

Разбор задачи «Качественный отдых»

В этой задаче отдельно нужно рассмотреть первую группу, когда все дни в графике выходные. Тогда при $k = 0$ или $k = 1$ ответ равен 0, а при $k \geq 2$ ответ равен k .

Во всех остальных случаях заметим, что каждый добавляемый выходной день всегда можно сделать днём качественного отдыха, если он будет соседствовать с каким-то другим выходным днём, поэтому каждый отгул увеличивает количество дней качественного отдыха как минимум на 1. Но если есть два изолированных выходных днях, между которыми есть один рабочий день, то взяв отгул в этот рабочий день количество дней качественного отдыха увеличивается на 3 — два существующих выходных и один новый. Разобьём все отдельные выходные дни на пары соседних, посчитаем количество таких пар n_3 . Также посчитаем отдельные выходные дни, не вошедшие в эти пары n_2 , и количество выходных дней, которые уже являются днями качественного отдыха.

Ответ для каждого данного k получается жадным алгоритмом. Сначала нужно выбирать отгулы между парой изолированных выходных дней, что увеличивает количество качественных дней отдыха на 3, таких отгулов может быть не более, чем n_3 . Следующие отгулы будем выбирать так, чтобы они увеличивали количество дней качественного отдыха на 2, присоединяя их к изолированным выходным дням, не вошедшим в пары. Каждый такой отгул будет увеличивать число дней качественного отдыха на 2, и таких отгулов может быть не более n_2 . Каждый из оставшихся отгулов увеличивают ответ на 1.

Если заранее подсчитать значения n_3 , n_2 и уже существующих дней качественного отдыха, то можно отвечать за один запрос за $O(1)$ и суммарная сложность будет $O(n + q)$.

Разбор задачи «Лягушки на болоте»

В задаче просят для каждой вершины графа построенного на точках ответить на вопрос: правда ли она лежит в не двудольной компоненте связности?

Подзадача 1.

Так как компонента связности не двудольная тогда и только тогда, когда в ней есть цикл нечётной длины, в этой подзадаче можно было проверить это любым полным перебором.

Подзадача 2.

В этой подзадаче можно обойти граф и раскрасить его в 2 цвета из каждой вершины за время $O(n^2)$. Граф можно было построить в явном виде.

Подзадача 3.

Здесь подойдет построение графа за $O(n^2)$ и любой обход за $O(n^2)$. Граф можно было построить в явном виде.

Подзадача 4.

Здесь требуется с оптимизировать решение из предыдущей группы по памяти. Самый простой способ это сделать — не хранить граф в явном виде.

Подзадача 5.

В этой подзадаче все точки на одной прямой. Можно показать, что в таком случае необходимо и достаточно проверить нужно ли провести ребра в 2 ближайших точки слева и справа в порядке сортировки по этой прямой. Затем обойти полученный граф за $O(n + m)$. Так как ребер будет $O(n)$ время работы решения $O(n) + O(sort)$.

Подзадачи 6-8.

В этих подзадачах можно обойти граф за $O(n \cdot r^2)$ рассматривая только точки на расстоянии не более r от текущей. В зависимости от эффективности реализация может набирать от 5 до 15 баллов.

Подзадача 9.

То что точки находятся на достаточно большом расстоянии гарантирует, что в графе линейное количество ребер. Для его построения воспользуемся следующей техникой: разобьём плоскость на квадраты со стороной $r/2$. Распределим точки по квадратам, в которые они попадают. Для каждой точки рассмотрим все точки попадающие в квадраты находящиеся на $\pm r$ по x и y от квадрата в котором она находится. Так как в каждом квадрате не более 2 точек, построение графа работает за $O(n)$ или $O(n \cdot \log(n))$ в зависимости от выбранного способа сохранения точек в квадратах. Обходить граф буде так же, как и в 5й подзадаче.

Полное решение.

