

Аннотация. Пусть на плоскости проведено несколько прямых, и мы хотим найти число (всех или только ограниченных) областей, на которые они делят плоскость. Если эти прямые в общем положении, то ответ не представляет труда. В общем же случае, когда через некоторые точки проходит много прямых, ответ даёт формула Робертса (1889). Правда, если заметная доля прямых проходит через одну и ту же точку, формула даёт разность двух сравнимых чисел, что можно считать её недостатком.

Это — отправная точка большого пути по конфигурациям прямых (и гиперплоскостей), которому посвящён курс. Мы определим характеристический многочлен конфигурации — и обсудим его связь с хроматическими многочленами. Теорема Т. Заславского (1975) выражает число областей, а также число ограниченных областей, как значения характеристического многочлена в точках -1 и 1 . Такая формула для числа всех клеток максимально эффективна — в ней нет вычитания, только сложение. После того, как формула сформулирована, доказать её — несложное упражнение. Другой вопрос — как до этой формулы или подобных вещей можно догадаться. Мы попробуем предложить ответ на этот вопрос, введя и рассматривая алгебру Варченко–Гельфанда (работа этих авторов вышла в 1987); размерности ее однородных компонент совпадают, с точностью до знака, с коэффициентами характеристического многочлена.

В целом, мы обсудим общие вопросы, которые имеют смысл в пространствах произвольной размерности — в этом случае речь идет о конфигурациях гиперплоскостей — и которые составляют предмет активных современных исследований, но проиллюстрируем их в двумерной ситуации, где основные идеи просты, но все же нетривиальны и интересны. Интересно также увидеть, какие именно геометрические предположения оказываются существенными, а какие нет. Например, вместо прямых на плоскости можно рассмотреть окружности на сфере или на торе. А можно даже заменить прямые кривыми линиями, наложив некоторые условия на их пересечения.

1. КОНФИГУРАЦИИ ПРЯМЫХ НА ПЛОСКОСТИ И СФЕРЕ

Конфигурацией прямых на плоскости $\mathbb{R}^2$ называется просто конечное множество прямых. Прямые конфигурации делят плоскость на конечное число областей, некоторые из которых ограничены, а некоторые простираются до бесконечности. Интересно, например, посчитать число этих областей, в зависимости от конфигурации. Если конфигурация обозначена через $\mathcal{A}$, то мы будем понимать под $\bigcup \mathcal{A}$ объединение всех прямых из $\mathcal{A}$. Дополнение $\mathbb{R}^2 \setminus \bigcup \mathcal{A}$ состоит из конечного числа открытых многоугольных кусков, которые и называются *областями*. Еще у этих многоугольных кусков есть стороны (=границы) и вершины. Будем писать f_0, f_1, f_2 для обозначения количества вершин, ребер и областей конфигурации $\mathcal{A}$, соответственно. Можно пытаться посчитать не только f_2 , то есть число областей, но и числа f_1, f_0 . См. рис. 1.

Разумеется, похожая постановка задачи имеет смысл в произвольной размерности — речь идет о *конфигурации гиперплоскостей*, то есть о конечном наборе гиперплоскостей в пространстве $\mathbb{R}^d$. Про точки пространства $\mathbb{R}^d$ можно думать как про наборы $(x_1, \dots, x_d)$ из d координат, а про *гиперплоскость* в $\mathbb{R}^d$ — как про множество точек, заданное одним уравнением на координаты вида

$$a_1x_1 + \dots + a_dx_d = a_0,$$

где $a_0, \dots, a_d$ — фиксированные действительные числа, причем не все коэффициенты $a_1, \dots, a_d$ равны нулю. Как и в случае прямых, $\bigcup \mathcal{A}$ обозначает объединение всех гиперплоскостей из $\mathcal{A}$, а компоненты множества $\mathbb{R}^d \setminus \bigcup \mathcal{A}$ называются *областями* конфигурации $\mathcal{A}$.¹ Число областей обозначается через f_d (или через $f_d(\mathcal{A})$, если конфигурация $\mathcal{A}$ требует уточнения).

Для любого подпространства $\mathbb{R}^k$, аффинно вложенного в $\mathbb{R}^d$ (то есть при данном вложении каждая координата в $\mathbb{R}^d$ выражается как линейная, не обязательно однородная, функция от координат в $\mathbb{R}^k$), обозначим через $\mathcal{A}|_{\mathbb{R}^k}$ конфигурацию гиперплоскостей в $\mathbb{R}^k$, состоящую из гиперплоскостей в $\mathbb{R}^k$ вида $H \cap \mathbb{R}^k$ при $H \in \mathcal{A}$. Обратите внимание: если $\mathbb{R}^k$ целиком лежит в некоторой гиперплоскости $H \in \mathcal{A}$, то $H \cap \mathbb{R}^k = \mathbb{R}^k$ не входит в $\mathcal{A}|_{\mathbb{R}^k}$, поскольку это пересечение не является гиперплоскостью в $\mathbb{R}^k$. *Плоскость* конфигурации

¹В этом описании мы воспользовались понятием «компонента». Это общее понятие из топологии, которое в данной конкретной ситуации можно определить так: *компонентой* множества $\mathbb{R}^d \setminus \bigcup \mathcal{A}$ называется класс следующего отношения эквивалентности: две точки из $\mathbb{R}^d \setminus \bigcup \mathcal{A}$ эквивалентны, если их можно соединить отрезком, целиком лежащим в $\mathbb{R}^d \setminus \bigcup \mathcal{A}$.

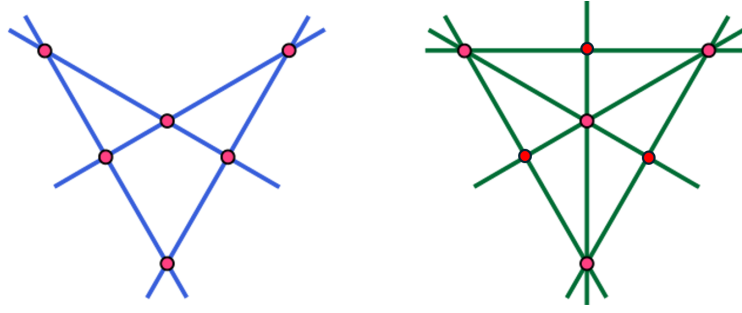


РИС. 1. Примеры конфигураций прямых на плоскости. Для левой конфигурации f -вектор (f_0, f_1, f_2) равен $(6, 16, 11)$; эта конфигурация находится в общем положении. У правой конфигурации f -вектор $(7, 24, 18)$.

fig:fvec

$\mathcal{A}$ — это, по определению, пересечение конечного подсемейства конфигурации $\mathcal{A}$, то есть множество вида $H_1 \cap \dots \cap H_k$ при $H_1, \dots, H_k \in \mathcal{A}$. Все пространство $\mathbb{R}^d$ тоже будем, чисто формально, считать плоскостью конфигурации. *Клетки* конфигурации $\mathcal{A}$ размерности d определяются как компоненты дополнения до $\bigcup \mathcal{A}$, а клетки меньших размерностей можно определить по индукции следующим образом. Если считать, что для конфигураций в размерности $d-1$ клетки всех размерностей уже определены, то k -клетки конфигурации $\mathcal{A}$ в $\mathbb{R}^d$ при $k < d$ определяются как клетки старшей размерности в каждой из конфигураций вида $\mathcal{A}|_L$, где L — плоскость конфигурации $\mathcal{A}$ размерности $< d$. Через $f_k = f_k(\mathcal{A})$ обозначим число всех k -клеток. Вектор с компонентами $f_0, \dots, f_d$ называется f -вектором данной конфигурации. Задача вычисления или оценки компонент f -вектора актуальна в любой размерности d .

1.1. Верхние оценки на f -вектор. Пусть $\mathcal{A}$ — конфигурация n прямых на плоскости. Каковы максимальные возможные значения для чисел f_0, f_1, f_2 в зависимости от n ? Максимальное значение f_0 достигается, очевидно, в том случае, когда любые две прямые из $\mathcal{A}$ пересекаются, и при этом никакие три не проходят через одну и ту же точку. Этот случай называется *случаем общего положения*. В случае общего положения вершин у конфигурации $\mathcal{A}$ столько же, сколько всевозможных неупорядоченных пар прямых, то есть $\binom{n}{2} := \frac{n(n-1)}{2}$. Итак, $f_0 \leq \binom{n}{2}$, причем равенство достигается только в случае общего положения.

Заметим теперь, что каждая прямая $L \in \mathcal{A}$ пересекается максимум с $n-1$ из оставшихся прямых, и образуемые таким образом точки пересечения делят прямую L на максимум n частей. Поскольку всего прямых n , получаем неравенство $f_1 \leq n^2$. Опять, равенство в этом неравенстве достигается только в случае общего положения.

Чтобы получить верхнюю оценку на f_2 , посмотрим, как связаны числа f_2 для конфигураций $\mathcal{A}$ и $\mathcal{A}^* := \mathcal{A} \cup \{L\}$, где L — прямая, не принадлежащая конфигурации $\mathcal{A}$. Мы будем записывать эти числа как $f_2(\mathcal{A})$ и $f_2(\mathcal{A}^*)$, чтобы подчеркнуть зависимость от конфигурации. Предположим, что $\mathcal{A}$ состоит из n прямых. Тогда прямая L делится этими прямыми на не более чем $n+1$ частей, будем называть эти части *звеньями*. С другой стороны, прямая L делит некоторые области конфигурации $\mathcal{A}$ на части; причем каждая область разбивается ровно на две части, если вообще разбивается. Если область разбивается прямой L , то пересечение прямой L с этой областью — это одно из упомянутых звеньев прямой L . Итак, при переходе от $\mathcal{A}$ к $\mathcal{A}^*$ число областей увеличится в точности на число звеньев в прямой L . Отсюда получаем неравенство $f_2(\mathcal{A}^*) \leq f_2(\mathcal{A}) + n + 1$, причем равенство достигается в случае общего положения. Если через $f_2(n)$ обозначить максимальное число областей, на которые n прямые могут делить плоскость, то получается, что $f_2(n+1) = f_2(n) + n + 1$, и что $f_2(n)$ достигается в случае общего положения. Так как

$f_2(1) = 2$, получаем

$$f_2(n) = 2 + 2 + 3 + \dots + n = \binom{n+1}{2} + 1, \quad \forall n \geq 2.$$

Подводя итоги приведенным выше рассуждениям, можно сформулировать следующую теорему.

Теорема. Для всякого фиксированного $n \geq 1$ все компоненты f -вектора конфигурации из n прямых достигают максимальных значений в случае общего положения, и эти максимальные значения равны

$$f_0 = \binom{n}{2}, \quad f_1 = n^2, \quad f_2 = \binom{n+1}{2} + 1.$$

Интересно также оценить компоненты f_k^b ограниченного f -вектора. По определению, f_k^b — это число ограниченных k -клеток рассматриваемой конфигурации. Очевидно, $f_0^b = f_0$ — не бывает неограниченных вершин. Таким образом, оценка сверху на f_0^b — та же самая, что уже была получена выше для f_0 , и она превращается в равенство в точности в случае общего положения. Максимальное значение числа f_1^b , как легко видеть, тоже реализуется в случае общего положения, и оно равно $n(n-2)$ (каждая прямая делится остальными прямыми рассматриваемой конфигурации на звенья, из которых $n-2$ ограниченных). Наконец, максимальное значение величины f_2^b может быть получено двумя способами. Во-первых, можно рассуждать по индукции так же, как и для f_2 . А во-вторых, можно заметить, что максимум опять реализуется в случае общего положения (при случайном шевелении прямых число областей не может уменьшиться), а в этом случае достаточно посчитать число неограниченных областей. Если отойти очень-очень далеко и посмотреть издалека на конфигурацию, то мы не разглядим никаких ограниченных областей — они все визуально сольются в одну точку. Увидим мы только n прямых, проходящих через одну точку, и $2n$ секторов, на которые эти прямые делят плоскость. Каждый из этих секторов представляет собой одну из неограниченных областей (вид издалека, при котором не видно ограниченных ребер, в частности, не видно, что на самом деле неограниченная область может и не быть сектором). Итак, чтобы получить максимальное значение f_2^b , достаточно из максимального значения f_2 вычесть $2n$:

$$\max f_2^b = \binom{n+1}{2} + 1 - 2n = \binom{n-1}{2}$$

См. рис. 2. Следующая теорема резюмирует полученный результат.

Теорема. Для всякого фиксированного $n \geq 1$ все числа f_k^b ($k = 0, 1, 2$) конфигурации из n прямых достигают максимальных значений в случае общего положения, и эти максимальные значения равны

$$f_0^b(n) = \binom{n}{2}, \quad f_1^b(n) = n(n-2), \quad f_2^b(n) = \binom{n-1}{2}.$$

Если конфигурация $\mathcal{A}$ не находится в общем положении, то все равно интересно знать числа f_0, f_1, f_2 ; только теперь эти числа зависят не только от n .

Замечание (Возможные пары (n, f_2)). Можно поставить такую задачу: на сколько частей могут n прямых делить плоскость? Иначе говоря, для данного числа n , каково множество всех возможных значений f_2 ? Выше мы нашли максимальное возможное значение. Самое маленькое возможное значение равно $n+1$; оно достигается в том случае, когда все прямые параллельны. Легко видеть, что следующее значение f_2 , которое можно реализовать конфигурацией из n прямых, равно $2n$, то есть в интервале между $n+1$ и $2n$ ни одно значение не может быть реализовано как f_2 . Это первая лакуна. При больших n возникают и другие лакуны — вторая, третья и т.д. Явные формулы для всех лакун

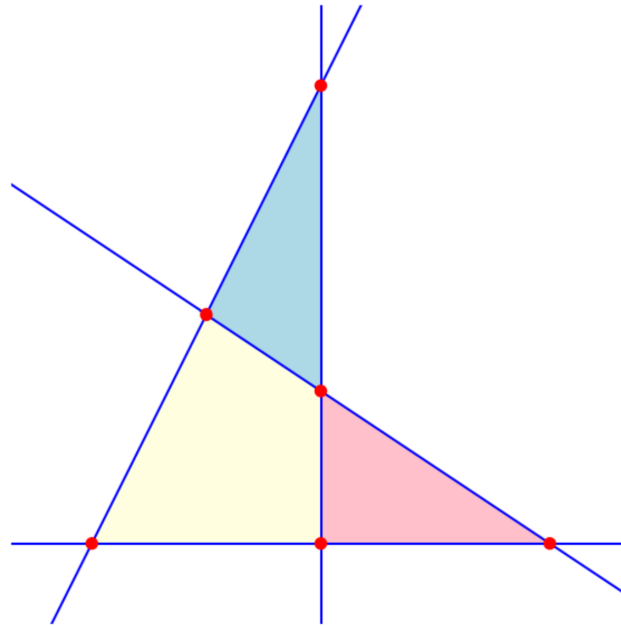


Рис. 2. Конфигурация из 4 прямых общего положения определяет $f_2^b(4) = 3$ ограниченные области.

fig:fb

получены в [7]. Владимир Игоревич Арнольд рассказывал про данную задачу на летней школе «Современная Математика» в 2007 году, его лекция была опубликована в [1].

ss:Eul

1.2. Эйлерова характеристика. Выражение $f_0 - f_1 + f_2$ называется *эйлеровой характеристикой* конфигурации. Как нетрудно видеть из результатов подсчета, проведенного в предыдущем разделе, в случае общего положения эйлерова характеристика вообще не зависит от конфигурации, и даже не зависит от числа прямых. Она всегда равна 1. Это частный случай общей теоремы, состоящей в том, что эйлерова характеристика есть у плоскости как у топологического пространства, а считать ее можно по произвольному разбиению, в частности, по произвольной конфигурации прямых — результат этого вычисления не зависит от выбора разбиения. Таким образом, эйлерова характеристика служит *топологическим инвариантом*. Мы можем проверить инвариантность эйлеровой характеристики для случая конфигураций произвольного вида, не обязательно общего положения; это будет сделано ниже.

Рассмотрим знакопеременную сумму $f_0^b - f_1^b + f_2^b$ только по ограниченным клеткам рассматриваемой конфигурации. В случае общего положения явные формулы для чисел f_k^b , $k = 0, 1, 2$, и непосредственное вычисление приводят к выводу, что эта знакопеременная сумма тоже равна единице. Топологическая интерпретация этого факта состоит в том, что мы посчитали эйлерову характеристику многоугольника, полученного как объединение всех ограниченных клеток (всех размерностей). Это ограниченный замкнутый многоугольник, который не обязан быть выпуклым (приведите пример, когда он невыпуклый!), но у которого нет дырок. Отсутствие дырок очевидно — дырки ограниченного многоугольника не могут быть неограниченными. Можно отсюда вывести (хотя формальное рассуждение потребует некоторого места и некоторой изобретательности), что рассматриваемый многоугольник имеет тот же топологический тип, что и любой выпуклый многоугольник, то есть что и замкнутый диск. Поэтому эйлерова характеристика должна быть равна единице. Итак, для конфигураций общего положения имеют место формулы

$$f_0 - f_1 + f_2 = f_0^b - f_1^b + f_2^b = 1.$$

Теорема. Для произвольной конфигурации $\mathcal{A}$ прямых на плоскости выполнено соотношение $f_0 - f_1 + f_2 = 1$, известное как формула Эйлера.

Доказательство. Мы уже проверили формулу Эйлера в случае общего положения. Теперь рассмотрим произвольную конфигурацию $\mathcal{A}$. Чуть-чуть пошевелим все прямые данной конфигурации, то есть каждую прямую немного повернем относительно случайно выбранного центра, не попадающего в вершину конфигурации $\mathcal{A}$. В результате такого шевеления получим новую конфигурацию, на этот раз общего положения. Остается проследить, что произойдет с эйлеровой характеристикой.

Если $\mathcal{A}$ не находилась в общем положении, то это значит одно из двух. Либо более трех прямых конфигурации сошлись в одной точке, либо есть хотя бы пара прямых, параллельных друг другу. В первом случае $\nu_a > 2$, где a — некоторая вершина конфигурации, а ν_a означает число прямых из $\mathcal{A}$, проходящих через вершину a . А во втором случае $\lambda_b > 1$, где b — некоторое направление (=класс параллельных прямых), а λ_b — число прямых из $\mathcal{A}$, представляющих это направление.

При малом шевелении каждая вершина a со свойством $\nu_a > 2$ порождает несколько маленьких областей, лежащих в окрестности точки a . Конечно, возникают не только новые области, но и новые ограниченные клетки всех размерностей. Чтобы посчитать, сколько клеток каждой размерности возникнет вблизи точки a , нужно забыть про все остальные прямые, не проходившие через a (эти прямые далеко, и в формировании рассматриваемых клеток не участвуют). Можно теперь представить себе такую картину: ν_a прямых, изначально проходивших через одну точку, теперь перешли в общее положение. Сколько получится ограниченных клеток каждой размерности, мы уже посчитали. Влияние этих новых клеток на эйлерову характеристику состоит вот в чем: раньше была только одна точка a , дававшая вклад 1, а потом появилось несколько ограниченных клеток, но посчитанная по ним эйлерова характеристика снова равна 1 (см. выше):

$$f_0^b(\nu_a) - f_1^b(\nu_a) + f_2^b(\nu_a) = 1.$$

Здесь, как и раньше, $f_k^b(n)$ обозначает число ограниченных k -клеток в конфигурации *общего положения* из n прямых. Таким образом, суммарная эйлерова характеристика не менялась.

Осталось посмотреть, не повлияет ли на значение эйлеровой характеристики то, что данное направление b , такое, что $\lambda_b > 1$, разрушится. Как и раньше, можно проигнорировать все прямые, направления которых до шевеления не совпадали с b — эти прямые никак не влияют на возникновение новых клеток за счет разрушения направления b . Прямые, ранее параллельные друг другу и представлявшие направление b , теперь начали пересекаться, но где-то вблизи бесконечности. Таким образом, достаточно рассмотреть такую ситуацию: было $\lambda_b > 1$ параллельных прямых, а получилось столько же прямых общего положения; каким образом менялась эйлерова характеристика? С одной стороны, к f_k добавилось $f_k^b(\lambda_b)$ за счет возникших ограниченных клеток. А с другой стороны, вместо $\lambda_b + 1$ неограниченных 2-клеток, которые мы видели в конфигурации из λ_b прямых, появилось $2\lambda_b$ неограниченных 2-клеток. Еще, вместо λ_b неограниченных 1-клеток, возникло $2\lambda_b$ неограниченных 1-клеток. Эйлерова характеристика в целом поменялась (за счет направления b) на величину

$$f_0^b(\lambda_b) - f_1^b(\lambda_b) + f_2^b(\lambda_b) - \lambda_b - 1 + 2\lambda_b + \lambda_b - 2\lambda_b = 0.$$

Значит, эйлерова характеристика вообще не поменялась, что и требовалось доказать. $\square$

Рассуждение, доказывающее формулу Эйлера, слегка осложняется одним обстоятельством. А именно, шевеление нельзя провести локально вблизи одной конкретной точки. Допустим, что a — *кратная вершина* конфигурации, то есть такая вершина, для которой $\nu_a > 2$. Можно сделать так, чтобы вершина a распалась на *простые вершины* (т.е. такие, через которые проходят ровно по две прямые конфигурации), пошевелив прямые, проходящие через a . Однако это, вообще говоря, нельзя сделать локально, то есть так, чтобы не распались и другие кратные вершины, лежавшие на прямых, которые подвергаются

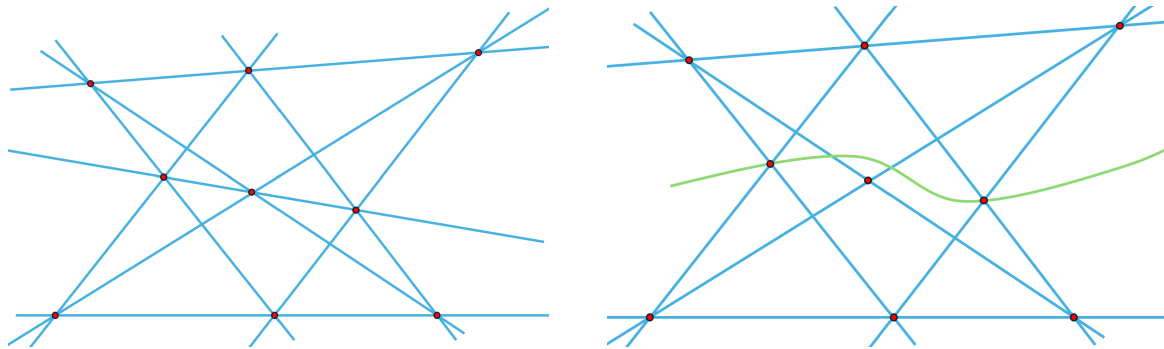


РИС. 3. Слева: конфигурация Паппа (на рисунке отмечены некоторые важные, хотя и не все, вершины этой конфигурации). Справа: конфигурация псевдопрямых, полученная «локальной» деформацией из конфигурации Паппа. Эта конфигурация не реализуется настоящими прямыми ни над каким полем, но все наши утверждения про f -вектор к ней применимы.

fig:papp

шевелению. Распадение кратных вершин и кратных направлений происходит, вообще говоря, одновременно. Впрочем, можно проследить за теми клетками, которые образуются при распадении каждого конкретного элемента. Например, если распалась кратная вершина a , то новые клетки образовались вблизи a , но далеко от других вершин и далеко от бесконечности. А если распалось кратное направление b , то образовались новые клетки около бесконечности в направлении b , а это далеко как от всех вершин, так и от точек вблизи бесконечности по другим направлениям.

