую
71 цену. Таким образом, обеспечивается адекватность этого процесса. Упомянутые
72 шесть операций над структурами хорошо известны и не приводятся здесь, см.
73 например, [1]. Итак, задача состоит в том, чтобы найти алгоритм, который по
74 двум данным структурам a и b и данным ценам операций выдаёт одну из цепочек
75 операций, которые последовательно преобразуют a в b , и имеет минимальную
76 суммарную цену операций.
77 В [2] получен алгоритм, по времени работы линейный от суммарного числа рёбер
78 в данных структурах a и b , который находит такую, минимальную при данных
79 ценах, цепочку операций. Вообще, линейная сложность алгоритма в реальных
прикладных задачах – большая редкость.

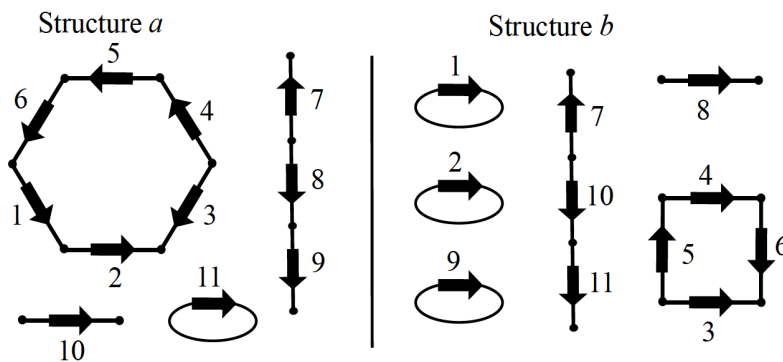


Рис. 1. Даны две структуры, нагруженные ориентированные графы a и b , и также цены каждой из шести операций над структурами. Требуется найти цепочку (последовательность) операций с их повторениями, суммарная цена которой минимальна по сравнению со всеми возможными цепочками, начинающимися с a и заканчивающимися на b .

80
81 2. В этом разделе речь идёт, скорее, о классе задач, для которых известно много эв-
82 ристических алгоритмов решения. Дано корневое дерево (иногда, ациклическая
83 сеть), для которого заданы длины рёбер, они соответствуют «дискретному вре-
84 мени» перехода от «предка» в начале ребра к «потомку» в конце ребра («начало»
85 расположено ближе к корню дерева, а «конец» – дальше от него). В листьях даны
86 современные однотипные данные из предметной области. Нелистовым вершинам
87 однозначно сопоставлены объекты того же типа, какой представлен в листьях, –
88 такое сопоставление называется *расстановкой* объектов по нелистовым верши-
89 нам; поскольку листьям объекты заранее сопоставлены, расстановка определена
90 на всех вершинах. В пространстве всех расстановок задан функционал, который

зависит от переменной (в нелистовых вершинах) расстановки и от длин рёбер. Задача состоит в минимизации функционала. Иными словами, задача состоит в описании дискретной эволюции «прапредка» – объекта в корне дерева вдоль всего дерева вплоть до современных данных в листьях. Примерами объектов служат последовательности в фиксированном алфавите, которыми, как известно, можно кодировать любой конечный объект; часто графы заданного типа. Иногда задача ставится так, что и само дерево является аргументом функционала или, наоборот, дано дерево и все рёбра считаются одинаковой длины, т.е. длины отсутствуют. Нами найдены эффективные компьютерные решения ряда задач такого рода и доказаны соответствующие теоремы, хотя в целом этот важный круг задач далёк от математического освещения, см., например, [3].

Один из важных подходов к определению такой эволюции состоит в следующем. Для данных точек в n -мерном пространстве задано ещё поле их скоростей, которое интерпретируется как потенциал развития точки-клетки (см. начало следующего пункта 3). Поле скоростей преобразуется в марковскую (или близкую к ней) цепь переходных вероятностей развития точек-клеток. По цепи образуются нечёткие множества клеток, которые образуют начальные, промежуточные и финальные макросостояния, и определяются вероятности перехода из одного макросостояния в другое. Отсюда находится эволюция макросостояний, от начального в финальные, которые интерпретируются как кластеры клеток [4].

3. Перейдём к задаче кластеризации точек в многомерном вещественном пространстве, решение которой составит оставшуюся часть статьи. Дана числовая матрица, строки которой называются *признаками* (features), а столбцы – *клетками*, не имея в виду именно биологическую клетку, а любую «изолированную часть мира». Далее столбцы будут рассматриваться как точки соответствующего пространства. Элемент матрицы – вещественное число, которое называется *экспрессией* (выраженностью) признака-строки i в клетке-столбце j . *Кластеризацией* столбцов матрицы называется разбиение её столбцов на непересекающиеся множества, каждое из которых называется *кластером*. Требуется найти кластеризацию, для которой столбцы внутри любого кластера наиболее похожи друг на друга в сравнении с похожестью столбцов между разными кластерами. «Похожесть» двух столбцов определяется функционалом, который не так просто сформулировать; это будет сделано по ходу решения задачи, что, увы, отступает от принципа, который пропагандировался во Введении.

Многочисленные эвристические решения этой задачи весьма сложны, в качестве примера мы сошлёмся на её решение в контексте прикладной биоинформатики, [5]. Мы изложим ниже в некоторых пунктах оригинальное решение, для которого разработали используемую нами эффективную компьютерную программу.

В заключение отметим, что в разделах 3-10 мы следуем общему плану метода Seurat-Louvain-Leiden. Наша цель – привлечь внимание математически ориентированного

131 читателя к многочисленным математическим проблемам, которые возникают в этом
 132 методе, далёком от математического не только обоснования, но даже отчасти и по-
 133 нимания. Это также относится и к задачам, сформулированным в разделах 1-2.

134 Следующий раздел 3 подробно излагается в разделах 4-10; читателю менее зна-
 135 комому с этим материалом можно вначале пропустить раздел 3.

136 § 3. План решение задачи кластеризации столбцов числовой матрицы.

137 Итак, клетка – точка в N -мерном пространстве вещественных чисел $\mathbb{R}^N$. Ины-
 138 ми словами, *координатами* клетки j называется столбец чисел в данной матрице,
 139 который удобно представить себе “расположенным под клеткой” j . Число столбцов
 140 в этой и последующих матрицах не меняется и равно M . Приведём план решения
 141 более или менее общий для разных подходов к этой задаче. Математическое обос-
 142 нование и оценка сложности предлагаемого ниже алгоритма далеки от решения,
 143 хотя активно изучаются. При желании можно прочесть пункты этого плана после
 144 остальных разделов статьи.

145 1. В данной числовой матрице $X = (x_{ij})$, размера $N \times M$, каждый её элемент x_{ij}
 146 заменим его долей относительно всего столбца j , так полученную матрицу обо-
 147 значим также X . Для каждой строки i матрицы X вычислим среднее x_i и выбо-
 148 рочную дисперсию σ_i^2 , и поэлементно преобразуем матрицу X в новую матрицу
 149 Y , в которой σ_i^2 уже слабо зависит от x_i .

150 Напомним, что для любой матрицы размера $N \times M$ среднее i -ой строки равно

$$x_i = \frac{1}{M} \sum_{j=1}^M x_{ij},$$

151 а дисперсия строки равна

$$\sigma_i^2 = \frac{1}{M-1} \sum_{j=1}^M (x_{ij} - x_i)^2$$

152 и стандартное отклонение строки равно σ_i .

153 2. По матрице Y образуем новую матрицу Z . Для этого по каждой строке $i \in$
 154 Y построим регрессию f для множества $\{(y_l, \sigma_l^2)\}$ точек на плоскости: точка –
 155 дисперсия σ_l^2 вместе со средним y_l , строки l , где y_l берётся только из заданной
 156 окрестности I_i среднего y_i . Тогда матрица Z определяется как

$$Z = \frac{y_{ij} - y_i}{f(y_i)}.$$

157 В Z выберем заданное число строк с наибольшей новой дисперсией $\delta = \frac{\sigma_i^2}{f(i)}$.
 158 Обозначим $\tilde{Z}$ часть матрицы Z , содержащую только эти строки.