Чтобы получить полный балл за задачу нужно объединить идеи предыдущих групп. Воспользуемся построением графа из 9й подзадачи и идеей из 5й подзадачи о том что достаточно проводить ребра в небольшое число соседей. Разобьем квадраты с точками на тяжелые, такие в которых хотя бы 3 точки и лёгкие, в которых менее 3. Если точка лежит в тяжелом квадрате, ответ для нее, очевидно, 1. Если точка лежит в легком квадрате, проведем из нее рёбра аналогично технике из 9й подзадачи. Так как для каждой клетки есть не более $25 \cdot 2$ точек из которых попробуют провести ребра в точки в ней, суммарно будет проведено $O(n)$ ребер. Для того, чтобы для решения было достаточно обойти граф аналогично предыдущим подзадачам, проведем петли для всех вершин в тяжелых клетках. Итого $O(n)$ или $O(n \cdot \log(n))$ времени на построение и $O(n)$ на обход графа.

Разбор задачи «Минимизация инверсий»

Обозначим за $M[a:b][c:d]$ подматрицу с левым верхним углом в (a, c) и правым нижним в (b, d) .

Идея динамического программирования

Пусть $dp[i][j]$ – ответ для прямоугольника с левым верхним углом в (i, j) и правым нижним в (n, m) .

Если первым действием была удалена первая строка, то минимальное количество инверсий в итоговом массиве это $D[i][j] + dp[i + 1][j]$, где $D[i][j]$ – количество инверсий внутри $M[i:i][j:m]$ плюс количество инверсий между $M[i:i][j:m]$ и $M[i+1:n][j:m]$ (то есть количество инверсий внутри удалённой части плюс количество инверсий, которые удалённая часть образует с оставшейся подматрицей).

Аналогично, если первым действием был удалён первый столбец, то минимальное количество инверсий в итоговом массиве это $R[i][j] + dp[i + 1][j]$, где $R[i][j]$ – количество инверсий внутри $M[i:n][j:j]$ плюс количество инверсий между $M[i:n][j:j]$ и $M[i:n][j+1:m]$.

При известных $D[i][j]$ и $R[i][j]$ динамика тривиально пересчитывается за $O(nm)$.

Решение за $O(n^2m^2(n + m))$

Значения $D[i][j]$, $R[i][j]$ могут быть вычислены напрямую из определения. Для вычисления одного значения нужно перебрать пару элементов внутри удаляемой строки/столбца (не более $n^2 + m^2$ пар для одного значения) и пару из элемента удаляемой строки/столбца и элемента оставшейся подматрицы (не более $nm(n + m)$ пар для одного значения). Всего нужно вычислить $2nm$ значений.

Итого время работы: $O(nm \cdot (n^2 + m^2 + mn(n + m))) = O(n^2m^2(n + m))$.

Решение за $O(n^2m^2)$

Можно действовать оптимальнее:

- Для вычисления числа инверсий внутри строки/столбца вычислять только разницу соседних значений. Одна разница вычисляется за линию от размера, поэтому суммарно на эту часть будет потрачено $O(nm(n + m))$ времени. Можно оптимизировать эту часть деревом Фенвика, тогда получится $O(nm \log(nm))$ времени на всю таблицу, но в этой подгруппе это не нужно.
- Подсчёт числа инверсий между удаляемой строкой/столбцом можно сделать за $O(nm)$. Достаточно насчитать массив префиксных сумм массива подсчёта одной из частей и пройти по другой. Каждое из действий выполняется за $O(nm)$.

Решение за $O(n^2m \log(nm))$

Оптимизируем предыдущее решение деревом Фенвика. Будем сразу считать все значения $D[i][j]$, $R[i][j]$ для строки. Для этого пройдемся вдоль длинной стороны (чтобы пересчёт работал за длину

короткой), поддерживая массив подсчёта каждой из частей в дереве Фенвика и вычисляя разницу значений за $O(\text{количество добавленных элементов})$.

- Для $R[i][j]$ пройдёмся по строкам, поддерживая остающуюся часть в дереве Фенвика, и вычисляя число инверсий удаляемого столбца проходом по нему. Суммарно будет произведено $O(n^2m)$ запросов прибавления и $O(n^2m)$ запросов суммы на префиксе.
- Для $D[i][j]$ пройдёмся по строкам, поддерживая обе части в дереве Фенвика. После каждого добавления аналогично вычислим число инверсий, образованных только что добавленными элементами (новым элементом удаляемой строки и новым столбцом оставшейся подматрицы). Здесь будет использоваться два дерева Фенвика, к одному будет произведено $O(nm)$ запросов прибавления и $O(n^2m)$ запросов суммы на префиксе, а к другому $O(n^2m)$ запросов прибавления и $O(nm)$ запросов суммы на префиксе.