Замечание. Имеет смысл рассматривать конфигурации более общего вида. Например, вместо прямых можно взять простые кривые на плоскости, уходящие обеими концами на бесконечность. Будем называть конфигурацию таких кривых *конфигурацией псевдопрямых*, а сами кривые — *псевдопрямыми*, если любые две различные псевдопрямые пересекаются максимум по одной точке, а отношение «не иметь общих точек или совпадать» является на данном конечном множестве псевдопрямых отношением эквивалентности. Все доказанные выше свойства конфигураций прямых переносятся дословно на псевдопрямые, при этом некоторые рассуждения даже упрощаются. Например, при доказательстве формулы Эйлера достаточно произвести локальные модификации псевдопрямых, меняя их только в маленькой окрестности кратной вершины, см. рис. 3.

Аналог формулы Эйлера для ограниченного f -вектора также имеет место.

Теорема. Для произвольной конфигурации (псевдо)прямых выполнено тождество $f_0^b - f_1^b + f_2^b = 1$. Это формула Эйлера для ограниченного f -вектора.

Доказательство аналогично проведенному выше, и поэтому мы не будем его приводить. Отметим, однако, топологическую интерпретацию полученной формулы. Левая часть формулы представляет собой эйлерову характеристику объединения всех ограниченных клеток конфигурации. Это объединение может не быть выпуклым многоугольником, но оно обязано иметь тот же топологический тип. Иначе говоря, в объединении ограниченных клеток не может быть «дырок» — что интуитивно очевидно, поскольку всякая дырка была бы ограниченной клеткой, и ее надо было бы добавить.

1.3. Формула Робертса (1889). Формула Робертса 1889 года позволяет посчитать f -вектор конфигурации прямых, не находящихся в общем положении. Разумеется, в этом случае мы не сможем выразить все компоненты f -вектора только через число n , то есть через число прямых конфигурации. Нам нужно знать еще ряд чисел, а именно, кратности ν_a и λ_b , которые нам уже встретились при доказательстве формулы Эйлера. Вспомним, что ν_a — это *порядок* вершины a , то есть число прямых данной конфигурации, проходящих через a , а число λ_b считает количество прямых *данного направления* b , то есть в данном

классе b по отношению параллельности (отношению эквивалентности, относительно которого две прямые эквивалентны, если они параллельны или совпадают). Вот формула Робертса:

$$f_2 = 1 + \binom{n+1}{2} - \sum_a \binom{\nu_a - 1}{2} - \sum_b \binom{\lambda_b}{2}.$$

Правая часть выписанной формулы выглядит как $f_2(n)$ минус поправка, где $f_2(n)$ — максимальное значение числа f_2 , реализующееся в случае общего положения (это число зависит только от n). А поправка учитывает тот факт, что данная конфигурация может не находиться в общем положении. Эта поправка всегда отрицательна или (в случае общего положения) равна нулю. Смысл поправки вот в чем: когда мы приведем конфигурацию в общее положение, слегка пошевелив все прямые, возникнет определенное число новых клеток. Число новых клеток и есть поправка, в чем мы сейчас убедимся.

Рассмотрим вершину a нашей конфигурации, в которой сходятся $\nu_a > 1$ прямых. После приведения в общее положение, на месте вершины a возникнет $f_2^b(\nu_a) = \binom{\nu_a - 1}{2}$ ограниченных областей. Тем самым объясняется первая сумма, фигурирующая в поправке. Осталось рассмотреть направление b , такое, что $\lambda_b > 1$. В результате приведения конфигурации в общее положение появится $f_2^b(\lambda_b)$ новых ограниченных областей и еще $\lambda_b - 1$ новых неограниченных областей (см. доказательство формулы Эйлера из раздела 1.2). Итого, за счет разрушения направления b добавится $\binom{\lambda_b - 1}{2} + \lambda_b - 1 = \binom{\lambda_b}{2}$ клеток. Это объясняет вторую сумму в поправке. Таким образом, каждое слагаемое в поправке получило объяснение как число клеток, образовавшихся в результате приведения конфигурации в общее положение, либо за счет разрушения кратной вершины, либо за счет разрушения кратного направления. Формула Робертса тем самым доказана.

Замечание. Можно заменить, что группы из ν_a пересекающихся в общей точке прямых и группы из λ_b параллельных друг другу прямых играют похожие роли в формуле Робертса. Хочется сказать, что эти λ_b прямых «пересекаются в одной точке на бесконечности». Тем не менее, на обычной, так называемой *аффинной*, плоскости, никаких бесконечно удаленных точек нет, и два случая необходимо рассматривать отдельно. Ситуация упрощается при переходе от аффинной плоскости к *проективной плоскости*. Можно получить проективную плоскость из аффинной плоскости добавлением некоторого множества точек, которое называется *прямой на бесконечности*. Всякая прямая будет пересекать прямую на бесконечности, причем параллельные прямые будут проходить через одну и ту же бесконечно удаленную точку. Проективная геометрия, то есть геометрия проективной плоскости, оказывается в некоторых отношениях проще аффинной геометрии. Например, любые две различные прямые на проективной плоскости пересекаются ровно в одной точке. Кроме того, все прямые выглядят одинаково, и в этом смысле бесконечно удаленная прямая ничем не отличается от остальных прямых. Любую прямую можно перевести в любую другую *проективным преобразованием*, то есть преобразованием, переводящим прямые в прямые. Формула Робертса на проективной плоскости упрощается — в ней останутся только члены с ν_a , правда, нужно будет рассматривать в том числе и те вершины конфигурации, которые находятся на бесконечности (мы считаем, что конфигурация содержит бесконечно удаленную прямую).

1.4. Конфигурации на сфере. *Двумерная сфера* S^2 может быть реализована в трехмерном евклидовом пространстве $\mathbb{R}^3$ с координатами x, y, z как множество точек, координаты которых удовлетворяют уравнению $x^2 + y^2 + z^2 = 1$. Плоскость $\mathbb{R}^2$ с координатами x, y можно вложить в $\mathbb{R}^3$ как множество точек вида $(x, y, -1)$, то есть как горизонтальную плоскость, заданную уравнением $z = -1$. Эта плоскость касается сферы S^2 в ее *южном полюсе* $(0, 0, -1)$.

Конфигурация больших окружностей на сфере S^2 — это просто конечное множество больших окружностей, то есть таких окружностей на сфере, центр которых совпадает с центром сферы, см. рис. 4. Имеется тесная связь между конфигурациями на сфере и

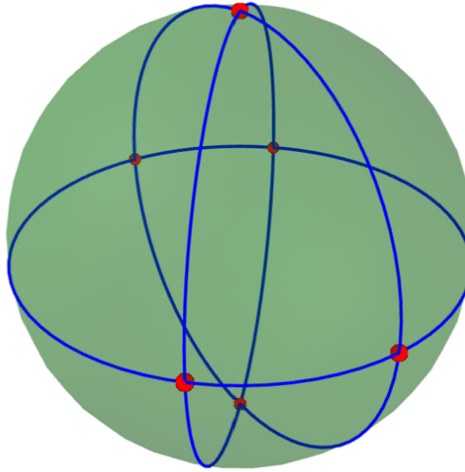


Рис. 4. Конфигурация больших окружностей на сфере.

fig:sph

конфигурациями на плоскости. В частности, для всякой конфигурации $\mathcal{A}$ прямых на плоскости можно определить соответствующую конфигурацию $\mathcal{A}^S$ на сфере S^2 , при помощи центральной проекции. А именно, реализуем плоскость $\mathbb{R}^2$ как описано выше, то есть будем считать ее касательной плоскостью к S^2 в южном полюсе. Каждой прямой $L \subset \mathbb{R}^2$ можно сопоставить плоскость P_L , проходящую через L и через точку $(0, 0, 0)$, т.е. через начало координат в $\mathbb{R}^3$, а по совместительству центр сферы S^2 . Пересечение $P_L \cap S^2$ представляет собой большую окружность на сфере, а если L пробегает все прямые из $\mathcal{A}$, то соответствующие большие окружности $P_L \cap S^2$ образуют некоторую сферическую конфигурацию $\mathcal{A}^S$. Очевидно, любая сферическая конфигурация, не содержащая экватора, может быть получена как $\mathcal{A}^S$ для подходящей конфигурации прямых $\mathcal{A}$. Если конфигурация на сфере содержит экватор, то это неприятное обстоятельство можно исправить незначительным вращением сферы относительно некоторой (невертикальной) оси.

Конфигурация $\mathcal{A}^S$ делит сферу на f_2^S областей, f_1^S дуг и f_0^S точек. Набор чисел (f_0^S, f_1^S, f_2^S) называется *f-вектором сферической конфигурации* $\mathcal{A}^S$. Обсудим связь между этим сферическим *f*-вектором и *f*-вектором исходной конфигурации прямых $\mathcal{A}$.

Предложение. Компоненты *f*-вектора сферической конфигурации $\mathcal{A}^S$ выражаются формулами

$$f_2^S = f_2 + f_2^b, \quad f_1^S = f_1 + f_1^b, \quad f_0^S = f_0 + f_0^b = 2f_0,$$

или, в векторных обозначениях, $f^S = f + f^b$.

Здесь f^b — это ограниченный *f*-вектор, а f — просто *f*-вектор конфигурации $\mathcal{A}$.

Доказательство. Посчитаем, например, число f_2^S областей сферической конфигурации. Ограниченной области на плоскости (имеется в виду область конфигурации $\mathcal{A}$) соответствует две области на сфере — эти две области антиподальны друг другу. С другой стороны, проецируя на сферу неограниченную область, мы получаем не сферическую область целиком, а только ее часть, доходящую до экватора. Вторую часть той же сферической области можно получить, проецируя другую неограниченную область на плоскости. Две плоские области $D_{\pm}$, чьи проекции образуют две разделенные экватором части одной и той же сферической области, «противоположны» друг другу в следующем смысле: неограниченные 1-клетки на границе обеих областей принадлежат одной и той же паре прямых. Если отойти очень далеко, то $D_{\pm}$ будут видны как два противоположных сектора из четырех, на которые две прямые делят плоскость. Как и в случае с ограниченной областью, из $D_+ \cup D_-$ возникает не одна, а две антиподальные друг другу проекции на сфере. Таким образом, каждая ограниченная плоская область дает две сферические области, а каждая пара неограниченных плоских областей дает пару сферических областей. Иначе говоря,

ограниченная плоская область вносит вклад $+2$ в число f_2^S , а неограниченная — вклад $+1$. Заметим, что точно такой же вклад они вносят в сумму $f_2 + f_2^b$.

Вычисление компоненты f_1^S вполне аналогично проведенному выше. Ситуация с f_0^S проще, поскольку каждая 0-клетка конфигурации $\mathcal{A}$ ограничена, и каждая дает две 0-клетки конфигурации $\mathcal{A}^S$. $\square$

Сферические конфигурации удовлетворяют варианту формулы Эйлера, а именно

$$f_0^S - f_1^S + f_2^S = 2.$$

Будем называть эту формулу *сферической формулой Эйлера*; ее можно вывести из доказанного выше предложения и формул Эйлера для f -векторов f и f^b . Топологическая интерпретация сферической формулы Эйлера состоит в том, что эйлерова характеристика 2-сферы равна двум. Если поверить в инвариантность и аддитивность эйлеровой характеристики, то это устанавливается совсем просто — сфера разбивается на точку и открытый диск, причем у каждого из этих двух подмножеств эйлерова характеристика равна единице.

1.5. Свойство Сильвестра–Галлаи. В 1893 году Джеймс Джозеф Сильвестр поставил такую задачу: верно ли, что среди любого конечного множества точек на плоскости, не лежащих на одной прямой, найдется пара точек, такая, что проходящая через них прямая не содержит никаких других точек данного множества? (Такая прямая, если она существует, называется *прямой Сильвестра*). Несмотря на элементарную формулировку, задача оставалась открытой 40 лет. Возможно, этой задачей просто никто не занимался. В 1933 году известный венгерский математик Эрдеш переоткрыл задачу Сильвестра, и после нескольких неудачных попыток ее решить, сообщил ее своему коллеге Тибору Грюнвальду (позже Грюнвальд сменил свою фамилию на Галлаи; он более известен под этой второй фамилией). Галлаи вскоре решил задачу. Однако широкую известность задача получила еще через 10 лет, в 1943 году, когда Эрдеш опубликовал ее в популярном американском математическом журнале *American Mathematical Monthly*. Одновременно с задачей, в редакцию было представлено и решение, полученное Галлаи. Вскоре в редакцию поступило еще несколько решений, полученных Баком, Келли, Штейнбергом и Стинродом. Как выяснилось впоследствии, еще до публикации Эрдеша в *Monthly*, решение задачи было опубликовано немецким математиком Мельхиором. Решение Мельхиора основано на той связи между плоскими и сферическими конфигурациями, которую мы обсудили выше. Сначала сформулируем результат для сферических конфигураций, а потом обсудим его связь с другими формулировками проблемы Сильвестра.

Теорема. Пусть на сфере проведено конечное число больших окружностей, не все из которых проходят через одну и ту же пару антиподальных точек. Тогда найдется точка, принадлежащая ровно двум из проведенных окружностей.

Доказательство. Обозначим через $\mathcal{A}^S$ сферическую конфигурацию, состоящую из всех проведенных окружностей. Рассуждая от противного, допустим, что через каждую вершину конфигурации $\mathcal{A}^S$ проходит по меньшей мере три прямые конфигурации. Это означает, что каждая вершина инцидентна не менее чем шести ребрам. Поскольку каждое ребро инцидентно ровно двум вершинам, получаем неравенство $f_1^S \geq 3f_0^S$. С другой стороны, каждая область конфигурации $\mathcal{A}^S$ представляет собой сферический многоугольник, у которого не менее трех сторон (если бы было только две стороны, то эти две стороны соединяли бы одну и ту же пару антиподальных точек, и вообще все прямые конфигурации проходили бы через эту пару). А так как каждое ребро инцидентно ровно двум областям, получаем неравенство $f_1^S \geq \frac{3}{2}f_2^S$.

Теперь воспользуемся соотношением Эйлера $f_0^S - f_1^S + f_2^S = 2$ и подставим в него полученные неравенства:

$$2 = f_0^S - f_1^S + f_2^S \leq \frac{1}{3}f_1^S - f_1^S + \frac{2}{3}f_1^S = 0.$$

Мы получили абсурдное неравенство $2 \leq 0$; полученное противоречие завершает доказательство. $\square$

Напомним понятие сферической двойственности. Пара антиподальных точек $a^\pm$ на сфере определяет так называемую *двойственную окружность*, а именно, пересечение сферы с плоскостью, перпендикулярной отрезку $[a^-; a^+]$ и проходящей через его середину. И наоборот, каждая большая окружность на сфере определяет *двойственную пару антиподальных точек*, а именно, пересечение сферы с прямой, перпендикулярной плоскости окружности и проходящей через центр окружности. Если пара антиподальных точек лежит на большой окружности, то двойственная (к данной паре антиподальных точек) окружность содержит двойственную (к данной окружности) пару антиподальных точек. *Сферическая двойственность* основана на этом утверждении; она позволяет переходить от утверждений про большие окружности и пары антиподальных точек к аналогичным утверждениям про двойственные объекты, если только данные утверждения используют лишь отношение инцидентности: такая-то пара точек лежит на такой-то окружности. Вот какое утверждение получается из доказанной теоремы при помощи сферической двойственности.

Следствие. Пусть на сфере отмечено конечное число пар антиподальных точек. Если не все эти пары принадлежат одной и той же большой окружности, то найдется большая окружность, проходящая через ровно две из отмеченных пар.

Отсюда уже несложно вывести исходную формулировку Сильвестра, рассмотрев центральную проекцию из плоскости на сферу (ту же самую проекцию, которую мы обсуждали выше именно в связи с переходом от плоской конфигурации к сферической).

Следствие. Пусть на вещественной плоскости отмечено конечное число точек. Если не все из отмеченных точек лежат на одной прямой, то найдется прямая, содержащая ровно две из отмеченных точек.

С другой стороны, из самой исходной теоремы, а не ее двойственного аналога, при помощи центральной проекции на сферу получается следующее утверждение — так называемая *двойственная теорема Сильвестра–Галлаи*.

Следствие. Пусть на вещественной плоскости проведено конечное число прямых. Если не все эти прямые проходят через одну точку и не все параллельны, то найдется точка, в которой пересекаются ровно две из проведенных прямых, или направление, которому параллельны ровно две из проведенных прямых.

Прямые, параллельные данному направлению, можно понимать неформально как прямые, проходящие через заданную точку на бесконечности. Эту интуицию можно довести до строгих формулировок, используя понятие *проективной плоскости*.

Интересно отметить, что свойство Сильвестра–Галлаи специфично для вещественных конфигураций. Можно рассматривать конфигурации комплексных прямых в комплексном двумерном пространстве $\mathbb{C}^2$, и для них это свойство уже не будет верно. Прямые в $\mathbb{C}^2$ — это множества точек, заданные линейными уравнениями на две координаты z, w , то есть уравнениями вида $az + bw + c = 0$, в которых a, b, c — постоянные коэффициенты, зависящие от выбора прямой, но не зависящие от z и w .

Рассмотрим, например, следующие 9 прямых в $\mathbb{C}^2$:

$$z = \zeta^k, \quad w = \zeta^l, \quad w = \zeta^m z, \quad k, l, m = 0, 1, 2.$$

Здесь $\zeta = e^{2\pi i/3}$ — примитивный корень третьей степени из единицы. Конфигурацию комплексных прямых удобно записывать одним уравнением $Q(z, w) = 0$, в левой части которого стоит произведение линейных функций, задающих прямые конфигурации. Легко видеть, что уравнение $Q(z, w) = 0$ в таком случае определяет объединение всех рассматриваемых

прямых, то есть сильно вырожденную кривую степени n . Такой подход к заданию конфигураций бывает полезен и в общем случае, а в нашем конкретном случае многочлен Q имеет особенно простой вид:

$$Q(z, w) = (z^3 - 1)(w^3 - 1)(w^3 - z^3).$$

Мы утверждаем, прямая через любая вершина рассматриваемой конфигурации $\mathcal{A}^{\mathbb{C}}$ инцидентна по меньшей мере трем прямым из конфигурации. Это же относится к бесконечно удаленным вершинам, то есть если есть две параллельные прямые из $\mathcal{A}^{\mathbb{C}}$, то им же параллельна и какая-то третья прямая из $\mathcal{A}^{\mathbb{C}}$. Таким образом, для $\mathcal{A}^{\mathbb{C}}$ нарушается двойственная теорема Сильвестра–Галлаи (мораль: данная теорема существенно использует то, что речь идет именно о вещественной плоскости).

Рассмотрим сначала конечную точку пересечения двух прямых из $\mathcal{A}^{\mathbb{C}}$. Эта точка имеет вид (ζ^k, ζ^l) для некоторой пары $k, l \in \{0, 1, 2\}$. Но через данную точку проходят следующие три различные прямые конфигурации:

$$z = \zeta^k, \quad w = \zeta^l, \quad w = \zeta^{l-k}z.$$

Осталось рассмотреть направления, представленные хотя бы двумя прямыми из $\mathcal{A}^{\mathbb{C}}$. Таких направления два — $z = \text{const}$ и $w = \text{const}$ — и каждое из этих направлений представлено тремя прямыми.

Замечание. Можно показать, что кривая третьей степени в $\mathbb{C}^2$, коэффициенты которой выбраны достаточно случайно, имеет ровно 9 точек перегиба, причем прямая, соединяющая любую пару из этих точек, проходит и через какую-либо третью точку. Тем самым, в $\mathbb{C}^2$ нарушается исходная формулировка теоремы Сильвестра–Галлаи. Более явный пример можно построить, исходя из рассмотренной выше конфигурации, с использованием двойственности.

1.6. Оценка снизу. Следующая простая оценка снизу на f_0 , зависящая только от n , принадлежит де Брейну и Эрдешу.

Теорема (de Bruijn–Erdős). *Для всякой конфигурации из $n \geq 3$ прямых на вещественной плоскости выполнено неравенство $f_0 \geq n$, если только не все прямые рассматриваемой конфигурации проходят через одну и ту же точку, и ни какие две прямые не параллельны друг другу.*

Доказательство. Приводимое ниже доказательство следует оригинальному рассуждению де Брейна и Эрдеша. Это рассуждение работает только над вещественными числами, поскольку использует теорему Сильвестра–Галлаи. Если $n = 3$, утверждение очевидно. Далее будем рассуждать по индукции, предположив, что для $n - 1$ прямых искомое утверждение уже доказано. Рассмотрим конфигурацию $\mathcal{A}$ из n прямых, не проходящих через одну и ту же точку и не параллельных друг другу. По двойственной теореме Сильвестра–Галлаи, найдется такая вершина v конфигурации, через которую проходит только две прямые из $\mathcal{A}$. Обозначим эти прямые через L и L' .