- 159 3. Для столбцов матрицы $\tilde{Z}$ найдём лучшие координаты – координаты по базису из
160 левых сингулярных векторов этой матрицы. Затем спроектируем столбцы этой
161 матрицы на подпространство из n наиболее информативных таких векторов, где
162 информативность зависит от соответствующих сингулярных чисел. Полученную
163 матрицу обозначим Z^* , её размер $n \times M$. Само n определяется по доле необъ-
164 яснённой дисперсии для $\tilde{Z}$. Заметим, что квадраты сингулярных чисел и левые
165 сингулярные векторы совпадают с собственными числами и собственными векто-
166 рами, соответственно, ковариационной матрицы $\frac{1}{M-1} \tilde{Z}^T \cdot \tilde{Z}$ признаков, размера
167 $M \times M$.
- 168 4. Столбцы матрицы Z^* одновременно являются точками в $\mathbb{R}^n$ (столбцов-точек M
169 штук) в n -ом вещественном пространстве $\mathbb{R}^n$ и одновременно ещё они будут *вер-*
170 *шинами* следующего графа G . Для каждой из этих точек j по евклидову рас-
171 стоянию до других точек, найдём список из k ближайших соседей точки j (этот
172 j -список начинается с самой j), и рассмотрим *ранги* точек из j -списка по возрас-
173 танию расстояния от j .
- 174 5. В графе G проведём ребро между вершинами j и l , если их списки имеют об-
175 щего k -соседа. Рёбрам припишем веса, что закончит определение нагруженного
176 неориентированного графа G . Выберем некоторую начальную кластеризацию C_0
177 вершин в G .
- 178 6. Итеративно найдём *максимум функции модулярности*, аргументом которой яв-
179 ляется текущая кластеризация C вершин в G ; в результате найдём финальную
180 кластеризацию вершин в G (вершины – столбцы в матрице Z^*).
- 181 7. Вернёмся к матрице Y , которая имеет n и M исходных строк и столбцов, но
182 теперь ещё и кластеризацию её клеток (назовём её *матрицей кластеризации*).
183 В каждом кластере найдём *дифференциально-экспрессированные строки* (DE-
184 строки или, то же самое, *DE-признаки*) с помощью Mann-Whitney U -теста и
185 процедуры Benjamini-Hochberg.
- 186 8. Граф G (с координатами в n -ом пространстве) можно приблизить другим гра-
187 фом G_1 , вершины которого находятся в K -ом пространстве, где K много меньше,
188 чем n , например, $K = 2$. Это делается с помощью нелинейной функции, называ-
189 емой *UMAP*; *граф* G_1 в K -ом пространстве находится как минимум *дивергенции*
190 *Кульбака-Лейблера*. Решаемая здесь задача понижения размерности данных (то
191 есть размерности n до размерности K) сама по себе – одна из центральных при
192 работе с большими данными. Вообще говоря, такое понижение размерности мож-
193 но выполнить от любой размерности, например, от N к K .

194 § 4. лог-Преобразование матрицы X

195 Часто дисперсия строки зависит от её среднего. Например, пусть $X = \alpha + \mu e^\eta + \varepsilon$,
196 где X – измеряемая случайная величина и μ – её точное значение, а ε и η – адди-
197 тивная и мультипликативная погрешности. Пусть параметры подобраны так, что

каждая строка является выборкой случайной величины X со средним u и дисперсией v . Тогда $u = \alpha + \mu \cdot \mathbf{E}(e^\eta)$, и $v = \mu^2 \text{Var}(e^\eta) + \text{Var}(\varepsilon)$. Отсюда дисперсия v зависит от среднего u ; а именно, $v = \frac{\text{Var}(e^\eta)}{\mathbf{E}(e^\eta)^2} \cdot (u - \alpha)^2 + \text{Var}(\varepsilon) = (c_1^2 u + c_2)^2 + c_3$ с тремя параметрами c_1, c_2, c_3 . Здесь, например, $c_3 = \text{Var}(\varepsilon) + \alpha(\alpha - \frac{\text{Var}(e^\eta)}{\mathbf{E}(e^\eta)})$.

Мы хотим избавиться от «скрытого параметра» u с тем, чтобы после преобразования матриц, $Y = g(X)$, дисперсия $\text{Var}(Y)$ слабо зависела от математического ожидания $\mathbf{E}(Y)$, то есть в существенном зависела только от самой строки.

Для этого разложим Y по степеням $X - u$:

$$Y = g(u) + g'(u)(X - u) + g''(u)(X - u)^2 + \dots$$

В разложении оставим только аффинную часть и получим $\mathbf{E}(Y) \approx g(u)$ и $\text{Var}(Y) \approx g'(u)^2 \cdot v$.

Пусть $x_{i1}, \dots, x_{ij}, \dots, x_{iM}$ — строка i в матрице X со средним $x_i = u$, а $y_{i1} = g(x_{i1}), \dots, y_{ij} = g(x_{ij}), \dots, y_{iM} = g(x_{iM})$ — значения функции g , строка со средним $y_i = \mathbf{E}(Y) \approx g(u)$. Более того, $g(x_i) \approx g(u) \approx y_i$ и $\text{Var}(Y) = \sigma_i^2(Y) \approx v(u)g'(u)^2$.

Положим $v(u)g'(u)^2 = 1$. Отсюда $g(u) = \int_0^u \frac{du}{\sqrt{v(u)}} + C$. Интегрирование приводит к гиперболо-арсинусному преобразованию, указанному ниже.

Итак, применим функцию g к каждой ячейке в X , учитывая $\sigma_i^2(Y) \approx 1$, а векторы $(1, \dots, 1)$ и $(\dots, y_i, \dots)$ можно рассматривать как слабо зависящие друг от друга.

Таким образом, в нашем алгоритме для каждого элемента x_{ij} исходной матрицы X сначала вычислим его долю в столбце j и затем применим гиперболо-арсинусное преобразование:

$$x_{ij} \rightarrow y_{ij} = \begin{cases} \frac{1}{c_1} \cdot \text{arsinh}(\frac{c_2}{\sqrt{c_3}} + \frac{c_1}{\sqrt{c_3}} \cdot x_{ij}), c_3 > 0 \\ \frac{1}{c_1} \cdot \ln(c_2 + c_1 \cdot x_{ij}), c_3 = 0, \end{cases}$$

где гиперболический арсинус определяется как

$$\text{arsinh}(z) = \ln(z + \sqrt{z^2 + 1}),$$

а константы c_1, c_2, c_3 подбираются по исходным данным.

§ 5. Выделение строк с большой дисперсией

Для матрицы Y и каждой её строки i вычислим среднее y_i и дисперсию σ_i^2 . Построим полиномиальную низкой степени регрессию $f(\cdot)$ для множества точек $\{(y_i, \sigma_i^2)\}$, где y_i в середине интервала $\mathcal{O}_j$, а интервалы $\{\mathcal{O}_j\}$ выбраны так, что число точек в каждом из них примерно одинаково, рис. 2. Положим $f(i) = f(y_i)$ — значение регрессии $f(\cdot)$ в y_i . Регрессия относится к интервалу $\mathcal{O}_j$, и в этом смысле называется

226 локальной, а число $f(i)$ можно назвать *регрессионной дисперсией* строки i (иногда
 227 её называют средне-дисперсионным отношением). Она характеризует «дисперсию»
 228 строки i без учёта её индивидуальности. Иными словами, нас будут интересовать
 229 строки, для которых дисперсия σ_i^2 не подчиняется общему правилу, заданному ре-
 230 грессией $f(\cdot)$, то есть σ_i^2 и $f(i)$ значительно отличаются, $f(i)$ существенно меньше
 231 σ_i^2 .