Можно добиться существенного ускорения, оптимизировав вторую часть: использовать не дерево Фенвика, выполняющее оба типа запросов за $O(\log nm)$, а корневую, которая выполняет один тип запросов за $O(1)$, а другой за $O(\sqrt{nm})$.

Идея симметричного решения

Рассмотрим разницу $R[i][j] - R[i+1][j]$. Это в точности количество инверсий между элементом (i, j) и $M[i:n][j:m]$, плюс количество инверсий между $M[i+1:n][j:j]$ и $M[i:i][j+1:m]$. Обозначим первое за $C[i][j]$, второе за $I[i][j]$.

Заметим, что разность $D[i][j] - D[i][j+1]$ тоже вычисляется аналогичным образом через $C[i][j]$ и $I[i][j]$. Только вместо $I[i][j]$ нужно использовать $(n-i)(m-j) - I[i][j]$, то есть количество пар, которые не образуют инверсию для $I[i][j]$ (мы вычли количество пар элементов, образующих инверсию, из количества всех пар элементов).

Обратите внимание, что значений $C[i][j]$, $I[i][j]$ достаточно для вычисления $R[i][j]$ и $D[i][j]$ за $O(nm)$. Никаких дополнительных тяжёлых (асимптотически больших $O(nm)$) вычислений производить не нужно. Причём значения $I[i][j]$ считаются один раз.

Решение за $O(nm(n + \sqrt{nm}))$

- Значения $C[i][j]$ можно вычислить сканлайном по значениям с 2d деревом Фенвика за $O(nm \log n \log m)$.
- Значения $I[i][j]$ можно вычислить аналогично проходу по строкам из решения за $O(n^2m \log(nm))$. Только на этот раз будет $O(nm)$ запросов изменения и $O(n^2m)$ запросов суммы на префиксе, что позволяет реализовать эту часть за $O(nm(n + \sqrt{nm}))$.

Оптимизация $C[i][j]$

Заметим, что с помощью $C[i][j]$ учитываются инверсии, которые будут присутствовать в итоговом массиве вне зависимости от порядка операций. Так как если клетка a была не левее и не выше клетки b , то клетка a будет идти после клетки b в итоговом массиве.

Можно вычислить суммарно число инверсий по всем таким парам клеток за $O(nm \log(nm))$:

- Выпишем табличку по строкам, по столбцам и посчитаем суммарное число инверсий в двух полученных массивах.
- Рассмотрим пару клеток a, b . Если ни одна из них не была (не строго) левее и выше другой, то в одном массиве a будет идти перед b , а в другом наоборот. Значит от этой пары клеток мы получим ровно одну инверсию к итоговой сумме. Заметим, что таких пар в точности $C_n^2 \cdot C_m^2$. Если же одна была (не строго) левее и выше другой, то если они образовывали инверсию, то

эта инверсия будет учтена два раза в итоговой сумме. Значит мы можем вычислить количество таких пар, образующих инверсию, просто вычтя из общего количества посчитанных инверсий число $C_n^2 C_m^2$ и разделив полученную разницу на два. Это в точности количество инверсий, которые гарантированно (в не зависимости от порядка удаления строк и столбцов) будут в итоговой последовательности.

Решение за $O(nm(n + \sqrt{nm}))$

Решение за $O(nm(n + \sqrt{nm}))$ имеет существенную константу, так как в нём осуществляется $O(n^2m)$ обращений к корневой, каждое из которых работает за $O(1)$, но является обращением к случайному месту в памяти.

Будем делать сканлайн по значениям в матрице. При обработке значения в клетке (i, j) хотим учесть все инверсии, которые оно образовало в $I[i-1][j], I[i-2][j], \dots, I[0][j]$. Пусть $T[i][j] = 0$, если элемент (i, j) ещё не был обработан, и 1 иначе. Пусть $P[i][j]$ – суффиксные суммы $T[i][j]$ в строках. Тогда вклад значения (i, j) в $I[i-1][j]$ это в точности $P[i-1][j]$, вклад в $I[i-2][j]$ это $P[i-2][j]$, аналогично для $I[i-3][j], \dots, I[0][j]$.