Чтобы воспользоваться предположением индукции, мы рассмотрим конфигурацию $\mathcal{A} \setminus \{L\}$. Если так случится, что в этой конфигурации все прямые проходят через одну и ту же точку, то мы вместо прямой L выкинем прямую L' , а все дальнейшее рассуждение проходит без изменений. Конфигурация $\mathcal{A}' := \mathcal{A} \setminus \{L\}$ состоит из $n - 1$ прямых, а следовательно, по предположению индукции, у этой конфигурации не менее $n - 1$ вершин. Все эти $n - 1$ точек являются также вершинами конфигурации $\mathcal{A}$. Но кроме того, у конфигурации $\mathcal{A}$ имеется вершина v , которая не есть вершина для $\mathcal{A}'$, поскольку L' — единственная прямая из $\mathcal{A}'$, содержащая точку v . Итак, у $\mathcal{A}$ не менее n вершин, что и требовалось доказать. $\square$

Рассуждение де Брейна и Эрдеша очень изобретательно и очень геометрично. Оно, однако, работает только над вещественными числами. Это последнее ограничение можно снять: сама оценка, как выясняется, верна над любым полем и даже, более общим образом,

в широком классе комбинаторных геометрий. Более комбинаторное и более абстрактное рассуждение, которое это доказывает, принадлежит Герберту Райзеру.

Доказательство Герберта Райзера. Рассмотрим конфигурацию $\mathcal{A}$ из n прямых на плоскости. Матрицей инцидентности для $\mathcal{A}$ называется прямоугольная таблица, строки которой отвечают прямым из $\mathcal{A}$, а столбцы которой отвечают вершинам конфигурации $\mathcal{A}$. На пересечении строки с номером i и столбца с номером j стоит 1, если i -ая прямая конфигурации содержит j -ую вершину, и 0 в противном случае. Мы, таким образом, считаем, что все прямые конфигурации перенумерованы, например, как $L_1, \dots, L_n$, и все вершины конфигурации тоже перенумерованы, как $p_1, \dots, p_m$, где $m = f_0$.

Обозначим i -ю строчку матрицы инцидентности через $\vec{l}_i$. Это строчка из нулей и единиц, причем единицы стоят ровно в тех столбцах, которые соответствуют точкам $p_j \in L_i$. Для всяких двух строчек $\vec{a} = (a_1, \dots, a_m)$ и $\vec{b} = (b_1, \dots, b_m)$ одинаковой длины m , состоящих из вещественных чисел, определено скалярное произведение

$$\vec{a} \cdot \vec{b} := \sum_{i=1}^m a_i b_i.$$

Оценим скалярные произведения двух строк матрицы инцидентности. Возьмем сначала две разные строки; они соответствуют двум разным прямым конфигурации. Поскольку две разные прямые имеют ровно одну точку пересечения, получаем, что скалярное произведение $\vec{l}_i \cdot \vec{l}_k$ равно единице при $i \neq k$. Если же $i = k$, то скалярное произведение $\vec{l}_i \cdot \vec{l}_i$, очевидно, равно числу вершин на прямой L_i . Это число не может быть меньше двойки при $n \geq 3$, согласно сделанным допущениям, что любая пара прямых пересекается и что не все прямые конфигурации проходят через одну и ту же вершину.

Мы утверждаем, что строки $\vec{l}_i$ линейно независимы, то есть соотношение вида $\alpha_1 \vec{l}_1 + \dots + \alpha_n \vec{l}_n = 0$, где α_i — действительные числа, возможно только при $\alpha_1 = \dots = \alpha_n$. Умножить строку на действительно число — значит, умножить каждый элемент данной строки на это самое (одно и то же) число, а складывать строки следует поэлементно, то есть

$$\alpha(a_1, \dots, a_m) + \beta(b_1, \dots, b_m) := (\alpha a_1 + \beta b_1, \dots, \alpha a_m + \beta b_m).$$

Рассмотрим строку вида $\vec{l} = \alpha_1 \vec{l}_1 + \dots + \alpha_n \vec{l}_n$ и посчитаем скалярное произведение этой строки с собой:

$$\vec{l} \cdot \vec{l} = \sum_{i=1}^n \alpha_i^2 (\vec{l}_i \cdot \vec{l}_i - 1) + \sum_{i,k=1}^n \alpha_i \alpha_k = \sum_{i=1}^n \alpha_i^2 (\vec{l}_i \cdot \vec{l}_i - 1) + (\alpha_1 + \dots + \alpha_n)^2 \geq \sum_{i=1}^n \alpha_i^2.$$

Отсюда следует, что если $\vec{l} = 0$, то тогда все коэффициенты α_i обязаны быть равны нулю.

Теперь утверждение теоремы вытекает из следующего общего факта, относящегося к линейной алгебре: если строки матрицы линейно независимы, то у этой матрицы столбцов не меньше чем строк. Иными словами: если система линейных уравнений с нулевыми правыми частями имеет только нулевое решение, то число неизвестных не превышает числа уравнений. $\square$

Доказательство Райзера не только работает над любым полем; оно годится также и для конфигураций псевдопрямых.

Замечание. Нижняя оценка $f_0 \geq n$ допускает глубокое обобщение на высшие размерности. Это так называемая гипотеза Даулинга–Вильсона, относящаяся к конфигурациям гиперплоскостей в $\mathbb{R}^d$. Обозначим через W_{k+1} число k -мерных плоскостей такой конфигурации. Гипотеза Даулинга–Вильсона утверждает, что $W_k \geq W_{d+1-k}$, если $2k \leq d+1$. В случае $d=2$ получаем $W_1 \geq W_2$, что совпадает с неравенством $f_0 \geq n$, поскольку $W_1 = f_0$ (число вершин) и $W_2 = n$ (число прямых). Хо Джун И (June Huh) и Ван Ботон (Boton Wang) доказали гипотезу Даулинга–Вильсона в 2017 [6], используя глубокие результаты

из алгебраической геометрии. Это один результат из целой серии, за которую Хо Джун И получил в 2022 филдсовскую медаль.

2. ХАРАКТЕРИСТИЧЕСКИЙ МНОГОЧЛЕН

В этом разделе мы обсудим конфигурации гиперплоскостей произвольной размерности и связанные с этими конфигурациями многочлены. Начнем мы, впрочем, с двумерного случая, а именно, с более экономным (по сравнению с формулой Робертса) выражением для числа областей f_2 .

2.1. Формулы Заславского (1975) на плоскости. Формула Робертса имеет следующий недостаток. Правая часть этой формулы представляется как разность двух неотрицательных чисел. Иногда бывает так, что оба эти числа намного больше, чем их разность, то есть f_2 . В этих случаях вычислительная эффективность формулы Робертса вызывает сомнения. С другой стороны, имеется более эффективная формула — даже самая эффективная в том смысле, что правая часть этой формулы сводится к суммированию неотрицательных чисел. Никаких вычитаний! Эффективная формула была получена Томасом Заславским в 1975. Вообще-то Заславский сформулировал свой результат сразу в любой размерности, но мы пока обсудим только двумерный случай. Как часто бывает, самое сложное — угадать ответ. После того, как ответ выписан, доказать его — несложное упражнение.

Теорема (Zaslavsky [8]). *Число областей конфигурации из n прямых выражается формулой*

$$f_2 = 1 + n + \sum_a (\nu_a - 1).$$

Здесь мы пользуемся ранее введенными обозначениями. А именно, сумма в правой части распространяется на все вершины a рассматриваемой конфигурации, а ν_a — количество прямых конфигурации, проходящих через вершину a .

Доказательство. Если $n = 1$, то, очевидно, и в левой, и в правой частях искомого равенства стоит двойка. Достаточно проследить за тем, что произойдет с обеими частями равенства при добавлении к рассматриваемой конфигурации $\mathcal{A}$ одной прямой L , то есть при переходе от $\mathcal{A}$ к $\mathcal{A} \cup \{L\}$. К левой части равенства добавится количество областей конфигурации $\mathcal{A}$, через которые проходит прямая L . А к правой части добавится единица (за счет замены n на $n + 1$) и количество точек пересечения прямой L с прямыми конфигурации $\mathcal{A}$. Осталось только заметить, что пересечения прямой L со всеми областями конфигурации $\mathcal{A}$ в точности совпадают с компонентами дополнения в L до точек пересечения с $\bigcup \mathcal{A}$, а t точек делят прямую на $t + 1$ компонент. $\square$

Еще одна теорема Заславского вычисляет количество ограниченных областей. В формуле встречаются минусы, но, по всей видимости, без них нельзя обойтись.

Теорема. *Число ограниченных областей конфигурации из n прямых выражается формулой $f_2^b = 1 - n + \sum_a (\nu_a - 1)$.*

Доказательство аналогично приведенному выше. Подчеркнем еще раз, что в такого рода результатах главное — додуматься до правильной формулировки. А после того, как результат сформулирован, доказать его не составляет труда. Все же хотелось бы понять также и то, как до такого результата додуматься. Мы предложим соответствующую концепцию, отличающуюся, впрочем, от авторской.

2.2. Меры и характеристический многочлен. Теперь рассмотрим конфигурацию $\mathcal{A}$ гиперплоскостей в пространстве $\mathbb{R}^d$ произвольной размерности. Все гиперплоскости из $\mathcal{A}$ и все множества, которые можно получить из пространства $\mathbb{R}^d$ и всех данных гиперплоскостей применением стандартных теоретико-множественных операций, образуют алгебру множеств $\mathbb{A}_{\mathcal{A}}$. Нетрудно видеть, что всякое множество из $\mathbb{A}_{\mathcal{A}}$ можно представить в виде объединения непересекающихся множеств вида $L^\circ := L \setminus \bigcup \mathcal{A}|_L$, где L — плоскость конфигурации $\mathcal{A}$ (напомним, что $\mathcal{A}|_L$ — это конфигурация гиперплоскостей в L , состоящая из пересечений $L \cap H$ плоскости L со всеми гиперплоскостями $H \in \mathcal{A}$ со свойством $L \not\subset H$).

Задача. Более общим образом, алгебра подмножеств произвольного множества S , порожденная подмножествами $A_1, \dots, A_n \subset S$, состоит из всевозможных объединений *атомарных множеств*, то есть множеств вида $A_1^{\varepsilon_1} \cap A_2^{\varepsilon_2} \cap \dots \cap A_n^{\varepsilon_n}$, где $\varepsilon_i \in \{\pm\}$, а под $A^\pm$ следует понимать множества $A^+ := A$ и $A^- := S \setminus A$.

Напомним, что *мерой* на $\mathbb{A}_{\mathcal{A}}$ называется такая функция μ на $\mathbb{A}_{\mathcal{A}}$, которая удовлетворяет условию аддитивности:

$$\mu(A \cup B) + \mu(A \cap B) = \mu(A) + \mu(B), \quad \forall A, B \in \mathbb{A}_{\mathcal{A}}.$$

Мера может принимать значения в действительных числах, или в каких-либо еще числах, или, более общим образом, во множестве, допускающем операцию сложения. В дальнейшем мы будем иметь дело с мерами, принимающими значения во множестве $\mathbb{Z}[q]$ всех многочленов с целыми коэффициентами от переменной q .

Лемма. *Существует единственная мера μ на $\mathbb{A}_{\mathcal{A}}$ со значениями в $\mathbb{Z}[q]$, такая что $\mu(L) = q^k$ для всякой плоскости L размерности k .*

Доказательство леммы мы приведем чуть позже, а пока, пользуясь леммой, дадим определение характеристического многочлена.

Определение (Характеристический многочлен). Пусть μ — та мера на $\mathbb{A}_{\mathcal{A}}$, существование и единственность которой утверждается в лемме. Значение $\mu(\mathbb{R}^d \setminus \bigcup \mathcal{A})$ является членом степени d от q . Этот многочлен обозначается через $\chi(q)$ и называется *характеристическим многочленом* конфигурации $\mathcal{A}$. Когда надо будет подчеркнуть зависимость характеристического многочлена от $\mathcal{A}$, будем писать $\chi_{\mathcal{A}}(q)$ вместо $\chi(q)$.

Теперь докажем лемму.

Доказательство леммы. Имеет место следующая формула включения–исключения:

$$\mu(A_1 \cup \dots \cup A_n) = \sum_i \mu(A_i) - \sum_{i < j} \mu(A_i \cap A_j) + \sum_{i < j < k} \mu(A_i \cap A_j \cap A_k) \mp \dots$$

Все индексы во всех суммах правой части пробегают целые значения от 1 до n . Формулу включения–исключения нетрудно вывести из свойства аддитивности, например, индукцией по n . Для каждой плоскости L конфигурации $\mathcal{A}$ формула включения–исключения дает меру множества $\bigcup \mathcal{A}|_L$, а следовательно, и множества L° . Тем самым доказана единственность. Существование основано на том, что выражение для $\mu(L^\circ)$, полученное по формуле включения–исключения, действительно удовлетворяет свойству аддитивности — в этом можно убедиться несложной и непосредственной проверкой. $\square$

Обсудим комбинаторный смысл характеристического многочлена. Конфигурации гиперплоскостей имеют смысл не только над действительными числами, но и над любым множеством, элементы которого можно складывать и умножать, причем сложение и умножение обладают хотя бы некоторыми из привычных свойств, например, коммутативностью и ассоциативностью. Пример такого множества — *кольцо вычетов* $\mathbb{Z}_q$ по модулю q . Элементы кольца $\mathbb{Z}_q$ — всевозможные остатки от деления на q , а арифметические операции в $\mathbb{Z}_q$ определяются следующим образом: берем два остатка по модулю q , складываем их или перемножаем как целые числа, а потом у получившегося результата берем остаток

по модулю q . Часто бывает так, что число элементов в каждой плоскости размерности k равно q^k (*размерность* плоскости определяется как $d - m$, где m — наименьшее число гиперплоскостей, дающее в пересечении рассматриваемую плоскость). Например, это всегда верно, если q — простое число. Даже если q не является простым числом, это может быть верно для определенных (хотя и не для всех) конфигураций. В этих случаях число $\mu(L) = q^k$ для k -мерной плоскости k приобретает ясный комбинаторный смысл. Значение характеристического многочлена $\chi(q)$ считает тогда точки во множестве $\mathbb{Z}_q^d \setminus \bigcup \mathcal{A}$.

Предложение. Для всякой конфигурации $\mathcal{A}$ гиперплоскостей в $\mathbb{R}^d$, имеет место формула $\chi(q) = q^d - nq^{d-1} + \dots$, в которой n — число гиперплоскостей, а точки обозначают члены с q^k при $k < d - 1$.

Доказательство. Во-первых, значение меры μ на множестве $\mathbb{R}^d \setminus \bigcup \mathcal{A}$ равно $q^d - \mu(H_1 \cup \dots \cup H_n)$, где $H_1, \dots, H_n$ — все гиперплоскости в $\mathcal{A}$. Достаточно, таким образом, установить, что $\mu(H_1 \cup \dots \cup H_n)$ записывается как nq^{d-1} плюс члены с меньшими степенями переменной q . И в самом деле, nq^{d-1} — это первая сумма в правой части формулы включения-исключения, а остальные суммы содержат только члены вида q^k при $k < d - 1$, поскольку пересечение как минимум двух различных гиперплоскостей имеет коразмерность > 1 . $\square$

Мы теперь для примера можем посчитать характеристический многочлен для $d = 2$.

Пример. Пусть $\mathcal{A}$ — конфигурация из n прямых $L_1, \dots, L_n$ на плоскости. Нам нужно посчитать значение меры μ на дополнении до $L_1 \cup \dots \cup L_n$; это, конечно, сводится к вычислению меры самого множества $L_1 \cup \dots \cup L_n$. В качестве первого приближения к $\mu(L_1 \cup \dots \cup L_n)$, как и в формуле включения-исключения, мы возьмем сумму мер всех прямых конфигурации, то есть nq . При этом мера каждой вершины a оказалась посчитанной ровно ν_a раз, поэтому нужно вычесть сумму величин $\nu_a - 1$ по всем вершинам:

$$\mu(L_1 \cup \dots \cup L_n) = nq - \sum_a (\nu_a - 1).$$

Отсюда получается такая формула для характеристического многочлена:

$$\chi(q) = q^2 - nq + \sum_a (\nu_a - 1).$$

Стоит обратить внимание на то, что сформулированные выше теоремы Заславского тесно связаны со значениями характеристического многочлена. А именно, количество всех областей конфигурации $\mathcal{A}$ равно $(-1)^d \chi(-1)$, а количество ограниченных областей равно $(-1)^d \chi(1)$. Именно в таком виде теорема была сформулирована для произвольной размерности.

2.3. Хроматический многочлен графа. *Графом* G называется структура, состоящая из множества $V(G)$, элементы которого называются *вершинами графа*, множества $E(G)$, элементы которого называются *ребрами графа*, и отображения ϕ_G , определенного на множестве $E(G)$ и сопоставляющего каждому ребру неупорядоченную пару вершин, называемых *концами ребра*. Рассматриваемые нами графы являются *неориентированными*, поскольку мы не указываем ни направления ребер, ни порядка вершин, являющихся концами данного ребра. В графах допускаются *петли*, то есть ребра, ведущие из вершины в эту же вершину, а также *кратные ребра*, то есть несколько ребер, соединяющие одну и ту же пару вершин.

Определение (Раскраски). Если каждой вершине графа однозначно сопоставлен некоторый элемент из q -элементного множества, например, из множества $\mathbb{Z}_q$, то говорят о *раскраске вершин графа в q цветов*. Более формально, раскраску можно определить как отображение ξ из $V(G)$ во множество $\mathbb{Z}_q$. *Правильной раскраской* графа G в q цветов называется такая раскраска ξ , при которой соседние вершины обязательно покрашены в разные цвета, то есть $\xi(a) \neq \xi(b)$, если $\{a, b\} = \phi_G(e)$ для некоторого $e \in E(G)$.

Обозначим через $\chi_G(q)$ число способов раскрасить граф G в q цветов. Ниже мы увидим, что это многочлен от q степени d , где d — число вершин графа, за исключением того случая, когда у G есть петли — в последнем случае функция χ_G тождественно равна нулю. Заметим, что если у G имеются кратные ребра, например, несколько ребер соединяет одну и ту же пару вершин a и b , то можно все эти ребра кроме одного удалить, что не повлияет на многочлен χ_G и даже не повлияет на множество правильных раскрасок. Функция χ_G называется *хроматическим многочленом* графа G .

Пример (Полный граф на d вершинах). *Полный граф* K_d , иногда называемый также *кликой* порядка d — это граф без петель, в котором любые две различные вершины соединены единственным ребром. Чтобы раскраска такого графа была правильной, необходимо, чтобы все вершины были покрашены в разные цвета. В частности, число цветов должно быть не меньше чем d . Количество правильных раскрасок, как нетрудно видеть, выражается формулой

$$\chi_{K_d}(q) = q(q-1) \dots (q-d+1).$$

Цвет первой вершины можно выбрать произвольно, а при выборе цвета каждой из последующих вершин, нужно избегать совпадений с уже выбранными цветами — отсюда и получается приведенная формула.

Пример (Линейный граф с d вершинами). *Линейный граф* P_d определяется как граф, вершины которого пронумерованы индексами от 1 до d , и ребра которого соединяют вершины с соседними индексами, то есть i и $i+1$. Хроматический многочлен линейного графа выражается формулой $\chi_{P_d} = q(q-1)^{d-1}$. В самом деле, цвет первой вершины можно выбрать произвольно, а далее, при определении цвета вершины номер $i+1$, нужно избегать только цвета i -й вершины; все остальные цвета подходят.

Следующий пример чуть менее тривиален в смысле нахождения характеристического многочлена.

Пример (Циклический граф с d вершинами). *Циклический граф* C_d , по определению, представляет собой d циклически упорядоченных вершин и d ребер, соединяющих соседние в смысле циклического порядка вершины. Для вычисления характеристического многочлена графа C_d , выделим одно ребро $e \in E(C_d)$ и (временно) удалим его. Останется граф P_d , характеристический многочлен которого мы уже посчитали. С другой стороны, не всякая правильная раскраска графа P_d является правильной для C_d . Из числа $\chi_{P_d}(q)$ мы должны вычесть количество таких правильных раскрасок графа P_d , для которых оба конца покрашены в один и тот же цвет — дело в том, что эти и только эти правильные раскраски графа P_d не являются правильными для C_d . Последнее количество раскрасок можно проинтерпретировать как $\chi_{C_{d-1}}(q)$, склеив оба конца графа P_d . Мы получаем, таким образом, следующее рекуррентное соотношение:

$$\chi_{C_d}(q) = q(q-1)^{d-1} - \chi_{C_{d-1}}(q),$$

откуда уже нетрудно вывести и явную формулу

$$\chi_{C_d}(q) = q((q-1)^{d-1} - (q-1)^{d-2} \pm \dots) = (q-1)^d + (-1)^d(q-1).$$

Определим теперь конфигурацию гиперплоскостей, связанную с графом G .

Определение (Графовые конфигурации). Рассмотрим граф G и свяжем с ним некоторую конфигурацию гиперплоскостей $\mathcal{A}_G$. Для этого перенумеруем все вершины графа как $v_1, \dots, v_d$, так что вершины соответствуют координатам $x_1, \dots, x_d$ пространства $\mathbb{R}^d$. Теперь каждому ребру $e \in E(G)$, соединяющему вершины v_i и v_j , сопоставим гиперплоскость H_e в $\mathbb{R}^d$, заданную уравнением $x_i = x_j$. Набор таких гиперплоскостей обозначим через $\mathcal{A}_G$ и назовем *графовой конфигурацией*, соответствующей графу G .

Мы теперь можем установить связь между хроматическим многочленом графа и характеристическим многочленом соответствующей графовой конфигурации.

Теорема. Допустим, что в графе G нет ни петель, ни кратных ребер. Хроматический многочлен графа G совпадает с характеристическим многочленом конфигурации $\mathcal{A}_G$.