232 образуем матрицу $Z = \frac{y_{ij} - y_i}{f(i)}$. Для её строк i среднее равно $z_i = 0$, а дисперсия
 233 равна $\delta_i = \frac{\sigma_i^2}{f(i)}$. В Z слишком большие дисперсии рассматриваются как артефакт
 234 исходных данных; поэтому строки, у которых $\delta_i^2 > \sqrt{M}$, удаляются из Z . В так
 235 полученной матрице Z оставим только строки с наибольшей регрессионной диспер-
 236 сией δ_i в заданном количестве. Обозначим $\tilde{Z}$ часть матрицы Z , содержащую только
 237 выбранные строки (например, 2000); тогда её размер $2000 \times M$.

238 Итак, строки матрицы $\tilde{Z}$ слабо зависят от среднего и не подчиняются регрес-
 239 сионному правилу, имея большую дисперсию. Это позволяет отчётливо различать
 240 столбцы по их координатам.

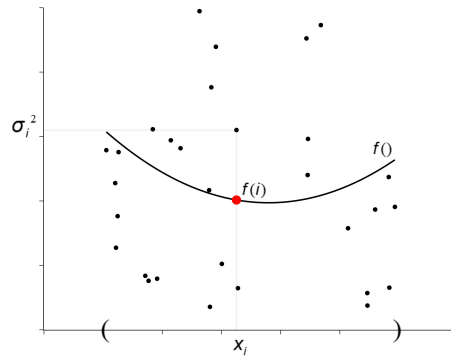


Рис. 2. Показана квадратичная регрессия, локальная на интервале $\mathcal{O}$ с серединой y_i , которая определяет регрессионную дисперсию $f(i)$ строки i матрицы Y . В этом интервале точки плоскости определяются как среднее y_l (абсцисса) строки l и её дисперсия σ_l^2 (ордината).

241 В результате каждая клетка-столбец получает 2000 координат, представляющих
 242 точку в 2000-ом пространстве вещественных чисел $\mathbb{R}^{2000}$. Получим множество $\mathcal{M}$ из
 243 M точек в этом пространстве.

§ 6. Выбор лучших координат для точек из множества $\mathcal{M}$

Выберем ещё лучшие, *новые координаты* клетки j , которые будут линейными комбинациями её старых координат, столбцов матрицы $\tilde{Z}$. Это выполняется в два шага.

1-ый шаг. Для матрицы $\tilde{Z}$ вычислим сингулярные числа и левые сингулярные векторы $u_1, u_2, \dots$ единичной длины, и перейдём к *ортонормированному* базису в $\mathbb{R}^{2000}$, состоящему из этих векторов, рис. 3. Удобно считать, что сингулярные числа расположены по убыванию $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_{2000}$, как и соответствующие им строки матрицы и новые координатные оси. Новые координаты клетки j (её старые координаты – элементы столбца $\tilde{Z}_j$) равны $(u_k, \tilde{Z}_j)$, где указана k -я координата, $1 \leq k \leq 2000$. Получен результат 1-го шага – матрица размера $2000 \times M$, которую обозначим $\hat{Z}$.

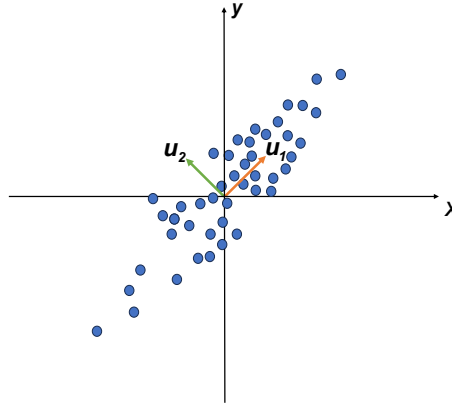


Рис. 3. Для клеток j показан переход от их старых координат – столбцов матрицы $\tilde{Z}$ к их новым координатам – столбцам матрицы $\hat{Z}$.

2-ой шаг. Оставим в $\hat{Z}$ только n новых и наиболее информативных координат, которые назовём *сокращёнными* (или наиболее информативными); они образуют матрицу Z^* размера $n \times M$, которая получена из матрицы $\hat{Z}$ 1-го шага размера $2000 \times M$.

Выбор числа n основан на следующем соображении информативности новых координат. Для краткости будем далее опускать знаки $\wedge$. Замена старых координат Z_j клетки j на её проекцию на первые n координат $Z_j \rightarrow Z_j^* = \sum_{k=1}^n (u_k, Z_j) \cdot u_k$ вносит средний квадрат ошибки в расчёте на одну клетку, равный $\frac{1}{M} \sum_{j=1}^M \|Z_j - \tilde{Z}_j^*\|^2 = \sum_{k=n+1}^N \lambda_k$, где $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_{2000}$. Эта сумма называется *остаточной дисперсией*.

266 Величина $\frac{1}{M} \sum_{j=1}^M \|\sum_{k=1}^n (u_k, Z_j)^2 \cdot u_k\|^2 = \frac{1}{M} \sum_{j=1}^M \sum_{k=1}^n (u_k, Z_j)^2 = \sum_{k=1}^n \lambda_k$ на-
 267 зывается *объяснённой дисперсией*. *Относительная ошибка* δ_n определяется как от-
 268 ношение остаточной дисперсии к выборочной дисперсии; эта доля называется *необъ-*
 269 *яснённой дисперсией*: $\delta_n = \frac{\lambda_{n+1} + \dots + \lambda_{2000}}{\lambda_1 + \lambda_2 + \dots + \lambda_{2000}}$.

270 Выберем n из условия $\delta_n \leq \varepsilon_0$, где ε_0 – заданный порог.

271 § 7. Кластеризация клеток по их новым сокращённым координатам (столбцам в 272 матрице Z^*)

273 образуем граф G : его вершины – клетки-точки j с их n сокращёнными и новыми
 274 (наиболее информативными) координатами – столбцами Z_j^* матрицы Z^* . Этот граф
 275 называется kNN -графом, а координаты клетки-точки-вершины j называются её
 276 *РСА-координатами*.

277 Для каждой вершины j по евклидову расстоянию между точками в n -мерном
 278 пространстве составим список k ближайших к ней вершин (k -соседей, сама j первая
 279 в этом j -списке, где k параметр). j -*Окрестностью* в G называется j -список вместе
 280 со всеми рёбрами, соединяющими любые две вершины из него. Заметим, что таким
 281 образом мы переходим к новым окрестностям для всех вершин j : к картам, которые
 282 представлены одномерными симплексами на многообразии всех клеток.

283 Список нумеруем подряд натуральными числами от 1 (которые называют *ран-*
 284 *гами*) по возрастанию евклидова расстояния. Обозначим $\text{rank}(v, j)$ – ранг вершины
 285 v в j -списке, который начинается с вершины j . Вершины, которые находятся на
 286 одинаковом расстоянии от j , нумеруются дробными числами: например, вершины,
 287 находящиеся на расстоянии 1, 2, 7, 7, 8, ... имеют ранги 1, 2, 3,5, 3,5, 4, ...

288 Вершины j и l соединим ребром, если $j \neq l$, а l является некоторым соседом в
 289 j -списке и j является некоторым соседом в l -списке.

290 Вес A_{jl} ребра $e = (j, l)$ положим равным максимуму по всем общим k -соседям v
 291 для j и l : $A_{jl} = \max\{k - 0,5(\text{rank}(v, j) + \text{rank}(v, l)) | v\}$.

292 Найдём некоторую начальную кластеризацию C_0 вершин графа G (разбиение его
 293 вершин). Выбор C_0 важен и существенно влияет на окончательный результат кла-
 294 стеризации, известно несколько способов такого выбора. Полученные в результате
 295 кластеры могут допускать рёбра между вершинами из разных кластеров; о таких
 296 рёбрах говорят, что они “*между кластерами*”; наоборот, о рёбрах, у которых оба
 297 края принадлежат одному кластеру, говорят, что они “*в кластере*”.