Если хранить матрицу по столбцам (то есть так, чтобы столбцы лежали последовательно в памяти), то операцию пересчёта значений $I[i-1][j], I[i-2][j], \dots, I[0][j]$ это поэлементное прибавление последовательного отрезка в памяти к другому последовательному отрезку в памяти. Что имеет сильно меньшую константу, чем обращение к случайному элементу (в пересчёте на элемент).

Значения $P[i][j]$ можно поддерживать явно, обновляя за $O(m)$ при обработке элемента матрицы. Эта часть не может быть одновременно с предыдущей последовательной в памяти (так как это требует хранения матрицы по строкам). Но можно хранить значения $P[i][j]$ в корневой, тогда обновление будет происходить за $O(\sqrt{m})$, что сделает её асимптотически легче предыдущей части при $n \approx m$.

Так как мы сделали самую асимптотически тяжёлую часть решения оптимальнее по константе, то всё решение будет работать значительно быстрее.

Решение за $O(nm(k \cdot n + k\sqrt{k}m))$

Вместо двухуровневой корневой будем использовать k -уровневую.

При почти квадратной табличке оптимально взять $k = 2$, при существенно отличающихся измерениях стоит взять большее k , например можно просто брать $k = 4$ (или использовать одно из предыдущих решений). Конкретный выбор не сильно важен, так как случай квадратной таблички самый тяжёлый.

Разбор задачи «Жизнь программистов»

В 1-й группе $n \leq 20$, поэтому в ней достаточно перебрать все разбиения и найти оптимальное для каждого k .

Во 2-й группе $k = 2$. В ней можно просто перебрать разбиение за $O(n)$ и выбрать оптимальное. Но можно заметить более важное для последующих подгрупп замечание: если первый элемент максимальный, то оптимально в первый отрезок взять все элементы без последнего, а иначе оптимально взять в первый отрезок только первый элемент.

В 3-й группе $k = 3$. Достаточно перебрать первый отрезок и оптимальным образом разбить оставшийся суффикс. Для этого достаточно использовать идею из предыдущей группы.

В 5-й подгруппе можно написать динамику за $O(n^3)$, а именно пусть $dp[i][j]$ — минимальный лексикографический массив, который можно получить, разбив префикс $[a_1, a_2, \dots, a_i]$ на j отрезков. Переходы — это или начать новый отрезок a_{i+1} элементом, или прорелаксировать максимум последнего подотрезка этим элементом.

Перейдем к ключевой идее в задаче, а именно, как для фиксированного k быстро получать оптимальное разбиение. Добавим обозначение $\text{greater}[i]$ — первая позиция $j > i$, такая что $a[j] > a[i]$. Если такой позиции нет, то $\text{greater}[i] = n$ (всё в 0 индексации). Будем жадно набирать разбиение

слева направо. Пусть в текущий момент нужно разбить суффикс i на k отрезков. Если $k = 1$, то последний отрезок уже определён. Если $k = 2$, то оптимально брать отрезок $[i, \min(n-1, \text{greater}[i]) - 1]$. Иначе, пусть в разбиение будет взят отрезок $[i, j - 1]$. Так как $k \geq 3$ максимум на следующем отрезке равен a_j . Максимум на первом отрезке равен a_i , поэтому на j накладываются следующие ограничения:

- $j \leq \text{greater}[i]$, так как иначе максимум на первом отрезке будет неправильным.
- $j \leq n - k + 1$, так как иначе в оставшемся суффиксе нельзя будет уместить оставшийся $k - 1$ отрезок.

Таким образом, в качестве оптимального j выгодно взять минимум на отрезке от $i + 1$ до $\min(n - k + 1, \text{greater}[i])$. Это можно реализовать за $O(n \log n)$ или за $O(n)$ для каждого k , что достаточно, чтобы сдать 6-ю подгруппу. С помощью данной идеи можно также сдать 7-ю и 8-ю группы.