Как следствие, получаем, что хроматический многочлен любого графа в самом деле является многочленом. Более того, этот многочлен имеет вид $\chi_G(q) = q^d - nq^{d-1} + \dots$, где n — число ребер в графе G . Что делать с петлями и кратными ребрами, мы уже обсуждали выше.

Доказательство. Главное, на что нужно обратить внимание — это на то обстоятельство, что конфигурация $\mathcal{A}_G$ может быть определена не только над любым полем, но и вообще над любым множеством, поскольку в ее определении не фигурируют никакие алгебраические операции, кроме отношения равенства. Обозначим через $\mathcal{A}_G(q)$ конфигурацию в $\mathbb{Z}_q^d$, заданную точно так же, как и $\mathcal{A}_G$, с той лишь разницей, что все координаты принимают значения в $\mathbb{Z}_q$, вместо того чтобы быть действительными числами. Для каждого конкретного q элементы множества $\mathbb{Z}_q^d \setminus \bigcup \mathcal{A}_G(q)$ соответствуют правильным раскраскам графа G , а значит количество этих элементов совпадает с $\chi_G(q)$. Если обозначить характеристический многочлен конфигурации $\mathcal{A}_G$ через $\chi(q)$, то мы получили $\chi(q) = \chi_G(q)$ для каждого конкретного целого положительного q . Но отсюда вытекает равенство многочленов. $\square$

2.4. Рекурсия типа «удаление – схлопывание». Характеристический многочлен $\chi_{\mathcal{A}}$ конфигурации $\mathcal{A}$ может быть описан рекурсивно. Выберем и зафиксируем одну из гиперплоскостей $H \in \mathcal{A}$. Тогда можно рассмотреть меньшую конфигурацию $\mathcal{A} \setminus \{H\}$, полученную из $\mathcal{A}$ удалением гиперплоскости H . А кроме того, можно рассмотреть конфигурацию $\mathcal{A}|_H$ в гиперплоскости H — эта конфигурация имеет меньшую размерность. Обе конфигурации $\mathcal{A} \setminus \{H\}$ и $\mathcal{A}|_H$ в определенном смысле проще, чем $\mathcal{A}$. Следовательно, имеет смысл сводить вычисление тех или иных величин, связанных с конфигурацией $\mathcal{A}$, к вычислению тех же величин для $\mathcal{A} \setminus \{H\}$ и $\mathcal{A}|_H$. В качестве главного и характерного примера, мы рассмотрим рекурсивную формулу для вычисления характеристического многочлена.

Чтобы упростить обозначения, положим $\mathcal{A}' := \mathcal{A} \setminus \{H\}$ и $\mathcal{A}'' := \mathcal{A}|_H$, а соответствующие этим конфигурациям характеристические многочлены будем обозначать через χ' и χ'' , соответственно. Тогда, согласно свойству аддитивности,

$$\mu(\bigcup \mathcal{A}) = \mu(\bigcup \mathcal{A}' \cup H) = \mu(\bigcup \mathcal{A}') + \mu(H \setminus \bigcup \mathcal{A}').$$

В терминах характеристических многочленов это равенство переписывается как $\chi = \chi' - \chi''$. Это последнее равенство можно использовать для доказательства различных свойств характеристического многочлена по индукции.

Пример (Удаление и схлопывание в теории графов). Для случая графовой конфигурации $\mathcal{A} = \mathcal{A}_G$ производные конфигурации $\mathcal{A}'$ и $\mathcal{A}''$ тоже имеют ясный смысл в терминах теории графов. Рассмотрим ребро $e \in E(G)$, и построим два новых графа по G и e . Во-первых, можно образовать граф G' , полученный из G удалением ребра e . А во-вторых, новый граф G'' можно получить, схлопнув ребро e в точку, то есть удалив это ребро и отождествив между собой его концы. Тождество $\chi(q) = \chi'(q) - \chi''(q)$ для хроматических многочленов в нашем случае называется *рекурсией удаления – схлопывания*. Мы эту рекурсию фактически уже использовали, когда считали хроматический многочлен для C_d .

В общем случае, то есть применительно к произвольным конфигурациям гиперплоскостей в $\mathbb{R}^d$, мы теперь можем обосновать общие формулы Заславского.

Теорема (Заславский). Пусть $\mathcal{A}$ — непустая конфигурация гиперплоскостей в $\mathbb{R}^d$, а χ — характеристический многочлен этой конфигурации. Тогда число областей конфигурации $\mathcal{A}$ равно $(-1)^d \chi(-1)$, а число ограниченных областей равно $(-1)^d \chi(1)$.

Доказательство. Сначала рассмотрим конфигурацию, состоящую из одной единственной гиперплоскости. Характеристический многочлен такой конфигурации равен $q^d - q^{d-1}$. Его

значение в точке $q = -1$ равно $(-1)^d 2$, а значение в точке $q = 1$ равно нулю. Это соответствует тому, что рассматриваемая конфигурация имеет 2 области и не имеет ограниченных областей. Будем теперь вести индукцию по числу гиперплоскостей; основание индукции мы уже проверили.

Можно также считать, что $d > 2$, поскольку случай $d = 2$ мы уже обсудили, а случай $d = 1$ очевиден. Пусть $\mathcal{A}$ — конфигурация гиперплоскостей в $\mathbb{R}^d$. Выберем и зафиксируем одну гиперплоскость $H \in \mathcal{A}$; получим две другие конфигурации $\mathcal{A}' := \mathcal{A} \setminus \{H\}$ и $\mathcal{A}'' := \mathcal{A}|_H$, для которых выполнено предположение индукции. Обозначим через χ' , χ'' соответствующие характеристические многочлены. Согласно сказанному выше, $\chi = \chi' - \chi''$. Это означает, что $(-1)^d \chi(-1) = (-1)^d \chi'(-1) + (-1)^{d-1} \chi''(-1)$; таким образом, нужно проверить, что $f_0 = f'_0 + f''_0$ (количество штрихов у f_0 очевидным образом ссылается на выбор конфигурации). В самом деле, плоскость H разбивает ровно f''_0 областей конфигурации $\mathcal{A}'$ на две области каждую.

Нам осталось проверить рекуррентную формулу для числа ограниченных областей, которая сводится к $f_0^b = f'_0{}^b + f''_0{}^b$. В самом деле, каждая ограниченная область конфигурации $\mathcal{A}''$ либо делит ограниченную область конфигурации $\mathcal{A}'$ на две ограниченные части, либо делит неограниченную область конфигурации $\mathcal{A}'$ на две части, одна из которых ограничена, а другая нет. Таким образом, в любом случае, каждая ограниченная область конфигурации $\mathcal{A}''$ отвечает за прибавление единицы к числу ограниченных областей конфигурации $\mathcal{A}'$. Неограниченные области конфигурации $\mathcal{A}'$ ничего в этом смысле не добавляют. $\square$

Многочлен $\pi(q) := (-1)^d \chi(-q)$ называется *многочленом Пуанкаре*. Разумеется, он несет не больше и не меньше информации, чем характеристический многочлен, но иногда рассматривать многочлен Пуанкаре оказывается удобнее. Например, теорема Заславского выражает число областей и число ограниченных областей как определенные значения многочлена Пуанкаре. Запишем $\pi(q)$ в виде

$$\pi(q) = h_0 + h_1 q + \cdots + h_{d-1} q^{d-1} + q^d.$$

Коэффициенты h_k многочлена Пуанкаре образуют так называемый *h -вектор конфигурации*. Согласно установленному выше, h_{d-1} совпадает с числом гиперплоскостей.

Рекурсия $\chi = \chi' - \chi''$ соответствует рекурсии $\pi = \pi' + \pi''$ в терминах многочленов Пуанкаре π , π' , π'' конфигураций $\mathcal{A}$, $\mathcal{A}'$, $\mathcal{A}''$, соответственно.

Предложение. *Компоненты h -вектора неотрицательны. Следовательно, знаки коэффициентов характеристического многочлена чередуются.*

Доказательство. Это утверждение верно для случая пустой конфигурации ($\pi(q) = q^d$) в произвольной размерности d , а также для размерности $d = 0$. Теперь достаточно воспользоваться формулой $\pi = \pi' + \pi''$, предположив по индукции, что утверждение верно для конфигураций $\mathcal{A}'$ и $\mathcal{A}''$. Многочлен в левой части имеет неотрицательные коэффициенты, как сумма двух многочленов с неотрицательными коэффициентами. $\square$

Коэффициенты многочлена $\pi(q)$ образуют *унимодальную* последовательность, то есть сначала нестрого возрастают, доходят до некоторого максимального значения, после чего нестрого убывают. Это теорема, доказанная Хо Джун И [5], закрыла давно стоявшую открытую гипотезу, и это еще один результат в серии блестящих достижений Хо Джун И, за которые ему была вручена филдсовская медаль.

ss:VG

2.5. Алгебра Варченко–Гельфанда (1987). Если S — конечное множество, то обозначим через $\mathbb{Z}^S$ множество всех функций на X со значениями в целых числах. Заметим, что $\mathbb{Z}^S$ — не просто множество; оно несет дополнительную алгебраическую структуру, а именно, функции можно складывать, вычитать и перемножать, причем операции сложения, вычитания и умножения функций удовлетворяют привычным алгебраическим свойствам. Такую структуру называют структурой *алгебры над $\mathbb{Z}$* . Делить функции, вообще говоря, нельзя, поскольку функции могут принимать нулевые значения в некоторых точках.

Определение (Алгебра Варченко–Гельфанда [2]). Пусть $\mathcal{A}$ — конфигурация гиперплоскостей в $\mathbb{R}^d$. Рассмотрим алгебру $\mathcal{V}$ всех функций со значениями в $\mathbb{Z}$ на множестве всех областей конфигурации $\mathcal{A}$. Эта алгебра называется *алгеброй Варченко–Гельфанда*.

Множество всех областей конфигурации $\mathcal{A}$ — это конечное множество; оно состоит из f_0 элементов. Функции на этом конечном множестве можно также понимать как локально постоянные функции на $\mathbb{R}^d \setminus \bigcup \mathcal{A}$, то есть как функции на $\mathbb{R}^d \setminus \bigcup \mathcal{A}$, постоянные на каждой области конфигурации $\mathcal{A}$. Кроме того, можно эти функции как угодно доопределить на множестве $\bigcup \mathcal{A}$ и считать, что речь идет о функциях, определенных на всем $\mathbb{R}^d$, причем функции, отличающиеся только на $\bigcup \mathcal{A}$, следует считать одинаковыми.

Определение (Функции Хэвисайда). Гиперплоскость делит пространство $\mathbb{R}^d$ на два полупространства. Задать *коориентацию гиперплоскости* H — все равно что выбрать одно из двух полупространств в $\mathbb{R}^d \setminus H$ и назвать его *положительной стороной* относительно H , а второе полупространство назвать *отрицательной стороной*. Для каждой гиперплоскости $H \in \mathcal{A}$ с фиксированной коориентацией определена функция θ_H , равная единице по положительную сторону от H и нулю по отрицательную сторону. Эта функция называется *функцией Хэвисайда* относительно H . Чему равна функция Хэвисайда на самой гиперплоскости H совершенно не важно, поскольку мы будем рассматривать θ_H как элемент алгебры $\mathcal{V}$. Последнее законно, так как функция Хэвисайда, конечно, постоянна на каждой области конфигурации $\mathcal{A}$.

Функции Хэвисайда составляют *набор образующих* алгебры $\mathcal{V}$ в следующем смысле.

Предложение. Любую функцию $f \in \mathcal{V}$ можно записать как $P(\theta_{H_1}, \dots, \theta_{H_n})$, где P — многочлен от n переменных, а $H_1, \dots, H_n$ — все гиперплоскости конфигурации $\mathcal{A}$, причем для каждой из этих гиперплоскостей выбрана и зафиксирована коориентация.

Доказательство. Обозначим все области конфигурации $\mathcal{A}$ через $e_1, \dots, e_{f_d}$. Для каждой области e_i , введем *индикаторную функцию* $\mathbb{1}_i \in \mathcal{V}$, равную единице на e_i и нулю на всех остальных областях конфигурации $\mathcal{A}$. Всякая функция $f \in \mathcal{V}$ может быть записана как

$$f = f(e_1)\mathbb{1}_1 + \dots + f(e_{f_d})\mathbb{1}_{f_d},$$

следовательно, искомое утверждение достаточно проверить для функции $\mathbb{1}_i$.

Если $H^\pm$ — одна и та же гиперплоскость, но с разными коориентациями, то $\theta_{H^+} + \theta_{H^-} = 1$. Следовательно, функции Хэвисайда гиперплоскостей из $\mathcal{A}$, взятых с произвольными коориентациями, выражаются искомым образом через функции Хэвисайда $\theta_1, \dots, \theta_n$ коориентированных гиперплоскостей $H_1, \dots, H_n$, соответственно. Зафиксируем e_i ; достаточно считать, что e_i лежит по положительную сторону от каждой из гиперплоскостей $H_1, \dots, H_n$. Но в таком случае $\mathbb{1}_i = \theta_1 \dots \theta_n$. $\square$

Для всякого неотрицательного целого k , определим $\mathcal{V}_{\leq k}$ как подмножество в $\mathcal{V}$, состоящее из всех функций, представимых полиномами степени $\leq k$ от функций Хэвисайда. Это подмножество — *аддитивная подгруппа* в $\mathcal{V}$ в том смысле, что сложение и вычитание не выводит за пределы этой подгруппы; однако умножение, вообще говоря, выводит. Если какую-то функцию из $\mathcal{V}$ можно представить многочленом степени k , то ее же можно представить и многочленом степени $k+1$, что следует из уже упомянутого равенства $\theta_{H^+} + \theta_{H^-} = 1$. Поэтому подгруппы $\mathcal{V}_{\leq k}$ можно также охарактеризовать как состоящие из всех функций, представимых многочленами степени k . Группы $\mathcal{V}_{\leq k}$ вложены друг в друга:

$$\mathcal{V}_{\leq 0} \subset \mathcal{V}_{\leq 1} \subset \dots \subset \mathcal{V}_{\leq d} \subset \dots,$$

то есть, как говорят, образуют *фильтрацию*. Будем называть ее *степенной фильтрацией*.

Определение (Факторы). Пусть A — *аддитивная абелева группа*, то есть такое множество, элементы которого можно складывать и вычитать, причем эти операции удовлетворяют привычным свойствам, включая коммутативность (переместительное свойство)

сложения. Рассмотрим *подгруппу* $B \subset A$, то есть подмножество, замкнутое относительно сложения и вычитания. *Факторгруппа* A/B получается из A отождествлением всех элементов из B с нулем. Более формально, элементами факторгруппы A/B служат классы отношения эквивалентности

$$a \sim a' \Leftrightarrow a - a' \in B.$$

Классы элементов $a + a'$ и $a - a'$ зависят только от класса элемента a и класса элемента a' , и по этой причине на множестве A/B корректно определены операции сложения и вычитания.

Возвращаясь к алгебре Варченко–Гельфанда, мы можем определить *факторы степенной фильтрации* как

$$\mathcal{V}_k := \mathcal{V}_{\leq k} / \mathcal{V}_{< k}, \quad \mathcal{V}_{< k} := \mathcal{V}_{\leq k-1}.$$

Мы здесь полагаем $\mathcal{V}_0 := \mathcal{V}_{\leq 0}$, то есть что $\mathcal{V}_{< 0} = \{0\}$. Связь между компонентами h -вектора и алгеброй Варченко–Гельфанда выражается следующей теоремой, доказательство которой мы отложим до следующего раздела.

Теорема. *Рассмотрим конфигурацию гиперплоскостей $\mathcal{A}$ в $\mathbb{R}^d$ и соответствующую алгебру Варченко–Гельфанда. Фактор $\mathcal{V}_k$ степенной фильтрации изоморфен группе функций на множестве из h_{d-k} элементов.*

Две группы *изоморфны*, если между ними можно установить такое взаимно однозначное соответствие, при котором сумма переходит в сумму. Само это соответствие называется (групповым) *изоморфизмом*. Более общее понятие — *гомоморфизм групп* (или *групповой гомоморфизм*); это отображение $f : A \rightarrow B$ между группами, которое сложение переводит в сложение, то есть $f(a + a') = f(a) + f(a')$ для любых $a, a' \in A$. Изоморфизм, таким образом, это гомоморфизм, являющийся одновременно биекцией.

Эта теорема, во-первых, дает более концептуальное объяснение неотрицательности компонент h -вектора. А во-вторых, первая часть теоремы Заславского тоже может быть объяснена с помощью степенной фильтрации алгебры $\mathcal{V}$ и сформулированной теоремы. Если группа A изоморфна группе функций на множестве из r элементов, то говорят, что A — *свободная абелева группа ранга r* (пишут $\text{rk}(A) = r$). Допустим теперь, что A допускает фильтрацию $A_{\leq 0} \subset \dots \subset A_{\leq d} = A$, у которой факторы $A_k := A_{\leq k} / A_{\leq k-1}$ тоже являются свободными абелевыми группами рангов r_k . В этом случае нетрудно показать, что $r = r_0 + \dots + r_d$. Применяя это общее наблюдение к степенной фильтрации алгебры $\mathcal{V}$ и пользуясь теоремой, получаем, что $\pi(1) = h_0 + \dots + h_d$ равно рангу группы $\mathcal{V}$, то есть f_2 (вспомним, что $\mathcal{V}$ — это алгебра функций на множестве из f_2 элементов).

Замечание. Группа всех функций на множестве из r элементов со значениями в $\mathbb{Z}$ изоморфна группе $\mathbb{Z}^r$, состоящей из наборов целочисленных координат длины r , с покоординатным сложением. В частности, $\mathcal{V}$, как группа по сложению, изоморфна группе $\mathbb{Z}^{f_2}$. А сформулированная выше теорема утверждает, что $\mathcal{V}_k$ изоморфна группе $\mathbb{Z}^{h_{d-k}}$.

3. СТРУКТУРА АЛГЕБРЫ ВАРЧЕНКО–ГЕЛЬФАНДА

Мы теперь более внимательно рассмотрим структуру алгебры Варченко–Гельфанда в том случае, когда $d = 2$, то есть $\mathcal{A}$ представляет собой конфигурацию прямых на плоскости. Наш подход к описанию алгебры $\mathcal{V}$ — *топологический*. Вообще, *топология* изучает наиболее грубые свойства геометрических фигур — такие, которые не меняются при искажении размеров и пропорций, а меняются только при разрывах и склейках. Если непрерывно деформировать фигуру, не допуская ни разрывов, ни склеивания, то будут получаться топологически эквивалентные или, как еще говорят, *гомеоморфные* фигуры. Очень часто топологи не делают разницы между гомеоморфными фигурами, например, между окружностью и кардиоидой (замкнутой кривой в виде сердечка).

3.1. Обзор результатов. Факторы алгебры Варченко–Гельфанда могут быть описаны относительно явно. Чтобы сориентировать читателя, приведем это описание сразу, а доказательства и необходимые для них топологические конструкции изложим позже. Мы всюду считаем, что $d = 2$, то есть рассматриваем конфигурацию $\mathcal{A}$ прямых на плоскости.

Напомним, что алгебра Варченко–Гельфанда снабжена степенной фильтрацией, то есть фильтрацией по степени многочленов, в которые подставляются функции Хэвисайда. В нашем двумерном случае степенная фильтрация выглядит так:

$$\mathcal{V}_0 = \mathcal{V}_{\leq 0} \subset \mathcal{V}_{\leq 1} \subset \mathcal{V}_{\leq 2} = \mathcal{V}.$$

Равенства $\mathcal{V}_{\leq 2} = \mathcal{V}$ мы, впрочем, еще не установили. Это содержательное и интересное утверждение: оказывается, что всякая функция на 2-клетках получается как многочлен от функций Хэвисайда степени не выше 2. Очевидно, группа $\mathcal{V}_0$ совпадает с множеством всех постоянных функций и, следовательно, изоморфна группе $\mathbb{Z}$. Следующий член степенной фильтрации $\mathcal{V}_{\leq 1}$ состоит из таких функций, которые можно представить как постоянный член плюс линейная комбинация функций Хэвисайда с целыми коэффициентами. В факторе $\mathcal{V}_1 := \mathcal{V}_{\leq 1}/\mathcal{V}_{\leq 0}$ постоянным членом можно пренебречь. Для описания группы $\mathcal{V}_1$ нам понадобится понятие коориентированной прямой.

Определение (Коориентация). Прямая L делит плоскость на две полуплоскости. Если выбрана одна из этих двух полуплоскостей и названа *положительной*, то говорят, что задана *коориентация* прямой L . Когда какая-либо геометрическая фигура лежит в положительной полуплоскости относительно прямой L , говорят, что она лежит *по положительную сторону* от L . Определим *коориентированную прямую* как прямую с фиксированной коориентацией. Всякая прямая, таким образом, определяет две различные коориентированные прямые, отличающиеся только выбором коориентации.

Функция g на множестве коориентированных прямых конфигурации $\mathcal{A}$ со значениями в $\mathbb{Z}$ называется *знакопеременной*, если она меняет знак каждый раз, когда коориентация прямой меняется на противоположную. Обозначим группу всех знакопеременных функций на коориентированных прямых из $\mathcal{A}$ со значениями в $\mathbb{Z}$ через $\mathcal{F}_1$. Групповая операция, разумеется, заключается в поточечном сложении функций. Если $\mathcal{A}$ состоит из n прямых, то, чтобы задать функцию из $\mathcal{F}_1$, достаточно выбрать на всех прямых $\mathcal{A}$ коориентации произвольным образом и задать значения функции на полученных коориентированных прямых. Таким образом, зафиксировав коориентации у всех прямых из $\mathcal{A}$, мы получаем изоморфизм между группой $\mathcal{F}_1$ и группой $\mathbb{Z}^n$. Важно, однако, помнить, что этот изоморфизм существенно зависит от выбранных коориентаций.

Теорема. Первый фактор $\mathcal{V}_1 := \mathcal{V}_{\leq 1}/\mathcal{V}_{\leq 0}$ алгебры Варченко–Гельфанда канонически изоморфен пространству $\mathcal{F}_1$. Следовательно, $\mathrm{rk}(\mathcal{V}_1) = n$.