298 Начиная с начальной кластеризации C_0 , итеративно найдём локальный макси-
 299 мум функции $H(C)$, которая называется *модулярностью* кластеризации $C = C(G)$
 300 графа G и характеризует качество разбиения вершин в G . Функция модулярности

301 определяются как

$$H(\mathcal{C}) = \frac{1}{2m} \sum_{j,l} \left(A_{jl}(e) - \gamma \cdot \frac{k_j \cdot k_l}{2m} \right) \rightarrow \max.$$

302 Здесь A_{jl} – вес ребра e , которое пробегает *все рёбра* во всех кластерах текущего
303 разбиения $\mathcal{C} = \mathcal{C}(G)$, а m – сумма весов всех рёбер в G , и k_j сумма весов всех рёбер
304 в G , инцидентных вершине j ; γ – параметр, называемый *гранулярностью*: минимум
305 суммарного веса рёбер в кластере (относительно суммарного веса всех рёбер) по
306 всем кластерам, делённый на максимум аналогичной доли суммарного веса рёбер
307 между всеми парами разных кластеров.

308 О параметре γ заметим следующее. При его увеличении число отрицательных
309 рёбер e может увеличиться, и становится выгодным разбить кластер на несколь-
310 ко кластеров, чтобы отрицательные рёбра сосредоточились между кластерами и,
311 тем самым, были исключены из H . Тогда число кластеров увеличится, а сами они
312 уменьшатся. Наоборот, при уменьшении γ число положительных рёбер может уве-
313 личиться, поэтому становится выгодным объединить несколько кластеров в один,
314 чтобы включить в H межкластерные рёбра, ставшие положительными. Тогда число
315 кластеров уменьшится, и они увеличатся.

316 Нетривиальный выбор вида функции $H(\mathcal{C}(G))$ принципиально влияет на резуль-
317 тат кластеризации; отметим следующие соображения по поводу выбора и обоснова-
318 ния функции модулярности. Обозначим $r \in \mathcal{C}$ – кластер данной кластеризации $\mathcal{C}$.
319 Веса A_{jl} рёбер в G можно интерпретировать как мультиграф без весов: вершины j и
320 l соединяются рёбрами в количестве A_{jl} . Обозначим той же буквой A_{jl} симметрич-
321 ную матрицу инцидентности полученного *мультиграфа* без петель, который так же
322 обозначим G (на диагонали находятся нули, остальные числа равны весу $\frac{A_{jl}}{2}$; как
323 обычно, у петли вес удваивается). Тогда $A_{jl} = A_{lj}$ и суммирование по всем рёбрам
324 $j, l \in r$ приводит к удвоенной сумме числа рёбер в r . В формуле для H перейдём
325 от суммирования по рёбрам $e = (j, l) \in G$ к суммированию по кластерам $r \in G$ и
326 получим

$$H(\mathcal{C}(G)) = \frac{1}{m} \sum_r \left[\left(\frac{1}{2} \cdot \sum_{(j,l) \in r} A(j,l) \right) = \frac{k(r)^2}{4m} \right],$$

327 где m – число рёбер в G и $\sum_{(j,l) \in r} A(j,l)$ – удвоенное число рёбер внутри кластера r .
328 Здесь $k(j)$ – степень инцидентности вершины j в графе G , вычисляемая по матрице
329 инцидентности A_{jl} , а $k(r) = \sum_{j \in r} k(j)$ – степень инцидентности кластера r . Легко
330 видеть, что $\sum_{j \in G} k(j) = 2m$. Итак, уменьшаемое – число рёбер в $r \in \mathcal{C}(G)$, и мы
331 хотим интерпретировать вычитаемое.

332 Это можно сделать как среднее число рёбер в семействе случайных графов с те-
 333 ми же вершинами (и параметрами), что у G и той же кластеризацией этих вершин.
 334 Однако приведём более простое комбинаторное пояснение. На множестве вершин
 335 мультиграфа G с его кластеризацией рассмотрим “идеальный” (не мульти, обычный)
 336 граф G' с той же кластеризацией и теми же степенями инцидентности $k(j)$, что у
 337 вершин мультиграфа G . Напомним, что в G и G' степень инцидентности определя-
 338 ется как сумма весов всех инцидентных ей рёбер. А именно, G' полный с петлями
 339 и в нём вес ребра (j, l) полагаем равным $\frac{k(j)k(l)}{2m}$, если $j \neq l$, и $\frac{k(j)k(l)}{4m}$, если $j = l$. В
 340 G' для любой вершины j степень инцидентности равна $\sum_{l \neq j} \frac{k(j)k(l)}{2m} + \frac{2k(j)^2}{4m} = k(j)$.
 341 И в G' сумма весов всех рёбер в r равна $\sum_{j \neq l, \{j, l\} \in r} \frac{k(j)k(l)}{2m} + \sum_j \frac{k(j)^2}{4m} = \frac{k(r)^2}{2m}$. По-
 342 следнее проверяется прямым раскрытием скобок в $k(r)$. Повторим, что графы G и
 343 G' имеют в качестве общих параметров саму кластеризацию $\mathcal{C}$, количества вершин
 344 и рёбер, и степени инцидентности $\{k(j)\}$ для всех вершин j . Величины $k(j)$ можно
 345 рассматривать как характеристики неоднородности графа, которые одинаковы у G
 346 и G' : у них в одних и тех же областях плотность рёбер повышена или понижена по
 347 сравнению со средней плотностью.

348 **7.1. Итеративный алгоритм кластеризации Leiden.** Изложенный выше алгоритм
 349 уточняется следующим образом. Для исходного графа G выбирается начальная кла-
 350 стеризация $\mathcal{C}_0$ и максимизируется функция модулярности $H(\mathcal{C}(G))$. По полученному
 351 графу строим новый граф G_1 , в котором все вершины одного кластера в G пред-
 352 ставлены одной новой вершиной в G_1 . Все рёбра внутри кластера в G заменяются
 353 петлёй в G_1 у этой новой вершины, а все рёбра между двумя разными кластерами
 354 в G заменяются одним новым ребром в G_1 между соответствующими новыми вер-
 355 шинами. Вес петли полагаем равным суммарному весу исходных соответствующих
 356 внутрикластерных рёбер, а вес нового ребра – суммарному весу исходных соответ-
 357 ствующих межкластерных рёбер. После этого максимизируем функцию $H(\mathcal{C}(G_1))$.
 358 Таким образом процесс продолжается, пока максимум $H(\mathcal{C}(G_k))$ на очередном гра-
 359 фе G_k увеличивается, рис. 4. Затем последовательными объединениями вернёмся к
 360 исходному графу G . А именно, если вершины в G_{k-1} входят в один и тот же кластер
 361 в G_k , то соответствующие им кластеры объединяются в один кластер в G_{k-1} . И об-
 362 ратным ходом придём к кластеризации исходного графа G . Получим кластеризацию
 363 вершин в исходном G . При каждом переходе от $k-1$ к k качество соответствующей
 364 кластеризации графа G_k не уменьшается.