Далее есть два пути. В первом из них можно просто переходить от оптимального разбиения для k отрезков к оптимальному для $(k + 1)$ -го. Разберёмся как именно отличаются оптимальные разбиения для k и $k + 1$. Сначала у них будут совпадать все подотрезки, а потом в какой-то момент, либо жадник для k придет в состояние, когда надо брать 1 или 2 отрезка, либо граница допустимого j из текущего i увеличится на один, из-за чего можно будет получить новый минимум. Если мы берем этот новый минимум, то это означает, что все следующие отрезки будут единичной длины. Скажем, что подотрезок (переход) длинный, если его длина больше единицы, и подотрезок (переход) короткий, если его длина один. Будем следить за всеми длинными подотрезками при изменении $k + 1 \rightarrow k$, какие-то длинные подотрезки удаляются, и добавляется не более одного, то есть всего их $O(n)$ для всех k .

Если сжимать все короткие подотрезки и эффективно их находить, то можно написать решение, работающее $O((n+q) \log n)$, так как сжатых коротких отрезков будет линейно. Их можно эффективно находить с помощью дерева отрезков и сетов. Но данная реализация не является самой простой.

Второй путь следующий: мы так же для всех k отдельно построим массив за $O(n)$. Для этого мы не будем искать минимум на отрезке, а будем переходить к ближайшему справа меньшему элементу, пока он левее чем $\text{greater}[i]$ и $n - k + 1$.

Теперь будем строить ответ параллельно для всех k . Для этого напомним функцию $\text{solve}(pos, lk, rk)$, которая предполагает, что префикс перестановки до pos разбит на подотрезки и это разбиение оптимально для всех k от lk до rk . Сначала отдельно обработаем случай разбиения суффикса начиная с pos на один или два отрезка. Теперь будем параллельно эмулировать жадник для всех оставшихся k на интересном отрезке. Для этого сначала посмотрим на все возможные разрезы, перебрав цепочку ближайших справа меньших элементов. Каждый разрез оптимальный для какого-то отрезка значений k , поэтому можно запускаться рекурсивно.

Если реализовать эту идею наивно, то параллельное построение массивов для всех k будет работать за $O(n^2)$. Но можно применить следующую оптимизацию: в рекурсии можно найти максимальный возрастающий подотрезок начинающийся в pos и добавить сразу много отрезков длины 1 в сжатом виде. Структуру, умеющую добавлять сразу много отрезков длины 1 и отвечать на k -й элемент, можно реализовать обычным стеком и для ответа на запрос использовать бинпоиск. Такая оптимизация позволяет сдать 10-ю группу.

Можно заметить, что rk всегда равно $n - pos$, поэтому чтобы сдать 9-ю группу достаточно отдельно обработать случай, когда суффикс разбивается только на единичные отрезки. Замечание про rk нужно для полного решения.

Давайте заметим, что решение бы работало быстро, если бы при каждом рекурсивном запуске отрезок $[lk, rk]$ разделялся, так как таких разделений не может больше чем $n - 1$. В случайной перестановке ожидаемо каждое деление происходит быстро, поэтому такое решение работает быстро, если эффективно обрабатывать суффикс отрезков длины 1.

Чтобы перейти к полному решению достаточно эффективно находить следующий момент рекурсии, когда отрезок $[lk, rk]$ разделится. Так как $rk = n - pos + 1$, чтобы найти ближайшее деление, достаточно найти минимальное $j \geq pos$, что для некоторых k из отрезка $[lk, rk]$ будет эффективнее

взять в разбиение не отрезок $[j, j]$, а отрезок $[j, \text{less}[j] - 1]$, где $\text{less}[j]$ — ближайший меньший справа элемент. Этот факт как раз следует из того, что $rk = n - pos$. После этого надо разом добавить много отрезков длины 1, после чего произойдёт разделение отрезка, что можно обработать уже явно.

Чтобы проверить, что разделение произойдёт в позиции j , надо проверить, что $\text{less}[j] < \text{greater}[j]$ и $(j - pos + 1) + (n - \text{less}[j] + 1) \geq lk$. Чтобы найти такое минимальное j , достаточно написать бинарные подёмы. Это позволяет находить такое j за $O(\log n)$, что даёт асимптотику $O((n + q) \log n)$.