Изоморфизм устроен следующим образом. Рассмотрим элемент из группы $\mathcal{V}_1$, представленный функцией $f \in \mathcal{V}_{\leq 1}$. По определению степенной фильтрации, f представляет собой линейную комбинацию функций Хэвисайда, плюс константа. Для каждой коориентированной прямой $L \in \mathcal{A}$ определен скачок, который значение функции f совершает при переходе из отрицательной стороны в положительную. Этот скачок корректно определен во всех точках прямой L , кроме вершин конфигурации, и постоянен вдоль L . Ясно, что при смене коориентации прямой L скачок меняет знак. Вот мы и получаем знакопеременную функцию g на коориентированных прямых из $\mathcal{A}$. Про переход от f к g можно думать как про дискретное дифференцирование.

Для описания второго фактора $\mathcal{V}_2$, нам понадобится рассматривать ориентированные флаги.

Определение (Ориентированные флаги). Рассмотрим прямую $L \in \mathcal{A}$ и некоторую вершину $a \in L$ конфигурации $\mathcal{A}$; пара (a, L) называется *флагом конфигурации $\mathcal{A}$* . *Ориентированный флаг* $[a, L]$ нашей конфигурации — это флаг (a, L) , снабженный ориентацией

прямой L (=направлением движения вдоль этой прямой), а также коориентацией прямой L в $\mathbb{R}^2$. Ориентированный флаг задает также ориентацию плоскости. А именно, если мы движемся в положительном направлении вдоль прямой L , и вдруг решили «сойти с дистанции», повернув в положительную сторону относительно L , то мы тем самым совершим поворот в положительном направлении. Если же мы знаем, что такое положительное направление вращения, то мы знаем ориентацию плоскости.

Функция g на ориентированных флагах конфигурации $\mathcal{A}$ называется *знакопеременной*, если, каждый раз, как у L меняется либо ориентация, либо коориентация (но что-то одно!), значение $g[a, L]$ меняет знак. Эквивалентным образом, можно еще так сказать: для знакопеременной функции g , значение $g[a, L]$ зависит только от геометрического флага (a, L) и от ориентации плоскости, которую определяет флаг $[a, L]$; обращение ориентации плоскости влечет обращение знака у $g[a, L]$. Обозначим через $\mathcal{F}_2$ аддитивную группу всех знакопеременных функций на ориентированных флагах конфигурации $\mathcal{A}$ со значениями в целых числах.

Теорема. *Группа $\mathcal{V}_2$ канонически изоморфна подгруппе в $\mathcal{F}_2$, образованной всеми функциями $g \in \mathcal{F}_2$, такими, что, для любой вершины a , сумма значений $g[a, L]$ по всем прямым $L \ni a$ равна нулю. В этой сумме, для каждой из прямых произвольным образом задается коориентация, а ориентация определяется так, чтобы из всех фигурирующих в сумме ориентированных флагов $[a, L]$ получалась одна и та же ориентация плоскости.*

Указанное в теореме соотношение на функцию $g \in \mathcal{F}_2$, связанное с вершиной a , называется *законом взаимности* в вершине a . Теорема говорит о том, что факторгруппу $\mathcal{V}_2$ можно отождествить с подгруппой в $\mathcal{F}_2$. А именно, можно считать, что $\mathcal{V}_2 \subset \mathcal{F}_2$ состоит из тех знакопеременных функций на ориентированных флагах, которые удовлетворяют законам взаимности по всем вершинам конфигурации. Чтобы вложить факторгруппу $\mathcal{V}_2 = \mathcal{V}/\mathcal{V}_{\leq 1}$ в $\mathcal{F}_2$, нужно определить такой гомоморфизм из $\mathcal{V}$ в $\mathcal{F}_2$, который отправлял бы в нуль все функции из $\mathcal{V}_{\leq 1}$. Этот искомый гомоморфизм состоит, неформально говоря, во взятии второй дискретной производной. Начнем с функции $f \in \mathcal{V}$ и сначала построим функцию f' , определенную на каждой коориентированной прямой, кроме вершин — это та же самая функция, первая дискретная производная функции f , которую мы рассматривали выше. Поскольку мы теперь не предполагаем, что $f \in \mathcal{V}_{\leq 1}$, то, вообще говоря, f' принимает разные значения на разных 1-клетках одной и той же прямой. Обозначим через f'_L ограничение функции f' на коориентированную прямую L ; при обращении коориентации у L функция f'_L меняет знак.

Рассмотрим ориентированный флаг $[a, L]$ и определим значение $g[a, L]$ как разность $f'_L(a+) - f'_L(a-)$, где $a\pm$ — это точки прямой L поблизости от a , выбранные так, что $a+$ находится по положительную сторону от a в смысле выбранной ориентации прямой L , а точка $a-$ находится по отрицательную сторону. Нетрудно видеть, что так определенная функция $g \in \mathcal{F}_2$ обращается в нуль тогда и только тогда, когда $f \in \mathcal{V}_{\leq 1}$. Следовательно, соответствие $f \mapsto g$ определяет некоторое изоморфное вложение группы $\mathcal{V}_2$ в группу $\mathcal{F}_2$. То, что образ этого вложения задается в точности законами взаимности во всех вершинах конфигурации, мы докажем позже. А сейчас отметим полезное следствие из приведенного описания группы $\mathcal{V}_2$. Чтобы удовлетворить всем законам взаимности, достаточно задать значения $g[a, L]$ для каждой точки a и для всех прямых $L \ni a$, кроме одной (то, какую прямую мы опускаем, может зависеть от точки a). Тогда все значения функции g однозначно восстановятся из законов взаимности. Получаем такое следствие.

Следствие. *Второй фактор $\mathcal{V}_2$ — свободная абелева группа ранга $\sum_a (\nu_a - 1)$, где суммирование ведется по всем вершинам конфигурации.*

Отсюда можно — еще раз, но более концептуально — вывести формулу Заславского. С одной стороны, группа $\mathcal{V}$ изоморфна $\mathbb{Z}^{f_2}$, и, стало быть, имеет ранг f_2 . А с другой стороны,

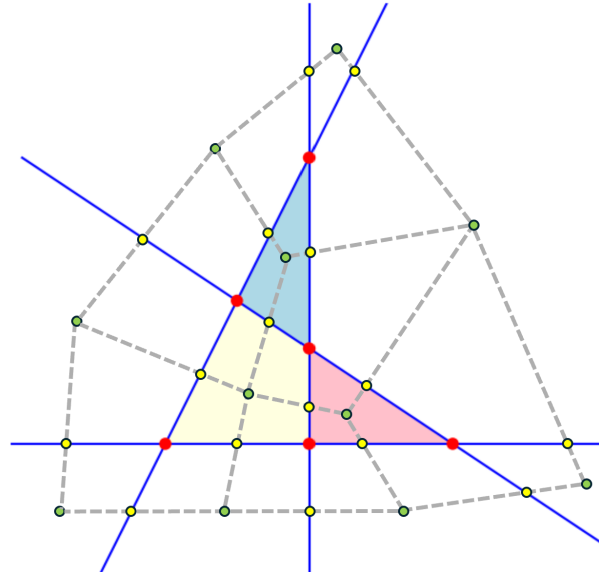


РИС. 5. Двойственные клетки для конфигурации прямых из рис. 2 (одномерные двойственные клетки показаны пунктиром). Особенность этого изображения в том, что двойственные 1-клетки показаны как прямолинейные отрезки; вообще говоря, они не обязаны быть отрезками, а могут быть двузвенными ломаными.

fig:dual

ее ранг можно посчитать как сумму рангов факторов:

$$f_2 = \text{rk}(\mathcal{V}) = \text{rk}(\mathcal{V}_{\leq 0}) + \text{rk}(\mathcal{V}_1) + \text{rk}(\mathcal{V}_2) = 1 + n + \sum_a (\nu_a - 1).$$

3.2. Двойственные клетки. Как уже обсуждалось выше, конфигурация $\mathcal{A}$ определяет разбиение плоскости на клетки размерностей 0, 1 и 2. Нам будет удобнее работать с другим, так называемым *двойственным*, разбиением на клетки. Будем называть клетки двойственного разбиения *двойственными клетками*. Двойственные k -клетки взаимно однозначно сопоставляются $(2 - k)$ -клеткам следующим образом. Внутри каждой 2-клетки отметим по точке, и назовем отмеченную точку *центром* клетки. Если клетка ограниченная, то можно, хотя и не обязательно, в качестве центра взять центр масс. Также отметим по одной точке в каждой 1-клетке, и назовем отмеченные точки *серединами* 1-клеток, хотя это может иметь дословный смысл только для ограниченных 1-клеток, для которых в самом деле можно отметить точку ровно посередине, а для неограниченных клеток это имеет только условный смысл. Отмеченные точки 2-клеток по определению являются *двойственными 0-клетками*. Центры двух соседних 2-клеток можно соединить двузвенной ломаной с вершинами в этих двух центрах и еще посередине ребра, инцидентного обеим 2-клеткам. Такие ломаные будут *двойственными 1-клетками*. Каждая двойственная 1-клетка, таким образом, пересекается с единственной 1-клеткой исходного разбиения. Наконец, объединение всех двойственных 0-клеток и двойственных 1-клеток составляет некоторую сетку на плоскости, ограниченные компоненты дополнения которой называются *двойственными 2-клетками*. В любой двойственной 2-клетке окажется ровно одна 0-клетка исходного разбиения, и наоборот, всякая 0-клетка содержится в единственной двойственной 2-клетке.

Обратите внимание: двойственные клетки покрывают не всю плоскость, а только некоторый ограниченный многоугольник. Этот многоугольник не обязан быть выпуклым, но в нем нет никаких «дырок», он гомеоморфен замкнутому диску; см. рис. 5.

Любая двойственная 0-клетка представляет собой точку. Двойственные 1-клетки гомеоморфны интервалам, не включающим концевых точек — чтобы увидеть этот гомеоморфизм, достаточно каждую 1-клетку «распрямить». А двойственные 2-клетки являются

открытыми многоугольниками, то есть многоугольниками, содержащими внутренность, но не содержащими граничные ребра и вершины. Эти многоугольники не обязательно выпуклы, но топологически они ничем не отличаются от выпуклых многоугольников; иными словами, они все гомеоморфны открытому диску. Кстати говоря, k -клетками вообще принято называть фигуры, гомеоморфные открытому k -мерному диску.

Определение (Группа C_0). Определим $C_0 = C_0(\mathcal{A})$ как группу всех формальных линейных комбинаций двойственных 0-клеток с целыми коэффициентами.

Эта группа, конечно, изоморфна группе $\mathbb{Z}^{f_2}$, как и группа $\mathcal{V}$. Таким образом, имеется групповой изоморфизм между C_0 и $\mathcal{V}$. Нам однако, будет полезней иметь дело не с этим изоморфизмом, а с функцией $\mathcal{V} \times C_0 \rightarrow \mathbb{Z}$, сопоставляющей каждой паре $(f, \alpha) \in \mathcal{V} \times C_0$ сумму $\langle f, \alpha \rangle$ значений функции f по всем точкам из α , посчитанных каждое со своей кратностью, равной коэффициенту, с которым данная точка входит в α . На языке формул, если $\alpha = \sum_b \alpha_b b$, где суммирование ведется по всем двойственным 0-клеткам b , то $\langle f, \alpha \rangle := \sum_b \alpha_b f(b)$. Будем говорить, что f и α *ортгоналичны*, если $\langle f, \alpha \rangle = 0$. Для всякой подгруппы $A \subset C_0$, обозначим через $A^\perp$ подгруппу в $\mathcal{V}$, состоящую всех $f \in \mathcal{V}$ со свойством $\langle f, \alpha \rangle = 0$ для всех $\alpha \in A$. Эта подгруппа называется *ортгоналичным дополнением* до группы A .

Нам еще понадобятся несколько вспомогательных групп. С точки зрения топологии, эти группы появляются очень естественным образом — они представляют собой группы *клеточных цепей*. Хотя общее понятие клеточной цепи мы уточнять не будем, приведем определения конкретных групп. Начнем с группы C_1 , которая определяется аналогично группе C_0 , но с использованием двойственных 1-клеток вместо двойственных 0-клеток. Если все элементы группы можно записать как целочисленные линейные комбинации некоторого фиксированного набора элементов, то члены этого набора называются *образующими* группы. В группе C_0 , например, образующими будут все двойственные 0-клетки, причем между этими образующими нет никаких нетривиальных соотношений. А в группах C_1 , C_2 , определяемых ниже, образующими служат ориентированные двойственные 1-клетки и, соответственно, ориентированные двойственные 2-клетки.

Определение (Группы C_1 и C_2). Определим $C_1 = C_1(\mathcal{A})$ как множество всех линейных комбинаций ориентированных двойственных 1-клеток с целыми коэффициентами; это группа по сложению. Под *ориентацией* двойственной 1-клетки понимается выбор направления на ней, или, что эквивалентно, упорядочение концов (один из которых, согласно выбранному упорядочению, именуется теперь *началом*). Если у двойственной 1-клетки поменять ориентацию, то это все равно что поменять знак. Кроме этого соотношения и соотношения коммутативности, других соотношений в группе C_1 нет. Группа $C_2 = C_2(\mathcal{A})$, по определению, состоит из целочисленных линейных комбинаций ориентированных двойственных 2-клеток. *Ориентация 2-клетки* — это направление обхода ее границы. Аналогично группе C_1 , обращение ориентации двойственной 2-клетки означает смену знака у соответствующего элемента группы C_2 .

Две ориентации одной и той же клетки можно условно назвать ориентацией по часовой стрелке и против часовой стрелки. Условность этих описаний состоит в том, что направление движения часовой стрелки, на самом деле, зависит от положения наблюдателя. В случае реальных часов наблюдатель смотрит на циферблат с одной определенной стороны, поскольку, если посмотреть на часы сзади, то циферблата вообще не будет видно. Однако, если, вместо циферблата, смотреть на идеальную плоскость, то можно на нее смотреть и сзади и спереди, а где какая сторона — неясно. Эти описания (сзади или спереди, по часовой стрелке или против нее) зависят от положения наблюдателя и не относятся к абсолютным свойствам плоскости.

Кроме уже описанных групп C_k при $k = 0, 1, 2$, нам понадобится еще одна группа $C_2^\times$, про которую мы тоже будем думать как про группу особого вида цепей. Образующие группы $C_2^\times$ отождествляются с ориентированными флагами $[a, L]$, причем соответствующая

образующая меняет знак каждый раз, когда меняется ориентация плоскости, заданная флагом $[a, L]$. Чтобы задать некоторый гомоморфизм из $C_2^\times$ в какую-либо абелеву группу G , достаточно задать знакопеременную функцию на ориентированных флагах, принимающую значения в G . Иногда удобно также будет представлять образующую $[a, L]$ как ориентированную двойственную 2-клетку e с центром в вершине a , помеченную прямой L — заметьте, что одна и та же двойственная 2-клетка, помеченная разными прямыми, задает разные образующие группы $C_2^\times$. Ориентация для e соответствует ориентации плоскости, заданной флагом $[a, L]$.

ss:bd-op

3.3. Граничные операторы. Между различными группами клеточных цепей, определенными выше, имеются естественные гомоморфизмы, называемые *граничными операторами*, или *граничными гомоморфизмами*. Начнем с определения граничного гомоморфизма из C_1 в C_0 .

Определение (Граничный оператор). Определим *граничный оператор* $\partial : C_1 \rightarrow C_0$ при помощи следующих двух свойств.

- (1) Во-первых, ∂ — *гомоморфизм групп*, то есть $\partial(\alpha \pm \beta) = \partial\alpha \pm \partial\beta$.
- (2) А во-вторых, $\partial(\ell) = b - a$ для каждого ориентированной двойственной 1-клетки ℓ , ведущей из a в b .

Перечисленные свойства, в самом деле, определяют оператор ∂ однозначно, поскольку каждый элемент группы C_1 можно записать как сумму или разность ориентированных ребер, причем по существу единственным способом. Единственная неоднозначность записи в C_1 состоит в том, что при смене ориентации меняется знак, но формула $\partial(\ell) = b - a$ эту неоднозначность учитывает: при смене ориентации ребра ℓ в левой части поменяется знак, а в правой части точки a и b поменяются местами, так что равенство не нарушится.

Иногда, чтобы подчеркнуть область определения оператора $\partial : C_1 \rightarrow C_0$ и чтобы отличить его от других граничных операторов, мы будем обозначать этот оператор через ∂_1 . Дальнейшие определения граничных операторов будут даны с сокращениями: мы просто будем указывать, как эти операторы действуют на образующих, подразумевая, что далее они продолжаются единственным образом с соблюдением свойства (1), то есть как групповые гомоморфизмы.

Следующий граничный оператор — это оператор $\partial : C_2 \rightarrow C_1$ (который мы иногда будем записывать как ∂_2 , чтобы отличать его от других граничных операторов). Для каждой ориентированной двойственной 2-клетки e определим ∂e как сумму всех двойственных 1-клеток на границе многоугольника e , ориентированных в соответствии с положительным направлением обхода границы. Легко видеть, что $\partial_2 \circ \partial_1 = 0$, потому что каждая двойственная 0-клетка на границе клетки e инцидентна ровно двум двойственным 1-клеткам, причем одна из этих 1-клеток входит, а другая выходит из рассматриваемой точки. Если $\partial\alpha = 0$, то цепь α называется *циклом*, а если $\alpha = \partial\beta$, то тогда α называется *границей*.

Задача. Докажите следующие утверждения.

- Если 1-цепь $\alpha \in C_1$ удовлетворяет условию $\partial\alpha = 0$, то найдется такая 2-цепь $\beta \in C_2$, для которой $\alpha = \partial\beta$.
- Если $\partial\beta = 0$ для $\beta \in C_2$, то $\beta = 0$.
- Цепь $\alpha \in C_0$ можно записать в виде $\alpha = \partial\beta$ для $\beta \in C_1$ тогда и только тогда, когда сумма коэффициентов цепи α при всех двойственных 0-клетках равна нулю.

Иными словами, каждый 1-цикл — граница, нетривиальных 2-циклов не бывает, а 0-цепь является границей тогда и только тогда, когда сумма ее коэффициентов равна нулю.

Приведенные в предыдущей задаче утверждения — топологические свойства плоскости, не зависящие от выбора конфигурации $\mathcal{A}$. Однако при замене плоскости более сложными пространствами (= фигурами произвольных размерностей) утверждения перестают быть

верными. То, насколько данные утверждения нарушены для данной фигуры, является важной топологической характеристикой; эти характеристики изучает *теория гомологий*.

Группы цепей, вместе со связывающими эти группы граничными операторами, активно изучают в топологии и гомологической алгебре под названием *цепные комплексы*. Структура цепных комплексов, которую мы описали, в принципе, имеет смысл в намного более общей ситуации, чем конфигурация прямых на плоскости. Во-первых, можно разбить плоскость на выпуклые многоугольные клетки как угодно. Не обязательно, чтобы из ребер этих клеток составлялись прямые какой-либо конфигурации. А во-вторых, разбиение на клетки возможно далеко не только для плоскости, но и для разных других поверхностей (случай тора, то есть поверхности бублика, мы рассмотрим ниже). Итак, определенные до сих пор граничные операторы носили весьма общий характер.

Еще один граничный оператор, который нам нужно определить — это $\partial^\times : C_2^\times \rightarrow C_1$; в отличие от предыдущих граничных операторов, $\partial^\times$ учитывает структуру конфигурации $\mathcal{A}$. На образующей $[a, L]$ этот оператор действует так. Рассмотрим двойственную 2-клетку e с центром в точке a , ориентированную в соответствии с флагом $[a, L]$, и определим цепь $\partial^\times[a, L] \in C_1$ как сумму тех двойственных 1-клеток на границе клетки e , которые пересекаются с L , то есть центры которых лежат в L .

Введенные выше группы цепей и связывающие их граничные гомоморфизмы можно изобразить следующей диаграммой:

$$\begin{array}{ccccc} C_2 & \xrightarrow{\partial} & C_1 & \xrightarrow{\partial} & C_0 \\ & & \uparrow \partial^\times & & \\ & & C_2^\times & & \end{array}$$

degfilt

3.4. Степенная фильтрация. Вспомним, что для всякого $\alpha = \sum_b \alpha_b b \in C_0$ и всякой функции $f \in \mathcal{V}$, определено целое число $\langle f, \alpha \rangle := \sum_b \alpha_b f(b)$, *билинейно* зависящее от α и f в том смысле, что при фиксированном f получается гомоморфизм из C_0 в $\mathbb{Z}$, а при фиксированном α — гомоморфизм из $\mathcal{V}$ в $\mathbb{Z}$. Ортогональность и ортогональные дополнения определялись по отношению к этому «скалярному произведению». Теперь мы готовы описать степенную фильтрацию на языке цепей и граничных операторов.

Теорема. *Степенная фильтрация выглядит так:*

$$\mathcal{V}_{\leq 0} = (\partial C_1)^\perp, \quad \mathcal{V}_{\leq 1} = (\partial \partial^\times C_2^\times)^\perp, \quad \mathcal{V}_{\leq 2} = \{0\}^\perp.$$

Самая последняя формула — это утверждение о том, что $\mathcal{V}_{\leq 2} = \mathcal{V}$, то есть что любая функция из $\mathcal{V}$ может быть записана как многочлен степени не выше второй от функций Хэвисайда. В самом деле, любая функция из $\mathcal{V}$ ортогональна нулевой 0-цепи. Мы будем доказывать утверждения теоремы в том же порядке, в котором они сформулированы. Части теоремы оформлены ниже в виде нескольких лемм.