365 **7.2. К обоснованию итеративного алгоритма кластеризации Louvain-Leiden.** В
 366 качестве начальной кластеризации можно выбрать *синглетонную кластеризацию*,
 367 которая состоит из всех вершин графа G_k , рассматриваемых как одноэлементные
 368 множества (“синглетоны”). Но обычно в качестве начальной кластеризации берут бо-
 369 лее крупную кластеризацию, называемую *окрестностной*; обозначим её $\mathcal{C}_0$. Граф

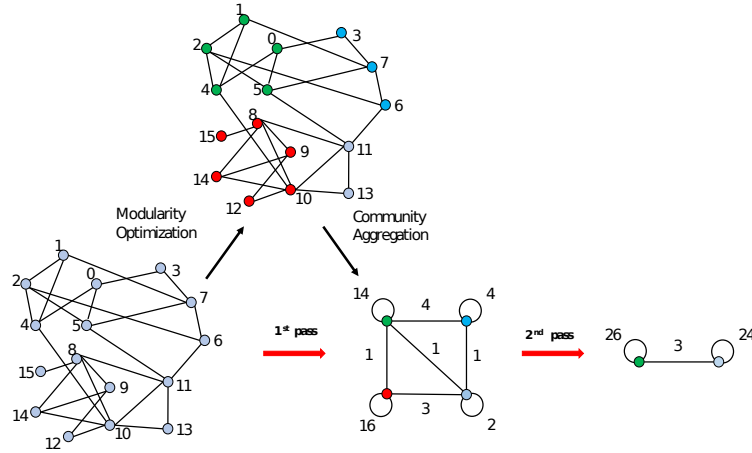


Рис. 4. Веса всех рёбер единичные. Показаны шаги итеративного алгоритма Leiden. Пример взят из [6].

370 G_k далее обозначим G , например, это граф G , определённый в начале этого разде-
 371 ла, с которого и начинает работу итеративный алгоритм. Для построения этой C_0
 372 используются описанные в начале этого раздела j -окрестности всех вершин в G , и
 373 веса рёбер графа G . Опишем кластер $c \in C_0$, каждый кластер образуется по неко-
 374 торой вершине $j \in G$ и в этом смысле обозначается c_j (кластер с “центром” в j).
 375 Кластеры c_j образуются итеративно, каждый начиная с j -окрестности. В подграфе
 376 c_j обозначим $d(v)$ сумму весов рёбер в c_j , инцидентных v ; это – степень вершины
 377 v , отнесённая к c_j . Проредим каждое c_j , удаляя из него вершины $v \in c_j$ вместе со
 378 всеми инцидентными им рёбрами, если $\frac{d(v)}{|c_j|} < r$, где $r \in (0, 1]$ – порог и $|c_j|$ – сумма
 379 весов всех рёбер в c_j . Удаление таких вершин из c_j зависит от порядка просмотра
 380 вершин в c_j ; продолжаем удаление, пока такая вершина находится. Удаления вы-
 381 полняются независимо для каждого j . Если $|c_j| = 0$, то вершину j сохраним; если v
 382 удаляется из всех j -окрестностей в G , то оставим v в качестве синглтона.

383 После этого в полученном наборе подграфов c_j устраним их пересечения по вер-
 384 шинам: пусть v принадлежит нескольким c_j . Для каждого c_j вычислим *средний вес*
 385 рёбер, соединяющих v со всеми вершинами в c_j . Вершину v сохраним в той c_j , для
 386 которой средний вес максимален, а из остальных c_j удалим её. Итоговые c_j образуют
 387 кластеры начальной кластеризации C_0 в графе G .

388 Наметим доказательство того, что при каждом переходе от k к $k + 1$ качество
 389 соответствующей кластеризации графа G_k не уменьшается.

390 Для краткости докажем здесь более высокое качество кластеризации C_0 по срав-
 391 нению с синглетонной в упрощённом случае, когда учитываются только k рёбер при
 392 каждой вершине, а вершины из окрестностей не удаляются и ближайший сосед двух
 393 вершин из j -окрестности – сама точка j , и ранги – целые попарно различные числа.
 394 Пусть G – полный граф с петлями: так будет, если отсутствующие рёбра положить

с весом 0. Напомним, что значение H на любой кластеризации в G (с точностью до мультипликативной константы) равно сумме разностей весов рёбер, лежащих внутри кластера в G и в идеальном графе G' . Петли лежат внутри кластера при любой кластеризации и, следовательно, приносят одинаковый вклад в H для синглетонной и окрестностной кластеризаций. Поэтому нам требуется оценить суммы весов в G и G' рёбер, которые лежат внутри кластера в окрестностной кластеризации, но не в синглетонной. Это все рёбра, не являющиеся петлями, в кластерах $c_j \in C_0$.

Поэтому в G сумма весов рёбер, которые не петли, из $c_j \in C_0$ равна $\frac{k^2(k-1)}{2} - \frac{1}{2} \sum_{i \neq l} (i+l) = \frac{k^2(k-1)}{2} - \frac{(k-1)k(k+1)}{4} = \frac{k(k-1)^2}{4}$, здесь i и l – ранги вершин в c_j , которые принимают значения от 1 до k . Оценим сумму k_i весов рёбер, инцидентных вершине ранга i из $c_j \in C_0$:

$$k_i = k(k-1) - \frac{1}{2} \sum_{l \neq i} (i+l) = \frac{3k^2 - 5k - 2ik + 4k}{4} \leq \frac{3k(k-1)}{4}.$$

При $\gamma = 1$ в идеальном графе G' сумма весов рёбер из c_j равна $\frac{1}{2m} \sum_{i \neq l \in c_j} k_i k_j \leq \frac{1}{2m} \left(\frac{3}{4} k(k-1) \right)^2 \frac{k(k-1)}{2} = \frac{9}{64m} k^3 (k-1)^3$. Итак, качество окрестностной кластеризации больше качества синглетонной, если $\frac{k(k-1)^2}{4} > \frac{9}{64m} k^3 (k-1)^3$, то есть $m > \frac{9}{16} k^2 (k-1)$. Обычно обрабатываются графы, у которых $m > 15000$, а k не больше 30, так что это неравенство выполняется.

Теорема 1. Пусть граф G полный с петлями и C_k – итоговая кластеризация k -го уровня, а C_{k+1} – кластеризация $(k+1)$ -го уровня, в которой все кластеры – сингletonы. Тогда $H(C_k) \leq H(C_{k+1})$.

Доказательство. Назовём кластеры кластеризации в C_{k+1} $(k+1)$ -го уровня крупными, кластеры кластеризации в C_k k -го уровня – средними, кластеры кластеризации в C_{k-1} $(k-1)$ -го уровня – мелкими. Аналогично назовём и сами кластеризации в этих графах.

Из определения веса ребра между кластерами следует, что суммарный вес m всех рёбер между кластерами не меняется при переходе к кластеризации следующего уровня (независимо от того, являются ли кластеры следующего уровня сингletonами). И также следует: если крупные кластеры – сингletonы, то суммарный вес всех рёбер, лежащих внутри такого сингletonа (а это вес петли при соответствующем среднем кластере) равен суммарному весу рёбер между мелкими кластерами, лежащими внутри среднего. Суммируя этот вес по сингletonам, получим, что сумма $\sum A_{jl}$ из формулы для H для крупной кластеризации равна этой сумме для средней.

Осталось рассмотреть сумму произведений $k_j \cdot k_l$ для пар (j, l) вершин, лежащих внутри кластера данной кластеризации. Любой крупный кластер-сингleton a состоит из одного среднего кластера, который обозначим a' . Для крупной кластеризации слагаемое в упомянутой сумме – квадрат веса петли при a' , назовём её первой сум-

430 мой. Для средней кластеризации это – сумма по всем парам (j, l) мелких кластеров
 431 (в a') произведений $k_j \cdot k_l$; назовём её второй суммой.