Лемма. *Функция $f \in \mathcal{V}$ постоянная тогда и только тогда, когда она ортогональна каждому элементу из ∂C_1 . Иначе говоря, $\mathcal{V}_{\leq 0} = (\partial C_1)^\perp$.*

Доказательство. Пусть сначала f — постоянная функция, принимающая значение c на всех 2-клетках. Достаточно взять произвольную двойственную 1-клетку I с концами a, b (ориентация клетки соответствует указанному порядку концов), и показать, что $\langle f, \partial I \rangle = 0$. Заметим, что $\partial I = b - a$, поэтому $\langle f, \partial I \rangle = f(b) - f(a) = 0$. Наоборот, предположим, что $\langle f, \partial I \rangle = 0$ для всякой двойственной 1-клетки I . Это значит, что f принимает одинаковые значения на концах каждой такой клетки. Однако из любой двойственной 0-клетки можно попасть в любую другую, пройдя по пути из двойственных 1-клеток, и поэтому f должна быть постоянной. $\square$

Прежде чем описывать ортогональное дополнение к $\partial \partial^\times C_2^\times$, опишем образующие этой группы. Зафиксируем ориентацию плоскости и выберем ориентированный флаг $[a, L]$;

отождествим его с образующей группы $C_2^\times$. Нетрудно понять, что представляет из себя элемент $\partial\partial^\times[a, L]$. А именно, это знакопеременная сумма четырех точек, две из которых взяты со знаком плюс, а две со знаком минус. Каких точек: нужно, начиная с точки a , двигаться сначала вдоль L до ближайшей 0-клетки прямой L , а потом уйти в сторону от прямой L в центр соседней 2-клетки; этот центр и будет одной из четырех точек, входящих в $\partial\partial^\times[a, L]$. Таких точек в самом деле четыре: можно двумя способами выбрать направление вдоль L и потом еще двумя способами выбрать клетку, примыкающую к L . Путь от a к описанной 0-клетке конфигурации $\mathcal{A}$ — двузвенная ломаная; ее можно назвать *крюком*. Наконец, знак полученной 0-клетки определяется ориентацией крюка: плюс, если крюк изломан против часовой стрелки и минус в противном случае.

Лемма. *Всякая функция $f \in \mathcal{V}_{\leq 1}$ ортогональна всякому элементу из $\partial\partial^\times C_2^\times$. Наоборот, если $f \in \mathcal{V}$ ортогональна всякому элементу из $\partial\partial^\times C_2^\times$, то $f \in \mathcal{V}_{\leq 1}$. Таким образом, имеет место равенство $\mathcal{V}_{\leq 1} = (\partial\partial^\times C_2^\times)^\perp$.*

Доказательство. Для доказательства первого утверждения достаточно предположить, что $f = \theta_H$ — функция Хэвисайда прямой $H \in \mathcal{A}$, снабженной определенной коориентацией. Рассмотрим одну из образующих $[a, L]$ группы $C_2^\times$, соответствующую произвольному ориентированному флагу $[a, L]$ данной конфигурации. Число $\langle f, \partial\partial^\times[a, L] \rangle$, по определению — сумма значений функции f в четырех точках, входящих в $\partial\partial^\times[a, L]$, со знаками. Из этих четырех значений отличными от нуля могут оказаться четыре, два или ноль. В первом случае результат равен 0 как сумма четырех чисел, два из которых равны единице, а другие два — минус единице. Если слагаемых два, то они, как нетрудно видеть, соответствуют точкам, расположенным по одну и ту же сторону от прямой L , но по разные стороны от точки a или, наоборот, по разные стороны от L , но по одну сторону от a (первый случай реализуется при $L = H$, а второй — если $H \neq L$ проходит через a). Поэтому сумма равна $1 - 1 = 0$. Наконец, если складывать вообще нечего, то есть все значения функции f в рассматриваемых четырех точках равны нулю, то сумма по определению равна 0.

Предположим теперь, что $\langle f, \partial\partial^\times[a, L] \rangle = 0$ для любой образующей $[a, L] \in C_2^\times$. Это значит, что для всех двойственных 1-клеток I , пересекающих прямую L и ориентированных в соответствии с коориентацией прямой L , величина $\langle f, \partial I \rangle$ принимает одно и то же значение. Но это значение составляет как раз скачок функции f при переходе через прямую L в положительном направлении. Данный скачок, таким образом, постоянен вдоль прямой. Иначе говоря, вычитая из f подходящее целочисленное кратное функции Хэвисайда θ_L , мы можем добиться того, что $\langle f, \partial I \rangle = 0$ для всех двойственных 1-клеток I , пересекающихся с L , то есть что f не меняется при переходе через L . Прodelывая описанную операцию для всех прямых $L \in \mathcal{A}$, произвольным образом коориентированных, получаем локально постоянную функцию f на областях конфигурации $\mathcal{A}$, значения которой не меняются при пересечении каждой из прямых $L \in \mathcal{A}$. Легко видеть, что полученная функция будет постоянной. А значит, исходная функция f лежала в $\mathcal{V}_{\leq 1}$. $\square$

Теорема, сформулированная в начале этого раздела, почти доказана. Осталось только убедиться в том, что всякая функция из $\mathcal{V}$ в самом деле представима многочленом от функций Хэвисайда степени не выше второй. Про это — в следующем разделе.

3.5. Конусное представление. В этом разделе мы завершим доказательство теоремы, сформулированной в разделе 3.4. Нам для этого понадобятся некоторые новые понятия, с обсуждения которых мы и начнем. Рассмотрим конфигурацию $\mathcal{A}$ прямых на плоскости. *Конусом* конфигурации $\mathcal{A}$ будем называть всякое непустое пересечение двух замкнутых полуплоскостей, ограниченных пересекающимися прямыми из $\mathcal{A}$. Пусть $\xi : \mathbb{R}^2 \rightarrow \mathbb{R}$ — ненулевая линейная функция на плоскости, а Λ — некоторый конус конфигурации $\mathcal{A}$. Будем говорить, что конус Λ *заострен* в направлении ξ , если функция ξ ограничена на Λ , причем максимум достигается только в вершине. *Индикаторной функцией* $\mathbb{1}_\Lambda$ конуса

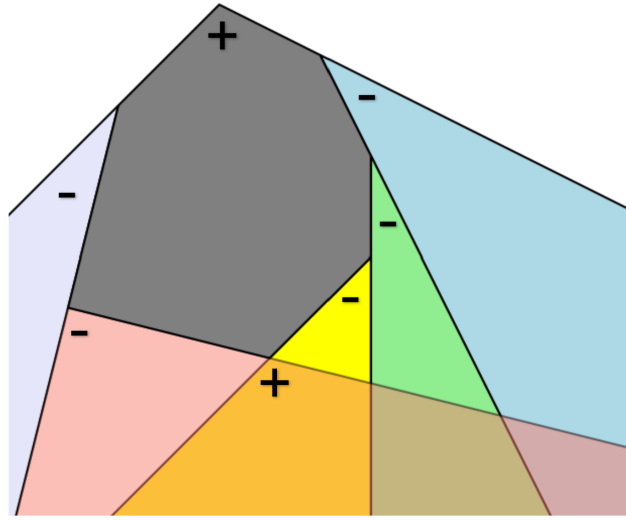


Рис. 6. Выражение индикаторной функции ограниченного многоугольника через индикаторные функции конусов, заостренных в вертикальном направлении.

fig:cone

Λ называется такая функция на плоскости, значения которой во всех точках из Λ равны единице, а во всех остальных точках равны нулю. Эту функцию можно также рассматривать как функцию на 2-клетках конфигурации $\mathcal{A}$, так что мы будем часто считать, что $1_\Lambda \in \mathcal{V}$. Данное только что определение индикаторной функции распространяется на любое множество, целиком состоящее из клеток, в частности, имеет смысл говорить про индикаторную функцию отдельной 2-клетки e и обозначать ее через 1_e . Напомним, что в алгебре $\mathcal{V}$ мы рассматриваем только функции на 2-клетках или, если про эти функции думать как про функции на плоскости, то мы пренебрегаем значениями рассматриваемых функций на клетках размерности < 2 , а на каждой клетке размерности 2 считаем эти функции постоянными. В частности, функции 1_e и $1_{\bar{e}}$ следует рассматривать как одинаковые в алгебре $\mathcal{V}$.

Теорема. *Всякая функция f из $\mathcal{V}$ представляется в виде линейной комбинации с целыми коэффициентами индикаторных функций конусов. Более того, если зафиксировать ненулевую линейную функцию ξ на плоскости, не постоянную ни на какой прямой из $\mathcal{A}$, то f можно представить в виде линейной комбинации индикаторных функций заостренных в направлении ξ конусов, к которой прибавлен определенный элемент из $\mathcal{V}_{\leq 1}$.*

Доказательство. Достаточно, очевидно, проверить последнее утверждение, то есть что f представляется линейной комбинацией индикаторных функций заостренных конусов, по модулю $\mathcal{V}_{\leq 1}$. Сначала рассмотрим ограниченную 2-клетку e . Ее индикаторная функция 1_e получается следующим образом. Возьмем вершину клетки e , в которой $\xi|_{\bar{e}}$ достигает максимума. Около этой вершины $\bar{e}$ совпадает с некоторым конусом Λ ; возьмем 1_Λ со знаком плюс. Далее, рассмотрим вершину, в которой $\xi|_{\bar{e}}$ достигает минимума, и пусть $\bar{e}$ около этой вершины совпадает с конусом V . Через V^* обозначим конус, симметричный конусу V относительно общей вершины этих двух конусов. Функцию 1_{V^*} возьмем тоже со знаком плюс. Для всех остальных вершин клетки e , в качестве соответствующих конусов будем выбирать «внешние углы» данной клетки; все эти конусы пойдут со знаком минус. Легко видеть (см. рис. 6), что полученная описанным образом линейная комбинация индикаторных функций конусов совпадает с $1_{\bar{e}}$, с точностью до значений на клетках размерностей < 2 . То же самое рассуждение работает и в том случае, когда e — неограниченная клетка, на которой функция ξ ограничена сверху. Единственная разница — не надо брать вершину, в которой достигается минимум, поскольку такой вершины просто нет.

Теорема будет полностью доказана, если всякую $f \in \mathcal{V}$ удастся привести по модулю $\mathcal{V}_{\leq 1}$ к такой функции, что на ее носителе ξ ограничена. Каждую прямую $L \in \mathcal{A}$ ориентируем таким образом, чтобы ξ возрастала и стремилась к плюс бесконечности при движении вдоль L в положительном направлении. Если еще выбрать ориентацию плоскости, то все прямые из $\mathcal{A}$ можно занумеровать как $L_1, \dots, L_n$ таким образом, чтобы при $i < j$ единичный вектор положительного направления прямой L_j получался из единичного вектора положительного направления прямой L_i вращением против часовой стрелки. Добавив к f подходящую функцию Хэвисайда, можно считать, что f обращается в нуль *слева от* L_1 , то есть в неограниченной области, примыкающей к L_1 слева, если смотреть из точки с очень большим значением функции ξ , повернувшись лицом в положительном направлении прямой L_1 . Прибавляя еще одну функцию Хэвисайда, можно также добиться того, чтобы f обращалась в нуль также и справа от L_1 . Предположим по индукции, что f уже приведена к такой функции, которая обращается в нуль слева от $L_1, \dots, L_k$. Если f принимает значение c слева от L_{k+1} , а коориентация прямой L_{k+1} выбрана так, что ее положительная сторона находится как раз слева, то $f - c\theta_{L_{k+1}}$ будет обращаться в нуль не только слева от $L_1, \dots, L_k$, но еще и слева от L_{k+1} . Мы таким образом добьемся в конце концов, что f станет нулевой слева и справа от $L_1, \dots, L_n$, то есть что на всех областях, на которых $f \neq 0$, функция ξ будет ограничена сверху, что и требовалось. $\square$

Идея представлять функции из $\mathcal{V}$ в виде линейных комбинаций индикаторных функций конусов присутствует уже в [2].

Замечание (Интегрирование экспоненты по многогранникам). Дальнейшее развитие эта идея получила в [3, 4]. Например, в [4] получено многомерное обобщение классической формулы Эйлера–Маклорена, которое, с другой стороны, служит комбинаторным аналогом формулы Римана–Роха–Хирцебруха, а с третьей стороны, обобщает формулу Пика. Ключевое соображение состоит в том, что экспоненту от $-\xi(q)$ можно легко проинтегрировать, если q пробегает конус, заостренный в направлении ξ . А это позволяет «явно» проинтегрировать экспоненту по любому ограниченному выпуклому многограннику, точнее говоря, свести интеграл к сумме по вершинам.

В качестве первого следствия из доказанной теоремы мы выведем, что $\mathcal{V}_{\leq 2} = \mathcal{V}$.

Лемма. *Любая функция из $\mathcal{V}$ представима как многочлен от функций Хэвисайда степени не выше второй.*

Доказательство. Произвольно выбранная функция $f \in \mathcal{V}$ представляется, согласно доказанной выше теореме, линейной комбинацией индикаторных функций конусов с целыми коэффициентами. С другой стороны, индикаторная функция конуса — это произведение двух функций Хэвисайда. $\square$

Доказательство основной теоремы раздела 3.4 завершено.

3.6. Группы коцепей. В этом разделе мы завершим доказательство основных результатов, анонсированных в разделе 3.1. Язык коцепей, который мы сейчас введем, хотя без него можно было бы и обойтись, полезен в ряде других топологических задач, и он же нам пригодится при рассмотрении конфигураций на торе. Определим k -коцепи как гомоморфизмы из соответствующей группы цепей C_k в $\mathbb{Z}$. Коцепи, конечно, тоже образуют группу по сложению. Мы будем, в соответствии с традицией, обозначать группу k -цепей через C^k . Группа $C_2^\times$, воспринимаемая как группа цепей, тоже задает соответствующую группу коцепей, то есть группу $C_2^\times$ всех гомоморфизмов из $C_2^\times$ в $\mathbb{Z}$. На языке коцепей можно говорить про аддитивную группу $\mathcal{V}$ алгебры Варченко–Гельфанда, а именно, это 0-коцепи, то есть $\mathcal{V} = C^0$. Спаривание $\langle f, \alpha \rangle$ между функциями $f \in \mathcal{V}$ и 0-цепями $\alpha \in C_0$ — частный случай спаривания между k -коцепями f и k -цепями α : в общем случае $\langle f, \alpha \rangle$ означает применение гомоморфизма f к цепи α .

С граничным оператором $\partial : C_k \rightarrow C_{k-1}$ можно связать *кограничный оператор* $d : C^{k-1} \rightarrow C^k$, действующий по формуле $\langle df, \alpha \rangle := \langle f, \partial \alpha \rangle$ (в правой части этой формулы иногда ставят знаки, зависящие от k , но мы этого делать не будем). Аналогичным образом, определенный ранее граничный гомоморфизм $\partial^\times : C_2^\times \rightarrow C_1$ задает соответствующий кограничный гомоморфизм $d_\times : C^1 \rightarrow C_2^\times$. Коцепь вида df называется *кограницей*, а если $df = 0$, то коцепь f называется *коциклом*. Следующая задача, по существу, равносильна задаче из раздела 3.3, но эту равносильность требуется осознать.

Задача. Имеют место следующие утверждения:

- Если 1-коцепь является кограницей, то она же — коцикл, и наоборот, всякий 1-коцикл представляет собой кограницу.
- Всякий 0-коцикл представляет собой постоянную функцию.
- Любая 2-коцепь автоматически является как коциклом, так и кограницей.

Первое из этих утверждений мы все же докажем ниже ввиду его важности. Рассмотрим 1-коцепь $g \in C^1$. Скажем, что g *постоянна вдоль прямых конфигурации* $\mathcal{A}$, если, для всякой коориентированной прямой $L \in \mathcal{A}$, значение $\langle g, I \rangle$ одно и то же для любой двойственной 1-клетки I , пересекающей прямую L , если только ориентация у I соответствует коориентации у L . Важное наблюдение состоит в том, что группа $\mathcal{F}_1$ отождествляется с множеством всех 1-коцепей, постоянных вдоль прямых конфигурации $\mathcal{A}$. Пользуясь этим отождествлением, мы будем считать, что $\mathcal{F}_1 \subset C^1$.

Лемма. Если 1-коцепь $g \in C^1$ лежит в $\mathcal{F}_1$, то есть если она постоянна вдоль прямых конфигурации, то g является коциклом, то есть $dg = 0$.

Доказательство. Достаточно считать, что g принимает нулевые значения на всех прямых из $\mathcal{A}$, кроме одной, скажем, L . Посчитаем теперь $\langle dg, e \rangle$ для произвольной ориентированной двойственной 2-клетки e . Если $e \cap L = \emptyset$, то g принимает нулевые значения на всех граничных двойственных 1-клетках области e , стало быть, $\langle dg, e \rangle = \langle g, \partial e \rangle = 0$. Иначе есть ровно две двойственные 1-клетки $I_\pm$ на границе области e , которые пересекают прямую L . Эти клетки, ориентированные так, как в ∂e , ориентированы в разные стороны относительно прямой L , то есть положительное направление вдоль одной из этих клеток, скажем, I_+ , указывает в положительную сторону от L , а положительное направление вдоль I_- указывает в отрицательную сторону. Тогда величины $\langle g, I_\pm \rangle$ отличаются только знаками, так как g постоянна вдоль L , и мы получаем

$$\langle dg, e \rangle = \langle g, \partial e \rangle = \langle g, I_+ \rangle + \langle g, I_- \rangle = 0,$$

откуда вытекает равенство $dg = 0$, что и требовалось. $\square$

Итак, $\mathcal{F}_1$ состоит из 1-коциклов. Заметим, что для $f \in \mathcal{V}$, кограница df постоянна вдоль прямых конфигурации тогда и только тогда, когда $f \in \mathcal{V}_{\leq 1}$. Это утверждение равносильно ранее установленному равенству $\mathcal{V}_{\leq 1} = (\partial \partial^\times C_2^\times)^\perp$.

Лемма. Всякий 1-коцикл из C^1 является кограницей.

Доказательство. Пусть $g \in C^1$ — коцикл, то есть $dg = 0$. Это значит, что $\langle g, \partial \alpha \rangle = 0$ для всякой 2-цепи $\alpha \in C_2$, то есть значение коцепи g на любой границе равно нулю. Но границы — то же самое, что циклы. Итак, значение коцепи g на любом цикле равно нулю.

Теперь зафиксируем произвольную двойственную 0-клетку q_0 . Для всякой двойственной 0-клетки q , рассмотрим некоторую ломаную γ , ведущую из q_0 в q и состоящую из двойственных 1-клеток. Ориентируем звенья ломаной $\gamma = \gamma_q$ от q_0 в сторону q . Положим $f(q) := \langle g, [\gamma] \rangle$, где $[\gamma] \in C_1$ — это сумма всех звеньев ломаной γ . Заметим, что выражение в правой части не зависит от выбора γ . В самом деле, если γ' — другая ломаная, соединяющая q_0 с q по двойственным 1-клеткам, то $[\gamma] - [\gamma']$ представляет собой цикл, а значит, согласно сказанному выше, $\langle g, [\gamma] \rangle = \langle g, [\gamma'] \rangle$.

Итак, f — корректно определенная функция на двойственных 0-клетках, то есть 0-коцепь. То, что $df = g$, достаточно проверить на произвольной ориентированной двойственной 1-клетке I с вершинами a, b , ориентированной от a к b . С одной стороны, $\langle df, I \rangle = f(b) - f(a)$ по определению кограничного оператора d , а, с другой стороны, выбирая γ_a и γ_b так, чтобы выполнялось $[\gamma_b] = [\gamma_a] + I$, получаем, что $f(b) - f(a) = \langle g, I \rangle$. $\square$

Мы теперь докажем результаты, анонсированные в разделе 3.1. Изоморфизм $\mathcal{V}_1 = \mathcal{F}_1$, фактически, уже был установлен сразу после формулировки — точнее говоря, имеющих там указаний достаточно, чтобы без труда восстановить полное доказательство. Но сейчас мы изложим доказательство по-другому, с более явным привкусом «теории гомологий» (сама эта теория здесь не излагается, но рассматриваемые нами структуры могут служить одной из многочисленных мотивировок для ее изучения).

Гомологическое доказательство изоморфизма $\mathcal{V}_1 = \mathcal{F}_1$. Рассмотрим гомоморфизм из $\mathcal{V}$ в C^1 , заданный формулой $f \mapsto df$. Образ подгруппы $\mathcal{V}_{\leq 1}$ при таком гомоморфизме состоит в точности из тех 1-коцепей, которые постоянны вдоль прямых конфигурации (см. выше). Другими словами, этот образ совпадает с $\mathcal{F}_1$. Осталось проверить, что ядро рассматриваемого гомоморфизма, то есть полный прообраз нуля, совпадает с подгруппой $\mathcal{V}_{\leq 0}$, состоящей из постоянных функций. Но это верно, поскольку $\langle df, I \rangle = 0$ для двойственной 1-клетки I с концами a и b равносильно равенству $f(a) = f(b)$, а функция, принимающая одинаковые значения в центрах любых двух соседних областей, обязана быть постоянной. $\square$

Еще нам нужно доказать описание фактора $\mathcal{V}_2$ как подгруппы в $\mathcal{F}_2$. Заметим, что $\mathcal{F}_2$ отождествляется с пространством коцепей $C^2_\times$, и при этом вложение из $\mathcal{V}_2$ в $\mathcal{F}_2$ выглядит как $f \mapsto d_\times df$. Для доказательства нам понадобится еще один общий факт из серии «коциклы = кограницы».

Лемма. *Всякая коцепь $g \in C^2_\times$ имеет вид $d_\times h$ для некоторой коцепи $h \in C^1$.*

Доказательство. Сначала зафиксируем прямую $L \in \mathcal{A}$ вместе с некоторой коориентацией. Эта прямая разбивается серединами 1-клеток на (ограниченные или неограниченные) интервалы, причем каждый из этих интервалов содержит единственную вершину конфигурации. Для каждого такого ориентированного интервала I прямой L , содержащего вершину a , положим $g_L(I) := g[a, L]$, причем ориентация прямой L в правой части этого равенства выбирается в соответствии с выбранной ориентацией интервала I . Всегда можно подобрать функцию h_L на содержащихся в L серединах 1-клеток таким образом, чтобы выполнялось условие $g_L(I) = h_L(b_+) - h_L(b_-)$ каждый раз, когда I — ограниченная клетка, ведущая из точки b_- в точку b_+ . Такая функция h_L определена с точностью до прибавления постоянной; она получается суммированием значений функции g_L по интервалам. Более точно функцию h_L можно определить так. Ориентацию прямой L (которую мы теперь фиксируем) будем рассматривать как направление слева направо. Значение $h_L(b)$ в левом конце b некоторого интервала I определим как сумму значений функции g_L на всех интервалах, расположенных левее чем I . Если обратить ориентацию у прямой L , то функция h_L поменяет знак.

Функции h_L , определенные сразу для всех коориентированных прямых, задают некоторую 1-коцепь h : значение этой коцепи на ориентированной двойственной 1-клетке равно значению функции h_L на середине этой клетки. Здесь L — это единственная прямая конфигурации, пересекающая данную двойственную 1-клетку, а коориентация прямой определяется ориентацией клетки. Условие $g_L(I) = h_L(b_+) - h_L(b_-)$ переписывается как $g = d_\times h$. $\square$

Теперь мы можем доказать анонсированное ранее описание факторгруппы $\mathcal{V}_2$ как подгруппы в $\mathcal{F}_2$.

Образ группы $\mathcal{V}_2$ в $\mathcal{F}_2$ задается законами взаимности: гомологическое доказательство. Сначала поясним еще раз, в терминах цепей и коцепей, почему 2-коцепь $d_\times df$ обязана

удовлетворять законам взаимности. Рассмотрим гомоморфизм $\varkappa^*$ из $C_\times^2$ в C^2 , который каждую коцепь $g \in C_\times^2$ отправляет в коцепь $\varkappa^*g$, заданную формулой

$$\varkappa^*g(e) := \sum_L g[a, L].$$

Здесь e в левой части — произвольная двойственная 2-клетка с центром в a , а суммирование в правой части ведется по всем прямым $L \in \mathcal{A}$. При этом коориентация прямой L выбирается произвольно, а ориентация выбирается таким образом, чтобы все ориентированные флаги $[a, L]$ определяли выбранную ориентацию клетки e . Нетрудно видеть, что каждое слагаемое в правой части корректно определено, в силу того, что функция g является знакопеременной. Легко видеть, что $\varkappa^*d_\times$, то есть композиция гомоморфизмов $d_\times$ и $\varkappa^*$, равна обычному кограницному оператору $d : C^1 \rightarrow C^2$. Следовательно, $\varkappa^*d_\times df = ddf = 0$ для всякой $f \in \mathcal{V}$, а это и есть законы взаимности для $d_\times df$. Мы здесь воспользовались тем, что $dd = 0$, что, в свою очередь, вытекает из тождества $\partial\partial = 0$.

Наоборот, предположим, что $\varkappa^*g = 0$ для некоторой $g \in C_\times^2 = \mathcal{F}_2$, и будем искать такую $f \in \mathcal{V}$, что $g = d_\times df$. Во-первых, $g = d_\times h$ для некоторой 1-цепи $h \in C^1$, по доказанной выше лемме. Осталось только проверить, что h имеет вид df для некоторой $f \in \mathcal{V}$, то есть является точной цепью. Для этого, в свою очередь, достаточно проверить, что $dh = 0$, так как 1-кограницы совпадают с 1-коциклами. Наконец, если a — центр некоторой двойственной 2-клетки e , то условие $\langle dh, e \rangle = 0$ переписывается как закон взаимности в точке a , то есть как условие $\varkappa^*g(e) = 0$, где e — двойственная 2-клетка с центром в a . Таким образом, $dh = 0$ вытекает из совокупности законов взаимности. $\square$

4. КОНФИГУРАЦИИ ПРЯМЫХ НА ТОРЕ

Зафиксируем систему координат на плоскости $\mathbb{R}^2$. *Стандартная целочисленная решетка* $\mathbb{Z}^2 \subset \mathbb{R}^2$, по определению, состоит из всех точек (или векторов — мы будем отождествлять точки и векторы, поскольку начало координат фиксировано) с целочисленными координатами. Такие точки или векторы называются *целочисленными*. Решетка $\mathbb{Z}^2$ действует на плоскости параллельными переносами, то есть каждый целочисленный вектор $v \in \mathbb{Z}^2$ определяет преобразование $T_v(z) = z + v$. Определим отношение эквивалентности $\sim$ на $\mathbb{R}^2$ следующим образом: две точки z и z' *эквивалентны*, если одна из другой получается переносом на целочисленный вектор. Фактормножество $\mathbb{R}^2/\mathbb{Z}^2$ по этому отношению эквивалентности можно наглядно себе представить следующим образом.

Рассмотрим единичный квадрат $[0; 1]^2$, состоящий из всех точек (x, y) , таких, что $0 \leq x, y \leq 1$. Любая точка плоскости эквивалентна некоторой точке из единичного квадрата. Следовательно, можно наше отношение эквивалентности ограничить на квадрат, и соответствующее фактормножество отождествляется с $\mathbb{R}^2/\mathbb{Z}^2$. Точки из $[0; 1]^2$, как правило, не эквивалентны друг другу. Если, скажем, взять точку строго внутри квадрата и сдвинуть ее на целочисленный вектор, то образ выйдет за пределы квадрата. По этой причине эквивалентными оказываются лишь точки на границе, а именно, точки $(0, y)$ и $(1, y)$, а также точки $(x, 0)$ и $(x, 1)$. Иначе говоря, нижняя сторона квадрата склеивается с верхней стороной, а левая сторона — с правой. Можно сначала склеить горизонтальные стороны, то есть как бы свернуть квадрат в трубочку. А потом концы этой трубочки нужно свести друг с другом и соединить. Получится *тор*, то есть поверхность бублика. Каноническую проекцию из $\mathbb{R}^2$ на тор $\mathbb{R}^2/\mathbb{Z}^2$ будем обозначать через p .

Ниже мы будем часто использовать обозначение X для тора $\mathbb{R}^2/\mathbb{Z}^2$. Мы специально выбрали такой ни на что не намекающий символ (по крайней мере, не намекающий на то, что речь идет именно о торе), поскольку многие вещи, которые будут сказаны про X , годятся не только для тора, но и для многих других топологических многообразий.

4.1. Прямые на торе. *Рациональной прямой* на плоскости $\mathbb{R}^2$ называется такая прямая, которая содержит хотя бы две различные точки из $\mathbb{Z}^2$. В этом случае на той же прямой, на самом деле, будет бесконечно много целых точек. Образы рациональных прямых при

канонической проекции $\mathbf{p} : \mathbb{R}^2 \rightarrow X$ называются *рациональными прямыми на торе*. Конечное множество рациональных прямых на торе называется *конфигурацией прямых на торе*, или просто *конфигурацией на торе*. Мы всегда будем предполагать, что конфигурация прямых на торе состоит не менее чем из двух прямых.

Пусть L — рациональная прямая на торе. Это означает, по определению, что $L = \mathbf{p}(\hat{L})$ для некоторой рациональной прямой $\hat{L}$ на плоскости. Выбор прямой $\hat{L}$ неоднозначен, поскольку, если сдвинуть прямую на любой целочисленный вектор, то от этого проекция не поменяется. Можно, впрочем, зафиксировать конкретную прямую $\hat{L}$, потребовав, например, чтобы она проходила через начало координат. В дальнейшем мы будем использовать обозначение $\hat{L}$ только для этой прямой. Заметим, что $L = \mathbf{p}(T_v \hat{L})$ для всякого вектора $v \in \mathbb{Z}^2$. Обратите еще внимание на то, что тор $\mathbb{R}^2/\mathbb{Z}^2$ можно рассматривать как факторгруппу аддитивной группы $\mathbb{R}^2$ по подгруппе $\mathbb{Z}^2$. При такой интерпретации каждая рациональная прямая на торе является подгруппой, в частности, она проходит через нейтральный элемент всего тора, то есть через $o := \mathbf{p}(0, 0)$.

Задача. Если разрезать тор вдоль рациональной прямой, то получится цилиндр.

Рассмотрим теперь две рациональные прямые на торе, обозначим их через L_1 и L_2 . По определению, $L_i = \mathbf{p}(\hat{L}_i)$ для некоторых рациональных прямых $\hat{L}_i$, $i = 1, 2$, проходящих через начало координат. Следовательно, точка o лежит в пересечении $L_1 \cap L_2$. Две прямые на плоскости пересекаются только в одной точке, но на торе это уже не так: у прямых L_1 и L_2 может быть несколько точек пересечения. Хотя $\hat{L}_1$ и $\hat{L}_2$ пересекаются только по точке $(0, 0)$, может случиться так, что $\hat{L}_1 \cap T_v \hat{L}_2 \not\subset \mathbb{Z}^2$ для некоторого $v \in \mathbb{Z}^2$, и в этом случае $\mathbf{p}(\hat{L}_1 \cap T_v \hat{L}_2)$ — точка пересечения прямых L_1, L_2 в торе, отличная от o .

Как и в случае плоскости, конфигурация $\mathcal{A}$ на торе задает некоторую комбинаторно-геометрическую структуру. Если выкинуть из тора все прямые конфигурации, то останутся компоненты, называемые *областями*. Рассмотрим какую-нибудь область U конфигурации $\mathcal{A}$ и произвольную компоненту $\hat{U}$ полного прообраза $\mathbf{p}^{-1}(U)$.

Предложение. Множество $\hat{U}$, определенное выше, представляет собой ограниченный выпуклый многоугольник, а отображение $\mathbf{p} : \hat{U} \rightarrow U$ взаимно однозначно.

Доказательство. Напомним, что, по нашему предположению, $\mathcal{A}$ содержит хотя бы две различные прямые. Выделим две любые прямые $L_1, L_2 \in \mathcal{A}$. Очевидно, можно подобрать целочисленные векторы v_1, w_1 таким образом, чтобы $\hat{U}$ лежало между прямыми $T_{v_1} \hat{L}_1$ и $T_{w_1} \hat{L}_1$, то есть чтобы эти прямые лежали по разные стороны от $\hat{U}$. Аналогично, можно рассмотреть пару прямых вида $T_{v_2} \hat{L}_2$ и $T_{w_2} \hat{L}_2$ по разные стороны от $\hat{U}$. Четыре прямые $T_{v_1} \hat{L}_1, T_{w_1} \hat{L}_1, T_{v_2} \hat{L}_2$ и $T_{w_2} \hat{L}_2$ ограничивают параллелограмм $\Pi \subset \hat{U}$. Все прямые, получающиеся из прямых вида $\hat{L}$ (где $L \in \mathcal{A}$), целочисленными сдвигами, разбивают параллелограмм Π на конечное число кусков, каждый из которых — ограниченный выпуклый многоугольник. Множество $\hat{U}$ — один из этих кусков.

Если бы отображение $\mathbf{p} : \hat{U} \rightarrow U$ не было взаимно однозначным, в $\hat{U}$ нашлись бы две точки z, w , такие, что $v := w - z \in \mathbb{Z}^2$. В последнем случае рассмотрим прямую $L \in \mathcal{A}$, не параллельную вектору v . Подходящим выбором целого числа m можно добиться, чтобы прямая $\hat{L} + mv$ разделяла бы точки z и w или проходила бы через одну из этих точек. С другой стороны, оба сценария невозможны, так как ни одна прямая вида $\hat{L} + mv$ не может пересекаться с областью $\hat{U}$. $\square$

Благодаря только что доказанному предложению, мы можем думать про области конфигурации $\mathcal{A}$ как про выпуклые многоугольники на плоскости, отождествляя U с $\hat{U}$. Рассмотрим множество прямых

$$\hat{\mathcal{A}} := \{T_v \hat{L} \mid L \in \mathcal{A}, v \in \mathbb{Z}^2\}$$

на плоскости. Это множество прямых бесконечно, поскольку, вместе с каждой прямой, оно содержит и все ее целочисленные сдвиги. Впрочем, $\hat{\mathcal{A}}$ локально конечно в том смысле, что каждое компактное подмножество плоскости пересекается только с конечным числом прямых из $\hat{\mathcal{A}}$. Многоугольник (не обязательно выпуклый), все стороны которых содержатся в прямых из $\hat{\mathcal{A}}$, будет называться $\hat{\mathcal{A}}$ -многоугольником. Интересно изучать $\mathbb{Z}^2$ -инвариантные меры на $\hat{\mathcal{A}}$ -многоугольниках, то есть такие функции μ на $\hat{\mathcal{A}}$ -многоугольниках со значениями в $\mathbb{Q}$, что $\mu(\emptyset) = 0$, и

$$\mu(P) + \mu(Q) = \mu(P \cap Q) + \mu(P \cup Q)$$

для любых $\hat{\mathcal{A}}$ -многоугольников P и Q , и при этом $\mu(T_v P) = \mu(P)$ для всякого $v \in \mathbb{Z}^2$. Скажем, что мера μ *простая*, если $\mu(P)$ для всякого $\hat{\mathcal{A}}$ -многоугольника P с пустой внутренностью.

Простые $\mathbb{Z}^2$ -инвариантные меры на $\hat{\mathcal{A}}$ -многоугольниках образуют *векторное пространство над $\mathbb{Q}$* в том смысле, что их можно не только складывать, но и умножать на рациональные числа, причем сложение и умножение на числа обладают привычными свойствами сложения векторов и их умножения на скаляры. А именно, векторное пространство над $\mathbb{Q}$ — это группа по сложению, элементы которой (называемые *векторами*) можно также умножать на рациональные числа (так что умножение на 1 действует тождественно, а умножение на $cc' \in \mathbb{Q}$ представляется как композиция умножения на c и умножения на c'), причем операции сложения векторов и умножения на числа согласованы в том смысле, что

$$c(\alpha + \beta) = c\alpha + c\beta, \quad (c + c')\alpha = c\alpha + c'\alpha$$

для всяких векторов α, β и всяких рациональных чисел $c, c' \in \mathbb{Q}$.

Как и в случае аффинной конфигурации, можно еще рассмотреть пространство $\mathcal{V}$ всех функций на областях конфигурации $\mathcal{A}$ со значениями в $\mathbb{Q}$. Разница тут, впрочем, не только между плоскостью и тором, но и между областью значений функций: в случае тора мы допускаем в качестве значений не только целые, но и произвольные рациональные числа. Возможность делить одно число на другое (любое ненулевое) нам будет важна.

Предложение. *Имеется естественный изоморфизм между $\mathcal{V}$ и пространством всех простых $\mathbb{Z}^2$ -инвариантных мер на $\hat{\mathcal{A}}$ -многоугольниках.*

Биекция между двумя векторными пространствами называется *изоморфизмом*, если она переводит сложение векторов в сложение векторов, а умножение вектора на скаляр — в умножение вектора на тот же скаляр.

Доказательство. Пусть $f \in \mathcal{V}$; про f можно думать как про локально постоянную функцию на $X \setminus \bigcup \mathcal{A}$. Рассмотрим *поднятие* $\hat{f} : \mathbb{R}^2 \setminus \bigcup \hat{\mathcal{A}} \rightarrow \mathbb{Q}$, определенное формулой $\hat{f}(q) := f(\mathbf{p}(q))$ для всех $q \in \mathbb{R}^2 \setminus \bigcup \hat{\mathcal{A}}$. Каждый $\hat{\mathcal{A}}$ -многоугольник P разбивается прямыми из $\hat{\mathcal{A}}$ на конечное число более мелких многоугольников — будем называть их *плитками* — а последние уже проецируются под действием $\mathbf{p}$ взаимно однозначно на области конфигурации $\mathcal{A}$. Определим $\mu_f(P)$ как сумму значений функции $\hat{f}$ на всех плитках, входящих в P . Легко видеть, что μ_f — простая $\mathbb{Z}^2$ -инвариантная мера на $\hat{\mathcal{A}}$ -многоугольниках. Более того, соответствие $f \mapsto \mu_f$ и задает искомый изоморфизм, поскольку, для однозначного задания простой $\mathbb{Z}^2$ -инвариантной меры на $\hat{\mathcal{A}}$ -многоугольниках, достаточно определить ее на компонентах прообраза любой плитки, причем на разных компонентах она должна принимать (в силу $\mathbb{Z}^2$ -инвариантности) одно и то же значение. $\square$

Вообще говоря, интересно рассматривать простые $\mathbb{Z}^2$ -инвариантные меры на $\hat{\mathcal{A}}$ -многоугольниках в том числе в тех случаях, когда множество прямых $\hat{\mathcal{A}}$ не является локально конечным. Например, естественно рассмотреть вообще все многоугольники с рациональными координатами, то есть выбрать в качестве $\hat{\mathcal{A}}$ множество всех рациональных прямых.

Чтобы решать такие задачи, впрочем, самый главный шаг — разобраться со случаем локально конечного $\hat{\mathcal{A}}$. Мы здесь только этим случаем и ограничимся.

Vc-desc

4.2. Описание пространства $\mathcal{V}$. Наша задача — описать $\mathcal{V}$ как векторное пространство. Как и в случае плоскости, в этом пространстве имеется фильтрация

$$\mathcal{V}_{\leq 0} \subset \mathcal{V}_{\leq 1} \subset \mathcal{V}_{\leq 2} := \mathcal{V}.$$

Однако, по сравнению с $\mathbb{R}^2$, есть существенная разница. В $\mathbb{R}^2$, прямая делит плоскость на две полуплоскости, и по этой причине определяет функцию Хэвисайда. Мы использовали многочлены от функций Хэвисайда, чтобы определить степенную фильтрацию в $\mathcal{V}$. С другой стороны, рациональная прямая на торе не делит тор на две части, и никакой функции Хэвисайда не получается. А степенную фильтрацию все равно можно определить, причем очень естественным способом.

Будем считать, что $\mathcal{V}_{\leq 0}$, по определению, состоит из постоянных функций, то есть из таких функций, которые принимают на всех областях конфигурации одинаковые значения. Это определение находится в полном соответствии с тем, что было для плоскости. Далее, скажем, что функция $f \in \mathcal{V}$ лежит в $\mathcal{V}_{\leq 1}$, если скачок, который функция f делает при переходе через прямую $L \in \mathcal{A}$, зависит только от направления перехода и от выбора самой прямой, но не зависит от того, в какой точке прямой совершился этот переход (за исключением, возможно, вершин конфигурации, в которых скачок вообще не является однозначно определенным). Такое описание подпространства $\mathcal{V}_{\leq 1}$ ранее было следствием определения через функции Хэвисайда, а сейчас оно является самостоятельным определением. Пространство $\mathcal{V}_{\leq 2}$ положим равным всему $\mathcal{V}$, опять по определению. Итак, степенную фильтрацию в $\mathcal{V}$ можно определить однозначным и естественным образом, несмотря даже на то, что про функции Хэвисайда говорить бессмысленно.

Рассмотрим факторгруппы

$$\mathcal{V}_1 := \mathcal{V}_{\leq 1} / \mathcal{V}_{\leq 0}, \quad \mathcal{V}_2 := \mathcal{V} / \mathcal{V}_{\leq 1},$$

которые в нашем случае являются также векторными пространствами (по этой причине их называют факторпространствами). Мы хотим получить описания этих факторпространств по аналогии с теми, которые имеются на плоскости. Следующее определение повторяет аффинный случай почти дословно.

Определение (Пространства знакопеременных функций). Определим $\mathcal{F}_1$ как пространство всех знакопеременных функций на коориентированных прямых конфигурации $\mathcal{A}$ со значениями в $\mathbb{Q}$. Пространство $\mathcal{F}_2$ состоит по определению из всех знакопеременных функций на ориентированных флагах конфигурации $\mathcal{A}$ вида $[a, L]$, где a — вершина, отличная от o , а прямая $L \in \mathcal{A}$ проходит через a и снабжена как ориентацией, так и коориентацией.

Вложение пространства $\mathcal{V}_1$ в $\mathcal{F}_1$ выглядит точно так же, как для плоскости. Каждой функции $f \in \mathcal{V}_{\leq 1}$ соответствует функция $g \in \mathcal{F}_1$, измеряющая скачки функции f при переходе через каждую прямую конфигурации $\mathcal{A}$. Эти скачки все одновременно равны нулю тогда и только тогда, когда $f \in \mathcal{V}_{\leq 0}$. Однако, в отличие от плоскости, образ пространства $\mathcal{V}_1$ в $\mathcal{F}_1$ не совпадает со всем $\mathcal{F}_1$, а задается нетривиальными линейными соотношениями. Происхождение соотношений легко объяснить. Рассмотрим произвольную замкнутую кривую Γ на торе, скажем, гладкую и пересекающую все прямые конфигурации под ненулевыми углами. Тогда можно просуммировать все скачки функции f вдоль кривой Γ . Ясно, что должен получиться нуль. На самом деле, если Γ можно стянуть в торе в точку, то соответствующее соотношение выполнено автоматически, и оно не накладывает никаких ограничений на подпространство $\mathcal{V}_1 \subset \mathcal{F}_1$. Однако на торе есть нестягиваемые кривые — в определенном смысле, независимых таких кривых можно выбрать две (например, параллель и меридиан тора, то есть образы в торе горизонтальной и вертикальной осей),

а все остальные через них выражаются. Следовательно, имеется два независимых линейных уравнения на $\mathcal{V}_1$. Линейное уравнение, выписанное по замкнутой кривой Γ так, как описано выше, называется *соотношением периодов вдоль Γ* .

Теорема. *Пространство $\mathcal{V}_1$ канонически изоморфно подпространству в $\mathcal{F}_1$, заданному двумя независимыми линейными уравнениями, а именно, соотношениями периодов вдоль двух различных произвольным образом выбранных и слегка сдвинутых прямых из $\mathcal{A}$.*

Более формальное определение вложения $\mathcal{V}_1 \hookrightarrow \mathcal{F}_1$ может быть сформулировано на языке цепей и коцепей. Как и в случае плоскости, рассматриваем цепи, порожденные двойственными клетками, и соответствующие коцепи. Только, в отличие от случая плоскости, мы теперь допускаем в качестве коэффициентов произвольные рациональные числа, а не только целые. Следовательно, пространства C_k при $k = 0, 1, 2$, а также пространство $C_2^\times$, являются теперь векторными пространствами над $\mathbb{Q}$. Пространство $\mathcal{F}_1$ отождествляется с подпространством в C^1 , состоящим из всех 1-коцепей, постоянных вдоль прямых конфигурации $\mathcal{A}$. Аналогично лемме из раздела 3.6, можно доказать, что все такие коцепи — коциклы. Вложение $\mathcal{V}_1 \hookrightarrow \mathcal{F}_1$ задано формулой $f \mapsto df$, где $d : C^0 = \mathcal{V} \rightarrow C^1$ — кограничный оператор, а $\mathcal{F}_1$.