432 Первая сумма не больше второй. Действительно, в первой сумме при раскрытии
 433 скобок возникают произведения $w(e_1) \cdot w(e_2)$ весов рёбер e_1 и e_2 между мелкими
 434 кластерами, лежащими внутри a' , причём каждое произведение встретится два ра-
 435 за, если $e_1 \neq e_2$ и один раз, если иначе. Во второй сумме при раскрытии скобок
 436 также возникнут произведения $w(e_1) \cdot w(e_2)$ весов рёбер e_1 и e_2 , уже не обязательно
 437 лежащих внутри a' . Однако каждое такое произведение для рёбер, лежащих внутри
 438 a' , встретится не меньше двух раз, если $e_1 \neq e_2$ и не меньше одного раза, если иначе.
 439 Поэтому часть суммы $\sum k_j k_l$ из формулы для H , относящаяся к одному кластеру-
 440 синглетону a , для крупной кластеризации не превосходит аналогичную часть суммы
 441 (относящуюся к a') для средней кластеризации. Суммируя по всем синглетонам и
 442 учитывая, что сумма $\sum k_j k_l$ входит в H со знаком минус, получим утверждение. ▲

443 Однако нет гарантии, что функция H не уменьшится на обратном ходе проце-
 444 дуры. Поэтому, мы начинали этот ход с итоговой кластеризации каждого, не обяза-
 445 тельно заключительного, уровня и из всех таким образом полученных кластериза-
 446 ций в G выбирали кластеризацию с максимальным значением H .

447 § 8. Дифференциально-экспрессируемые признаки.

448 Уже получена кластеризация (разбиение) $\{G_s\}$ клеток-столбцов исходной мат-
 449 рицы Y размера $N \times M$, которая прошла log-нормирование, но не шкалирование-
 450 центрирование; матрица состоит из старых координат, до перехода к PCA-координатам
 451 и сокращения менее информативных из них; здесь s пробегает все кластеры. Для
 452 каждого кластера G клеток и каждой строки i определим дифференциально экс-
 453 прессируемые признаки-features. Такие признаки-строки сокращённо называются
 454 *DE-признаками* (или *DE-строками*). Говоря интуитивно, признак-строка i назы-
 455 вается *дифференциально-экспрессируемой*, если экспрессии в клетках $j \in G$ стати-
 456 стически значимо отличаются от экспрессий во всех остальных клетках $l \notin G$ той
 457 же строки i . Это понимание *DE-строк* для данного кластера G ниже уточняет-
 458 ся: сначала с помощью Mann-Whitney U -теста [7–9], затем с помощью процедуры
 459 Benjamini-Hochberg [10], а также ещё ряда условий.

460 Итак, фиксирована строка i (её имя и соответствующий индекс можно не упоминать)
 461 и фиксирован кластер G , который индуцирует разбиение строки на два мно-
 462 жества, которые будем обозначать теми же буквами G и $\neg G$, по-прежнему называя
 463 их кластерами. Расположим все экспрессии строки i в порядке их возрастания; и
 464 присвоим числам в полученной последовательности экспрессий *ранги* – натураль-
 465 ные числа также по их возрастанию, начиная с 1. Одинаковым числам присвоим
 466 средний по ним ранг. Пусть R_1 – сумма рангов экспрессий в данном кластере G , и
 467 R_2 – сумма рангов экспрессий вне этого кластера, т.е. в $\neg G$; пусть m_1 и m_2 – мощно-

сти множеств G и $\bar{G}$, соответственно; $m_1 + m_2 = M$. Поскольку $R_1 + R_2 = M \cdot \frac{1+M}{2}$, можно оперировать только с R_1 или только с R_2 .

Предполагается, что G и $\bar{G}$ – *независимые* выборки из двух непрерывных распределений $\mathfrak{F}$ и $\mathfrak{L}$, которые *совпадают* или *не совпадают*. Для экспериментально-математических выборок G и $\neg G$ вопрос об их независимости трудно разрешим, так что он остаётся на усмотрение вычислителя; в нашей ситуации это является предположением. Если $\mathfrak{F} = \mathfrak{L}$, то для экспрессий $y_j \in G$ и $y_l \in \neg G$ выполняется $\mathbf{P}(y_j < y_l) = 1/2$. Это свойство называется отсутствием “доминирования” G и $\neg G$ друг над другом. Обратное неверно: например, если $\mathfrak{F}$ и $\mathfrak{L}$ – два нормальных распределения с одинаковыми математическими ожиданиями, но разными дисперсиями, то доминирование отсутствует: $\mathbf{P}(y_j - y_l < 0) = \frac{1}{2}$, так как разность этих распределений имеет нулевое математическое ожидание.

Итак, мы хотим проверить, выполняется ли θ -ая гипотеза, которая состоит в “отсутствии доминирования G и $\neg G$ друг над другом”. В случае отрицания (говорят: отвержения) 0-ой гипотезы, строка i называется *DE-признаком*. Если 0-ая гипотеза отвергается, то на том же уровне достоверности выполняется $\mathfrak{F} \neq \mathfrak{L}$.

8.1. Критерий Манна-Уитни-Вилкоксона (Wilcoxon rank-sum test, Mann-Whitney U test и определение p -value; сокращённо МУВ). По строке i матрицы Y образуем случайную величину $U = R_2 - \frac{m_2(m_2+1)}{2} = \sum_{j=1}^{m_1} \sum_{l=1}^{m_2} I(x_j < x_l)$, где x_j и x_l – экспрессии в кластере G и вне него, то есть в $\neg G$, соответственно; а $I(\cdot)$ – характеристическая функция условия, которое приведено в скобках. Если 0-ая гипотеза выполняется, то её математическое ожидание и дисперсия равны $\mathbf{E}(U) = \frac{m_1 \cdot m_2}{2}$ и $\sigma(U)^2 = m_1 \cdot m_2 \frac{M+1}{12}$.

В разделе 6 выполнено центрирование и нормирование матрицы Y . Аналогично этот общий приём применяется и к U , чтобы получить статистику (случайную величину) $Z = \frac{U - m_1 m_2 / 2}{\sqrt{m_1 m_2 (M+1) / 12}}$. Если в последовательности рангов наблюдается асимметрия (относительно линейного порядка чисел) экспрессий в пользу G или, наоборот, в пользу $\neg G$, то Z *отклоняется* вправо или влево от нуля, больше или меньше. Если это отклонение Z превосходит заданный порог Z_0 , то строку-признак i считаем *DE*; уточним этот подход.

Повторим, пусть G и $\neg G$ – выборки двух некоторых случайных величин. Повторим, 0-я гипотеза состоит в том, что $\mathbf{P}(y_j < y_l) = \frac{1}{2}$ для $y_j \in G$ и $y_l \in \bar{G}$ (говорят: « y_j не доминирует y_l , а y_l не доминирует y_j »). Эта гипотеза выражает “отсутствие порядковой асимметрии” между G и $\bar{G}$. Заметим, что распределение для Z задаётся мерой на $\mathbb{R}$, которая определяется как Z^{-1} от меры на M -ом пространстве, индуцированной распределениями $\mathfrak{F}$ и $\mathfrak{L}$.

При выполнении 0-й гипотезы и возрастающих $m_1, m_2 \rightarrow \infty$, распределение статистики Z стремится к нормальному 0-1 распределению, обозначаемому $N(0, 1)$. В

силу этого приближённо случайная величина Z имеет распределение $N(0, 1)$, то есть вместо обычно неизвестных распределений $\mathfrak{F}$ и $\mathfrak{L}$ вероятность, далее везде обозначаемую $\mathbf{P}$, будем вычислять по нормальному распределению $N(0, 1)$.

Хотя это дальше не используется, заметим: если экспрессия в G доминирует экспрессию в $\neg G$, то есть $\mathbf{P}(y_j < y_l) > 1/2$, то Z стремится к $+\infty$, а если экспрессия в G в том же смысле доминирует экспрессию в $\neg G$, то Z стремится к $-\infty$ при также возрастающих m_1, m_2 .

Выберем Z_0 из условия $2\mathbf{P}(x > Z_0) = \alpha$, тогда Z_0 называется *квантилем уровня значимости* α ; само α задаётся вычислителем, например, $\alpha = 0, 1$.

Значение p -value для данной строки i , обозначаемое $p_{\text{val}}(i)$, определяется по распределению $N(0, 1)$ как удвоенная площадь под распределением правее числа $|Z|$, то есть $p_{\text{val}}(i) = 2\mathbf{P}(x > |Z|)$; а эмпирическое значение $Z = Z(i)$ вычисляется по данной строке i .

Заметим $|Z(i)| > Z_0 \Leftrightarrow p_{\text{val}}(i) < \alpha$. Число $p_{\text{val}}(i)$ называется *текущим уровнем значимости* с уровнем значимости α ; например, $\alpha = 0, 05$.

Если $p_{\text{val}}(i) < \alpha$, то строка-признак i включается в *предварительный список DE-признаков* данного кластера G клеток (и 0-я гипотеза отвергается). Иначе для i принимается 0-я гипотеза.

Дальнейшее сокращение *предварительного списка DE-признаков* связано с возможной случайностью самого значения $p_{\text{val}}(i)$.

8.2. Дополнение МУВ процедурой Беньямини-Хохберга. Определение p_{adj} -значения.

Переставим строки-features i , $1 \leq i \leq N$, в матрице Y по возрастанию самих чисел $p_{\text{val}}(i)$: $p_{\text{val}}(1) \leq \dots \leq p_{\text{val}}(n)$, где бывшая строка i получает какой-то свой номер k .

По определению:

$$p_{\text{adj}}(k) = \min \left\{ \frac{n}{k} \cdot p_{\text{val}}(k), p_{\text{adj}}(k+1) \right\}, p_{\text{adj}}(n) = p_{\text{val}}(n).$$

Тогда $p_{\text{adj}}(k) \leq p_{\text{adj}}(k+1)$ и $0 < p_{\text{val}}(k) \leq p_{\text{adj}}(k) \leq 1$.

Итак, *DE-признаком* ещё раз предварительно называется строка i (с новым номером k), которая удовлетворяет в качестве 1-го шага в пункте 8.1 условию $p_{\text{adj}}(k) < \alpha$. Очень слабым вариантом является $\alpha = 0, 1$.

Обратной индукцией легко проверить: для строки k , если выполняется $p_{\text{val}}(k) > \alpha$ (то есть строка k не-DE-признак относительно p_{val}), то $p_{\text{adj}}(k) > \alpha$ (эта же строка – не-DE-признак относительно p_{adj}).

На так отобранные DE-признаки накладываются ещё три дополнительных условия соответственно с параметрами q, q_1, q_2 (например, $q = 2, q_1 = q_2 = 0, 25$ или очень слабый вариант параметров $q = 0, 1, q_1 = 0, 01, q_2 = 0$).

540 Первое условие на DE-признаки:

$$\log_2 \text{FC}(i) = \left| \log_2 \left(\frac{1}{m_1} \cdot \sum_{j \in G} x_j \right) - \log_2 \left(\frac{1}{m_2} \cdot \sum_{l \notin G} x_l \right) \right| > q.$$

541 Второе и третье условия на DE-признаки: для признака i обозначим pct.1 и pct.2
542 доли клеток с ненулевой экспрессией в кластере G и в его дополнении $\neg G$, соответ-
543 ственно; и положим: $\max\{\text{pct.1}, \text{pct.2}\} \geq q_1$, $|\text{pct.2} - \text{pct.1}| \geq q_2$.

544 Окончательное определение DE-признака состоит в выполнении всех перечис-
545 ленных условий.

546 Маркером называется DE-признак, для которого $p_{\text{adj}}(i) < \beta$ с особо малым значе-
547 нием $\delta \ll \alpha$ (например, $\beta = 10^{-20}$ или меньше); значение β подбирается по данным.

548 § 9. Пояснение к процедуре Беньямини-Хохберга.

549 До конца этого пункта предположим 0-ую гипотезу, обозначив её (*).

550 Напомним, что распределение для Z лишь приближается к нормальному распре-
551 делению; а $\mathbf{P}(x)$ здесь удобнее понимать как $\mathbf{P}(x) = \int_x^\infty e^{-\frac{t^2}{2\pi}} dt$.

552 Тогда вероятность *отвержения* 0-ой гипотезы равна $\mathbf{P}(p_{\text{val}}(i) < \alpha) = \mathbf{P}(x >$
553 $\mathbf{P}^{-1}(\alpha/2))$, где $\mathbf{P}^{-1}(\alpha/2) = Z_0$ и $\mathbf{P}^{-1}$ – обратная функция к $\mathbf{P}$. Итак, если 0-
554 ая гипотеза верна, то $\mathbf{P}(p_{\text{val}}(i) < \alpha) = \alpha$. Матрица Y имеет N строк, поэтому
555 $\mathbf{P}$ (“для всех строк 0-ая гипотеза принимается”) = $(1-\alpha)^N$, а $\mathbf{P}$ (“для некоторых строк
556 0-ая гипотеза отвергается”) = $1 - (1-\alpha)^N$, что при большом N кажется парадок-
557 сальным в предположении (*). Упомянутая в пункте 8.2 перенумерация строк (от i
558 на k) не влияет на эту кажущуюся парадоксальность.

559 Для процедуры Беньямини-Хохберга доказано (здесь не приводится), что при
560 замене условия $p_{\text{val}}(i) < \alpha$ на условие $p_{\text{adj}}(i) < \alpha$ получим: $\mathbf{P}(\{k | p_{\text{adj}}(i) < \alpha\} = \emptyset) \geq$
561 $1 - \alpha$ что уменьшает соединение маловероятного события с предположением (*).
562 Заметим, что замена p_{val} на p_{adj} , с одной стороны, удаляет ложно дифференциально-
563 экспрессируемые признаки (“ложно-положительные”), а с другой стороны, может
564 привести к потере реально дифференциально-экспрессируемых признаков (“ложно-
565 отрицательных”).

566 § 10. Понижение размерности данных.

567 Если кластеризованное множество точек (данные) в пространстве $\mathbb{R}^n$ большой
568 размерности n отобразить в пространство $\mathbb{R}^K$, где $K < n$, например, на плоскость
569 ($K = 2$), обычной (линейной) проекцией, то проекции кластеров в $\mathbb{R}^n$, т.е. кластеры
570 в $\mathbb{R}^K$, могут перекрывать друг друга или находиться в проекции ближе, чем по ев-
571 клидову расстоянию в $\mathbb{R}^n$. Такой неадекватной ситуации стремится избежать нели-

нейная «проекция», называемая UMAP, рис. 5. Обычно данные в $\mathbb{R}^n$ и $\mathbb{R}^K$ представляются графами G и G_1 , вершины которых однозначно соответствуют друг другу и точкам соответствующих пространств. Итак, задача понижения размерности состоит в переходе от данных в $\mathbb{R}^n$ к соответствующим данным в $\mathbb{R}^K$. А точнее, от графа G к графу G_1 , у которых одинаковые клетки-вершины.

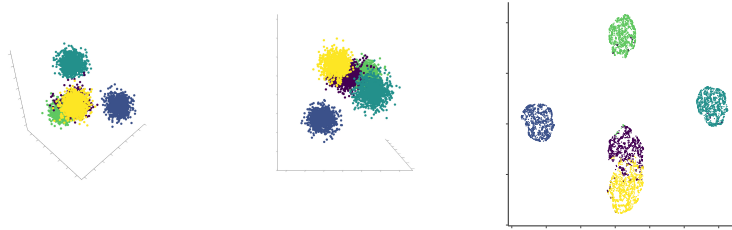


Рис. 5. Показан пример нелинейной проекции UMAP данных (здесь не приводятся) в большой размерности на плоскость ($K = 2$).

Повторим точнее, даны клетки-точки j в n -ом пространстве с их PCA-координатами (и одновременно j – вершины графа G). Как и в пункте 7, для каждой точки-клетки j строится список k -соседей и $\rho(j, l)$ – евклидово расстояние от j до её k -соседа l . Пусть ρ_j – наименьшее ненулевое расстояние от j до её ближайшего k -соседа. Для данного j вычислим нормировочное число σ_j , заданное уравнением

$$\sum_l \exp \left(-\frac{\rho(j, l) - \rho_j}{\sigma_j} \right) = \log_2 k,$$

где l пробегает всех k -соседей точки-клетки j .