Следствие. *Векторное пространство $\mathcal{F}_1$ имеет размерность $n - 2$, где n — число прямых в конфигурации $\mathcal{A}$.*

Используя язык коцепей, можно сказать, что вложение $\mathcal{V}_2 \hookrightarrow \mathcal{F}_2$ выглядит как композиция отображения $f \mapsto d_\times df$ и проекции $C_\times^2 \rightarrow \mathcal{F}_2$ (в отличие от случая плоскости, $\mathcal{F}_2$ отождествляется не с пространством $C_\times^2$, а с некоторым его факторпространством, полученным путем забывания значений коцепи на флагах вида $[o, L]$). Это вложение реализует пространство $\mathcal{V}_2$ как подпространство в $\mathcal{F}_2$. На первый взгляд, описание этого подпространства ничем не отличается от случая плоскости. Уравнения, выделяющие $\mathcal{V}_2$ как подпространство в $\mathcal{F}_2$ — это *законы взаимности* во всех вершинах конфигурации, то есть соотношения на $g \in \mathcal{F}_2 = C_\times^2$ вида $\sum_L g[a, L] = 0$, где суммирование ведется по всем прямым $L \in \mathcal{A}$, проходящим через вершину a .

Теорема. *Отображение $f \mapsto d_\times df$ определяет вложение $\mathcal{V}_2$ в $\mathcal{F}_2$, образ которого задается законами взаимности во всех вершинах конфигурации.*

Это описание пространства $\mathcal{V}_2$ вполне аналогично случаю плоскости. Может показаться, что разница зашита в определение пространства $\mathcal{F}_2$ — мы ведь отказались от рассмотрения флагов вида $[o, L]$, где o — специальная вершина конфигурации, которой соответствуют все целые точки и через которую тем самым проходят все рациональные прямые. Однако, с другой стороны, мы могли бы такую специальную точку добавить и к плоскости. А именно, можно к плоскости добавить одну бесконечно удаленную точку ∞ и получить сферу $\mathbb{S}^2 = \mathbb{R}^2 \cup \{\infty\}$ (понять это равенство можно, например, в терминах стереографической проекции из сферы в плоскость — точкам плоскости соответствуют все точки сферы, кроме одной, а эта последняя — центр проекции — соответствует точке ∞). При таком подходе точка ∞ тоже оказывается специальной вершиной конфигурации, через которую проходят все прямые, и мы отказываемся от рассмотрения флагов вида $[\infty, L]$. Что же получится, если не отказываться от рассмотрения таких флагов? Про это мы поговорим чуть позже.

Следствие. *Пространство $\mathcal{V}_2$ имеет размерность $\sum_a (\nu_a - 1)$, где суммирование ведется по всем вершинам a конфигурации $\mathcal{A}$, за исключением точки o .*

Мы теперь знаем размерности всех факторов $\mathcal{V}_k$ степенной фильтрации. Складывая эти размерности, получим размерность пространства $\mathcal{V}$:

$$\dim(\mathcal{V}) = n - 1 + \sum_{a \neq o} (\nu_a - 1).$$

Здесь в правой части суммирование происходит по всем вершинам a конфигурации $\mathcal{A}$, кроме o . Поскольку $\mathcal{V}$, очевидно, изоморфно $\mathbb{Q}^{f_2}$ (это пространство всех функций на множестве из f_2 элементов), то мы получаем также формулу для f_2 .

Следствие. Число f_2 областей конфигурации из n рациональных прямых на торе $\mathbb{R}^2/\mathbb{Z}^2$ можно посчитать по формуле $n - 1 + \sum_{a \neq o} (\nu_a - 1)$.

Это непосредственный торический аналог формулы Заславского. Его можно доказать, конечно, и не прибегая к рассмотрению пространства $\mathcal{V}$ — достаточно повторить те же рассуждения, которыми мы воспользовались для доказательства оригинальной формулы Заславского. Разница только в «основании индукции»: в торическом случае таким основанием служит утверждение, что любые две различные рациональные прямые выделяют в торе только одну область (ее можно изобразить на плоскости в виде фундаментального параллелограмма группы $\mathbb{Z}^2$).

4.3. Гомологии и когомологии тора. Доказательство сформулированных в предыдущем разделе теорем можно, в принципе, провести по той же схеме и на том же языке, что и в случае плоскости. Что касается языка, то мы рассматриваем пространства цепей C_k при $k = 0, 1, 2$, а также $C_2^\times$, и соответствующие пространства коцепей C^k , $C_\times^2$. Между этими пространствами действуют граничные $(\partial, \partial^\times)$ или кограничные $(d, d_\times)$ операторы. Выше мы видели, что $\partial \circ \partial = 0$, или, эквивалентным образом, что $d \circ d = 0$ — это верно в очень общей ситуации, и в этом месте нет никакой разницы между плоскостью и тором. Равносильное утверждение состоит в том, что всякая граница является циклом, а каждая кограница — коцикл. Для плоскости, по крайней мере в размерности один, это утверждение было верно и в обратную сторону: любой 1-цикл является границей, а 1-коцикл — кограницей. Это уже не так для тора, что связано с нетривиальной топологией тора и что несколько меняет как утверждения, так и доказательства.

Определение (Гомологии и когомологии). Данное определение имеет смысл для C_1 или C^1 в торе, но оно дословно переносится на случай очень широкого класса топологических пространств. Пространство k -границ является подпространством в пространстве k -циклов (сейчас осмысленно думать про $k = 1$, но мы сразу сформулируем определение общим образом). Соответствующее факторпространство называется *пространством k -гомологий* и обозначается через $H_k = H_k(X; \mathbb{Q})$. Аналогичным образом, *пространство k -когомологий* $H^k = H^k(X; \mathbb{Q})$ определяется как факторпространство пространства всех k -коциклов по подпространству всех k -кограниц.

Рассмотрим произвольную прямую L и снабдим ее произвольным образом выбранными ориентацией и коориентацией. Скажем, что двойственная 1-клетка *параллельна* прямой L , если она пересекает ребро конфигурации, примыкающее к L с положительной стороны. Ориентируем эту параллельную двойственную 1-клетку в соответствии с ориентацией прямой L . Далее, составим 1-цепь $[L]$, просуммировав все параллельные прямой L двойственные 1-клетки с указанными ориентациями. Легко видеть, что $[L]$ — цикл, а следовательно, он определяет класс гомологий, то есть элемент пространства H_1 . Назовем $[L]$ *классом прямой L* . Приведенная ниже теорема — результат вычисления когомологий тора в размерности один; это единственная нетривиальная размерность.

Теорема. Для двумерного тора пространство H^1 двумерно: когомологический класс коцикла $\xi \in C^1$ равен нулю тогда и только тогда, когда $\xi[L] = 0$ для всякой прямой L или, что равносильно, для двух геометрически различных прямых.

Доказательство. Выберем две различные прямые $L_1, L_2 \in \mathcal{A}$ и зафиксируем для них какие-либо ориентации и коориентации. Теперь рассмотрим такой коцикл $\xi \in C^1$, что $\xi[L_1] = \xi[L_2] = 0$. Мы хотим найти такую функцию $f \in C^0$ на двойственных 0-клетках, для которой $\xi = df$. Зафиксируем произвольную двойственную 0-клетку q_0 . Любую другую двойственную 0-клетку q можно соединить с q_0 путем Γ , проходящим по двойственным

1-клеткам. Ориентировав Γ по направлению от q_0 к q , заменим эту ломаную суммой $[\Gamma]$ «звеньев», то есть всех входящих в Γ двойственных 1-клеток, ориентированных в соответствии с выбранной ориентацией пути Γ . Положим $f(q) := \xi[\Gamma]$.

Нам нужно прежде всего проверить, что такая функция f корректно определена, т.е. не зависит от выбора ломаной Γ , соединяющей две данные точки q_0 и q . Рассмотрим некоторое поднятие пути Γ , то есть такой путь $\hat{\Gamma}$ на плоскости, который проецируется в Γ при проекции p . Это поднятие соединяет точки $\hat{q}_0$ и $\hat{q}$, такие, что $p(\hat{q}_0) = q_0$ и $p(\hat{q}) = q$. Пусть теперь Γ' — другой путь по двойственным 1-клеткам, выходящий из q_0 и заканчивающийся в q . Можно поднять этот путь так, чтобы соответствующий поднятый путь $\hat{\Gamma}'$ выходил из $\hat{q}_0$. Конечно, нет никаких оснований считать, что конец $\hat{q}'$ пути $\hat{\Gamma}'$ совпадет с $\hat{q}$. Однако, в любом случае точки $\hat{q}'$, $\hat{q}$ проецируются в одну и ту же точку q на торе, а значит, отличаются на целочисленный вектор. Соединим $\hat{q}$ с $\hat{q}'$ ломаной $\hat{B}$ специального типа: сначала эта проекция этой ломаной на тор делает несколько оборотов (в положительном или отрицательном направлении) параллельно прямой L_1 , в потом еще несколько оборотов параллельно прямой L_2 . Легко видеть, что такая ломаная, соединяющая указанные точки, в самом деле существует.

Заметим, что если пройти сначала по $\hat{\Gamma}$, потом по $\hat{B}$ и, наконец, по $\hat{\Gamma}'$ в обратную сторону, то получится замкнутая ломаная на плоскости. Ее можно представить как границу некоторой линейной комбинации плоских многоугольников (упражнение: проведите аккуратно доказательство этого утверждения). Опуская эту замкнутую ломаную обратно на тор, получим 1-границу $[\Gamma] + [B] - [\Gamma']$, где $B := p(B)$. А с другой стороны, поскольку $d\xi = 0$, значение коцикла ξ на этой границе равно нулю. Условие $\xi[B] = 0$ вытекает из сделанных предположений $\xi[L_1] = \xi[L_2] = 0$ и из выбора ломаной $\hat{B}$. Следовательно, $\xi[\Gamma] = \xi[\Gamma']$, что доказывает корректность определения.

Наконец, осталось проверить равенство $df = \xi$. Рассмотрим произвольную двойственную 1-клетку I , ведущую из точки a в точку b . Тогда $f(a) = \xi[\Gamma_a]$, $f(b) = \xi[\Gamma_b]$ для двух ломаных Γ_a , Γ_b , соединяющих точку q_0 с точками a и b , соответственно. Поскольку значение $f(b)$ не зависит от выбора ломаной Γ_b , будем считать, что Γ_b получается добавлением к Γ_a одного звена, а именно, двойственной 1-клетки I . Далее, по определению кограниченного оператора d , получаем

$$df(I) = f(b) - f(a) = \xi([\Gamma_a] + I) - \xi[\Gamma_a] = \xi(I).$$

Это верно для любой двойственной 1-клетки I , а значит, $df = \xi$, что и требовалось доказать. $\square$

Вычисление гомологий нетрудно свести к вычислению когомологий, и наоборот.

Задача. Из доказанной теоремы, выведите в качестве следствия, что пространство H_1 тоже двумерно, и что оно порождается классами любых двух геометрически различных прямых конфигурации.

Теперь мы можем сформулировать описание пространства $\mathcal{V}_1$ более точно. Напомним, что отображение $f \mapsto df$ вкладывает это пространство в пространство $\mathcal{F}_1$ знакопеременных функций на коориентированных прямых конфигурации или, что равносильно, в пространство 1-коцепей, постоянных вдоль прямых конфигурации (последнее пространство мы договорились отождествлять с $\mathcal{F}_1$, так что $\mathcal{F}_1 \subset C^1$). Зафиксируем две геометрически различные прямые $L_1, L_2 \in \mathcal{A}$ вместе с какими угодно их ориентациями и коориентациями.

Теорема. Векторное подпространство в $\mathcal{F}_1$, состоящее из всех таких $\xi \in \mathcal{F}_1$, для которых $\xi[L_1] = \xi[L_2] = 0$, совпадает с $\mathcal{V}_1$.

В самом деле, $\xi \in \mathcal{F}_1$ лежит в $\mathcal{V}_1$ тогда и только тогда, когда ξ — кограница. Последнее условие равносильно уравнениям $\xi[L_1] = \xi[L_2] = 0$, как было доказано выше.

Для примера, укажем некоторые элементы из $\mathcal{F}_1$ явно (на самом деле, указанные элементы даже порождают все $\mathcal{F}_1$ как векторное пространство). Выберем три попарно геометрически различные прямые $L_1, L_2, L_3 \in \mathcal{A}$. Положим $m_i := |L_j \cap L_k|$ для каждой тройки попарно различных индексов $i, j, k = 1, 2, 3$. Рассмотрим три прямые $\hat{L}_i$, которые отображаются на L_i проекцией $\mathbf{p}$, и которые образуют продолжения сторон некоторого треугольника Δ на плоскости. Коориентируем все три прямые так, чтобы Δ находился по положительную сторону от каждой прямой. Ориентируем все три прямые в соответствии с некоторым направлением обхода границы треугольника Δ . Будем теперь искать такую функцию $g \in \mathcal{F}_1$ на коориентированных прямых конфигурации $\mathcal{A}$, которая обращается в нуль на всех прямых, геометрически отличных от L_1, L_2, L_3 . Если положить $g_i := g(L_i)$ для $i = 1, 2, 3$, то соотношения на функцию g из доказанной теоремы выглядят как

$$m_2 g_3 - m_3 g_2 = m_3 g_1 - m_1 g_3.$$

Иначе говоря, тройка (g_1, g_2, g_3) должна получаться из тройки (m_1, m_2, m_3) умножением на некоторое рациональное число.

Пример. В качестве примера можно взять $(g_1, g_2, g_3) = (m_1, m_2, m_3)$; обозначим соответствующую функцию g через θ_{L_1, L_2, L_3} . Эти функции, построенные по тройкам прямых, порождают векторное пространство $\mathcal{F}_1$; их же можно рассматривать как торические аналоги функций Хэвисайда. Впрочем, соответствующие элементы из $\mathcal{V}_{\leq 1}$ определены только с точностью до прибавления констант.

4.4. Законы взаимности. Данное выше определение степенной фильтрации можно в терминах цепей переписать так:

$$\mathcal{V}_{\leq 0} = (\partial C_1)^\perp, \quad \mathcal{V}_{\leq 1} = (\partial \partial^\times C_2^\times)^\perp.$$

Это вполне аналогично случаю плоскости, с той лишь разницей, что теперь это определение — или переформулировка определения — а не следствие из определения в терминах функций Хэвисайда. Выше уже было описано пространство $\mathcal{V}_1$ как подпространство в $\mathcal{F}_1$. А теперь мы опишем $\mathcal{V}_2$ как подпространство в $\mathcal{F}_2$. Напомним, что $\mathcal{F}_2$ состоит из знакопеременных функций на ориентированных флагах вида $[a, L]$, в которых $a \neq o$, а отображение вложения $\mathcal{V}_2 \hookrightarrow \mathcal{F}_2$ выглядит как $f \mapsto d_\times df$. Будем отождествлять пространство $\mathcal{V}_2$ с образом при этом вложении. Мы теперь можем доказать теорему из раздела 4.2, описывающую $\mathcal{V}_2$ как подпространство в $\mathcal{F}_2$, заданное определенными линейными соотношениями.

Доказательство того, что $\mathcal{V}_2 \subset \mathcal{F}_2$ задается законами взаимности. Напомним, что имеется гомоморфизм $\varkappa^* : C_\times^2 \rightarrow C^2$. Его определение для тора ничем не отличается от определения в случае плоскости. Формула $\varkappa^* d_\times d = d^2 = 0$ доказывается точно так же, как и для аффинных конфигураций. А следовательно, любая 2-цепь вида $d_\times df$ обязана удовлетворять всем законам взаимности.

Теперь рассмотрим $g \in \mathcal{F}_2$ и предположим, что для g имеет место закон взаимности в каждой вершине конфигурации $\mathcal{A}$, за исключением o . Утверждается, что в этом случае $g = d_\times h$ для некоторой 1-коцепи $h \in C^1$ (точнее говоря, g является ограничением 2-цепи $d_\times h$ только на те образующие $[a, L]$ пространства $C_2^\times$, для которых $a \neq o$). Это тоже совершенно аналогично лемме из раздела 3.6; доказательство проходит дословно. Единственная разница в том, что делать дальше. В силу законов взаимности, h представляет собой коцикл. Но если раньше (в аффинном случае) из этого сразу вытекало, что h — кограница, то теперь уже не вытекает. Чтобы h было кограницей, дополнительно должны выполняться соотношения $h[L_1] = h[L_2] = 0$ для двух геометрически различных прямых конфигурации $\mathcal{A}$. Вообще говоря, указанные условия могут нарушаться, если h выбрана неудачно. С другой стороны, у нас есть определенная свобода в выборе h . А именно, можно добавить к h такую 1-цепь, которая будет постоянна вдоль каждой прямой конфигурации, а следовательно, которая перейдет в 0 под действием оператора $d_\times$. Легко выбрать эту цепь так, чтобы обеспечить условия $h[L_1] = h[L_2] = 0$ — достаточно даже взять такую цепь из $\mathcal{F}_1$,

которая принимает ненулевые значения только на L_1 и L_2 , а на всех остальных прямых конфигурации обращается в нуль.

Итак, в результате подходящей модификации цепи h , мы можем считать, что $h = df$ для некоторой $f \in C_0$. Тогда $g = d_{\times}h = d_{\times}df$, что и требовалось. $\square$

Альтернативное описание пространства $\mathcal{V}_2$ может быть основано на вложении не в $\mathcal{F}_2$, а в более широкое пространство $C_{\times}^2$. Само вложение задается той же формулой $f \mapsto d_{\times}df$, но, поскольку теперь $\mathcal{V}_2$ вкладывается в большее пространство, его образ задается большим количеством соотношений.

Определение (Законы взаимности на прямых). Зафиксируем прямую $L \in \mathcal{A}$. *Закон взаимности на прямой L* — это условие на функцию $g \in C_{\times}^2$, состоящее в том, что $\sum_{a \in L} g[a, L] = 0$, где суммирование ведется по всем вершинам a конфигурации, лежащим на данной прямой L . Предполагается, что у L выбраны ориентация и коориентация, однако содержание самого закона взаимности от этих выборов не зависит, поскольку при смене ориентации или коориентации у левой части может поменяться только знак.

Образ пространства в $\mathcal{V}_2$ можно теперь описать с использованием комбинации двух типов законов взаимности.

Теорема. *Элемент $g \in C_{\times}^2$ лежит в образе пространства $\mathcal{V}_2$ тогда и только тогда, когда g удовлетворяет законам взаимности во всех точках и на всех прямых конфигурации.*

Заметим сразу же, что такую же теорему можно было бы сформулировать и в евклидовом случае, только вместо точки o нужно добавить точку на бесконечности. Имеется четкая аналогия между этими точками — каждая из них выделяется тем свойством, что лежит на всех прямых рассматриваемой конфигурации.

Схема доказательства. Всех деталей доказательства мы приводить не станем, но обрисуем план. То, что 2-цепь вида $d_{\times}df$ удовлетворяет законам взаимности во всех точках, включая точку o , и на всех прямых, проверяется непосредственно. Первое мы уже проверили, а второе утверждение — дискретный аналог теоремы о том, что если на окружности задана гладкая функция, то интеграл по всей окружности от ее производной равен нулю. Суть здесь в том, что прямые на торе компактны, и с топологической точки зрения представляют собой окружности. Технически, нам надо просуммировать разности, и в этой сумме все члены взаимно уничтожатся.

Обратно, пусть $g \in C_{\times}^2$ удовлетворяет всем законам взаимности, как в точках, так и на прямых, и мы хотим представить g в виде $d_{\times}df$ для некоторой $f \in C^0$. Конечно, для этого нужно два раза проделать «дискретное интегрирование». В первый раз оно осуществляется вдоль прямых, а осуществимость этой операции вытекает из законов взаимности на прямых. Результат этого первого интегрирования — некоторая 1-коцепь $h \in C^1$, которая является коциклом в силу законов взаимности в точках. Неверно, однако, что h автоматически будет кограницей. Чтобы получилась кограница, нужно еще прибавить к h подходящим образом выбранную 1-коцепь, постоянную вдоль прямых конфигурации; этот трюк был описан в предыдущем доказательстве. $\square$

СПИСОК ЛИТЕРАТУРЫ

- | | |
|-------|---|
| [Arn] | [1] В. И. Арнольд, «На сколько частей делят плоскость n прямых?», Матем. просв., 2008, выпуск 12, 95–104. |
| [VG] | [2] А. Н. Варченко, И. М. Гельфанд, «О функциях Хевисайда конфигурации гиперплоскостей», Функц. анализ и его прил., 21:4 (1987), 1–18. |
| [PK] | [3] А. В. Пухликов, А. Г. Хованский, «Конечно-аддитивные меры виртуальных многогранников», Алгебра и анализ, 4:2 (1992), 161–185. |
| [PK2] | [4] А. В. Пухликов, А. Г. Хованский, «Теорема Римана–Роха для интегралов и сумм квазиполиномов по виртуальным многогранникам», Алгебра и анализ, 4:4 (1992), 188–216. |

- Huh12 [5] J. Huh, “Milnor numbers of projective hypersurfaces and the chromatic polynomial of graphs”, J. Amer. Math. Soc. 25 (2012), 907–927.
- HW [6] J. Huh, B. Wang, “Enumeration of points, lines, planes, etc.”, Acta Mathematica 218 (2017), 297–317.
- Mar [7] N. Martinov, “Classification of arrangements by the number of their cells”, Discrete Comput. Geom. 9 (1993), 39–46.
- Zas [8] T. Zaslavsky, “Facing up to Arrangements: Face-Count Formulas for Partitions of Space by Hyperplanes”, Memoirs of the AMS 1:154, 1975.