Строится новый граф G_1 , который имеет такие же вершины-клетки, рёбра и кластеризацию, какие у графа G ; но вершины у G_1 однозначно соответствуют точкам K -мерного пространства. Тем самым, точки-вершины у G переходят в точки-вершины у G_1 ; это отображение и называется UMAP, [11].

Перейдём к построению *весов* рёбер и *координат* вершин у графа G_1 . Но сначала изменим веса в графе G , и полученный граф обозначим G_{01} . В нём вершины-клетки j и l , соединяются неориентированным ребром, если они являются k -соседями друг друга, как и было в пункте 7. Но вес ребра $e = (j, l)$ в G_{01} положим равным $W(j, l) = w(j \rightarrow l) + w(l \rightarrow j) - w(j \rightarrow l) \cdot w(l \rightarrow j)$, $0 < W(j, l) < 1$, где вспомогательный односторонний вес равен

$$w(j \rightarrow l) = \exp \left(-\frac{\rho(j, l) - \rho_j}{\sigma_j} \right).$$

593 Вес $W(j, l)$ можно представить себе как вероятность существования хотя бы од-
594 ного ребра между вершинами j, l .

595 Таким образом, получен новый в части весов неориентированный нагруженный
596 граф G_{01} в n -мерном пространстве “большой размерности”, где те же самые клетки
597 j являются вершинами с теми же рёбрами e и новыми весами $W(e)$. Пусть для
598 краткости $K = 2$.

599 Перейдём к построению *самого графа* G_1 на плоскости: он имеет те же вершины
600 j и рёбра e , что у графов G и G_{01} , но опять новые веса. Обозначим $x(j)$ координаты
601 на плоскости клетки-вершины $j \in G_1$. Определим в G_1 вес того же, что в G_{01} , ребра
602 $e = (j, l)$, полагая $S_{jl} = (1 + a\|x(j) - x(l)\|_2^{2b})^{-1}$, $0 < S_{jl} < 1$, где a и b – параметры
603 (“УМАР-веса”). Веса S_{jl} в G_1 близки к расстоянию, обратному к евклидову между
604 $x(j)$ и $x(l)$.

605 Осталось вычислить координаты $x(j)$ всех вершин j так, чтобы G_{01} и G_1 бы-
606 ли в некотором смысле близки друг к другу. Задача УМАР состоит в том, чтобы
607 найти все координаты $x(j)$ из условия близости двух распределений на одних и тех
608 же рёбрах в G_{01} и G_1 . А именно, найти их из условия минимума следующей функ-
609 ции, называемой дивергенцией Кульбака-Лейблера: $L(x) = \sum_e L(x(j_1), \dots, x(j_M)) =$
610 $\sum_e \left[W(e) \log \left(\frac{W(e)}{S(e)} \right) + (1 - W(e)) \log \left(\frac{1 - W(e)}{1 - S(e)} \right) \right] \rightarrow \min$, минимизируя по координатам
611 $x = x_1, x_2; \dots; x_{M-1}, x_M$ всех вершин в G_1 .

612 Результат минимизации определит искомый 2-мерный граф G_1 на плоскости.
613 Повторим: на него переносится кластеризация с графа G в n -мерном пространстве.

614 Как обычно, численная минимизация существенно зависит от выбора начальной
615 точки. Начальными координатами вершин графа G_1 можно выбрать две первые,
616 наиболее информативные координаты точек в n -мерном пространстве.

617 Минимизируемый функционал перепишем:

$$L(x) = \sum_j \sum_{l \neq j} \left(W_{jl} \cdot \log \frac{W_{jl}}{S_{jl}} + (1 - W_{jl}) \cdot \log \frac{1 - W_{jl}}{1 - S_{jl}} \right) \rightarrow \min$$

618 Получим задачу минимизации:

$$\min_x L(x) = \min_x \left(- \sum_j \sum_{i \neq j} (W_{jl} \cdot \log S_{jl} + (1 - W_{jl}) \cdot \log(1 - S_{jl})) \right) = \min_{x,y} \sum_j \sum_{l \neq j} (W_{jl} \cdot A_{jl}(S) + (1 - W_{jl}) \cdot R_{jl}(S)),$$

619 где

$$A_{jl} = -\log S_{jl}, R_{jl} = -\log(1 - S_{jl}).$$

Для минимизации по $x(j)$ можно использовать алгоритмы градиентного спуска или стохастического градиентного спуска, в которых итерации заканчиваются, если минимизируемая функция начинает мало меняться.

СПИСОК ЛИТЕРАТУРЫ

1. Gorbunov K. Yu, Lyubetsky V.A. An almost exact linear complexity algorithm of the shortest transformation of chain-cycle graphs //Eprint, Apr 29 2020. arXiv:2004.14351 [math.CO]. <https://doi.org/10.48550/arXiv.2004.14351>
2. Gorbunov K. Yu, Lyubetsky V.A. Linear time additively exact algorithm for transformation of chain-cycle graphs for arbitrary costs of deletions and insertions // Mathematics 2020, V.8. N.11. Art. 2001. P. 1–28. <https://doi.org/10.3390/math8112001>
3. Gorbunov K. Yu, Lyubetsky V.A. Algorithms for the reconstruction of genomic structures with proofs of their low polynomial complexity and high exactness // Mathematics. 2024. V. 12. N. 6, Art. 817. P. 1–26. <https://doi.org/10.3390/math12060817>
4. Weiler P., Lange M., Klein M., Pe'er D. Theis F. CellRank 2: unified fate mapping in multiview single-cell data //Nat. Methods. 2024. V. 21. Art. 1196–1205. P. 1–32 <https://doi.org/10.1038/s41592-024-02303-9>
5. Satija Lab. Seurat - Guided Clustering Tutorial. https://satijalab.org/seurat/articles/pbmc3k_tutorial.html, 2023 (дата обращения: 13 июня 2025).
6. Blondel V. D., Guillaume J.-L., Lambiotte R., Lefebvre E. Fast unfolding of communities in large networks //Journal of Statistical Mechanics: Theory and Experiment, 2008. V.10. P10008. P.1–12. <https://doi.org/10.1088/1742-5468/2008/10/P10008>.
7. Боровков А. А. Математическая статистика. Дополнительные главы. М.: Наука 1984. – 472 с.
8. Гаек Я., Шидак З. Теория ранговых критериев. М.: Наука 1971. – 376 с.
9. Чубисов Д. М. Лекции по асимптотической теории ранговых критериев. Вып. 14 М.: МИАН, 2009. – 186 с.
10. Benjamini Y., Hochberg Y. Controlling the false discovery rate: a practical and powerful approach to multiple testing //Journal of the Royal Statistical Society. Series B (Methodological), 1995. V.57. I.1. P.289–300. <https://doi.org/10.2307/2346101>
11. Hardle W. K., Simar L., Fengler M. R. Applied Multivariate Statistical Analysis. 6th Edition. – Springer, 2024. – 637 p.

	<i>Любецкий Василий Александрович</i>	Поступила в редакцию
	<i>Горбунов Константин Юрьевич</i>	09.09.2025
	<i>Пирогов Сергей Анатольевич</i>	После доработки
	<i>Хазиев Георгий Андреевич</i>	26.09.2023
	<i>Агламазова Анастасия Ильинична</i>	Принята к публикации
651	Институт проблем передачи информации им. А.А. Харкевича Российской академии наук, Москва	20.10.2023
	lyubetsk@iitp.ru	
	gorbunov@iitp.ru	
	pirogov@iitp.ru	
	khaziev@iitp.ru	
	aglamazova.an@iitp.ru	