

А. Что посмотреть?

Условие

В онлайн-кинотеатре пользователю нужно рекомендовать один фильм из нескольких возможных вариантов.

Для этого система подбирает для пользователя пять фильмов-кандидатов и задаёт правило, по которому нужно выбрать один из них. Такую задачу выбора будем называть **запросом**.

Каждый запрос описывается следующими данными:

- `user_id` — идентификатор пользователя;
- `query_type` — тип запроса, то есть правило выбора;
- `c1`, `c2`, `c3`, `c4`, `c5` — идентификаторы пяти различных фильмов-кандидатов.

Каждая строка файла `queries_A.csv` задаёт один запрос: для указанного пользователя из пяти предложенных фильмов нужно выбрать один.

Для каждого запроса нужно вывести `item_id` выбранного фильма.

Формат ввода

Используйте приложенные к задаче файлы:

1. `queries_A.csv`
2. `items_A.csv`
3. `events_A.csv`
4. `item_meta_A.json`
5. `sessions_A.json`

Примечания

К задаче приложен файл `baseline_A.ipynb` с примером чтения входных файлов, фильтрацией запросов по типу для каждой подзадачи и подробным описанием данных.

Формат вывода

Нужно вывести CSV-файл с заголовком

`query_id,item_id`

Для каждого `query_id`, относящегося к соответствующей подзадаче, в ответе должна быть ровно одна строка. Лишних строк быть не должно.

Система оценивания

Максимальный балл за каждую подзадачу — 20.

Решение считается верным, если для каждого `query_id`, относящегося к соответствующей подзадаче, указан правильный `item_id` по правилам задачи.

Итоговые баллы за задачу выставляются по **лучшей** посылке.

A1. Подзадача 1

При решении этой подзадачи следует отфильтровать строки `queries_A.csv`, оставив только запросы типа `favorite_genre`, и вывести ответы только для них.

Для запроса `favorite_genre` для каждого из пяти кандидатов берётся его жанр, после чего для этого жанра вычисляется сумма весов всех событий данного пользователя по фильмам того же жанра, где

- `open` = 1
- `finish` = 2
- `like` = 3

Выбирается кандидат, для жанра которого эта сумма максимальна.

При равенстве:

1. фильм с меньшей `duration`;
2. фильм с меньшим `item_id`.

A2. Подзадача 2

При решении этой подзадачи следует отфильтровать строки `queries.csv`, оставив только запросы типа `actor_match`, и вывести ответы только для них.

Для каждого запроса `actor_match` формируется множество `watched_actors` актёров фильмов, которые пользователь

1. Или посмотрел — `finish`;
2. Или отметил как понравившиеся — `like`.

Рассматриваются только актёры из поля `actors` файла `item_meta_A.json`.

Если поле `actors` отсутствует, его нужно считать пустым списком.

Для кандидата `score` равен числу общих элементов множеств:

множества актёров этого кандидата; множества `watched_actors`.

При равенстве выбирается фильм с минимальным `item_id`.

A3. Подзадача 3

При решении этой подзадачи следует отфильтровать строки `queries_A.csv`, оставив только запросы типа `next_in_session`, и вывести ответы только для них.

Для пользователя рассмотрим все соседние пары фильмов $f_1 \rightarrow f_2$, извлечённые из всех списков `path` в файле `sessions_A.json`.

Если `path` = $[x_1, x_2, \dots, x_k]$, то из него извлекаются пары:

$$\begin{aligned}x_1 &\rightarrow x_2 \\x_2 &\rightarrow x_3 \\&\vdots \\x_{k-1} &\rightarrow x_k\end{aligned}$$

Для кандидата f_2 его `score` — это количество таких пар $f_1 \rightarrow f_2$, где фильм f_1 пользователь когда-либо досмотрел, то есть для него есть событие `finish` в файле `events_A.csv`.

Каждое вхождение пары в `sessions_A.json` учитывается отдельно.

Если у пользователя нет подходящих пар, то `score` всех кандидатов считается равным 0.

Выбирается кандидат с максимальным `score`.

При равенстве выбирается фильм с минимальным `item_id`.

В. Сейсмоактивный остров

Условие

Исследователи из центра сейсмологического мониторинга анализируют записи землетрясений на острове Террамотус. Для каждого события данные собирались на двух станциях наблюдения, расположенных в разных концах острова. На каждой станции установлены два измерительных прибора, размещённые в разных частях станции.

Данные передавались со станции в лабораторию по двум каналам: по оптоволоконному кабелю и по беспроводной связи. Но в беспроводной связи были помехи, и некоторые моменты времени в данных были утеряны.

Каждый прибор фиксирует 10 величин каждую минуту в течении 100 минут.

Чтобы ускорить разбор архива, исследователи привлекли ИИ-агента. Исследователи неаккуратно написали запрос и дали агенту слишком много прав. Затем, в процессе работы агент зачем-то переименовал все файлы. После этого исследователям стало непонятно, какие записи относятся к одному и тому же землетрясению: файлы перемешаны, а исходная группировка утрачена. Исследователи запаниковали, но затем взяли себя в руки и обратились за помощью к школьнику Олегу — эксперту в машинном обучении. Помогите ему кластеризовать данные и починить последствия неконтролируемых экспериментов с ИИ.

Известно, что каждому землетрясению соответствуют ровно 8 наблюдений: по 2 прибора на каждой из 2 станций и по 2 способа передачи. Вам нужно понять, какие наблюдения относятся к одному и тому же землетрясению.

Формат ввода

К задаче прикреплены файлы:

- `data_B.npy`, содержащий массив наблюдений $240 \times 10 \times 100$ — 240 наблюдений по 10 значений в течение 100 минут.
- `baseline_B.ipynb` — ноутбук с базовым решением задачи.
- `submission_B.csv` — пример решения, которое нужно отправить в тестирующую систему.

Формат вывода

Для проверки необходимо загрузить архив `solution_B.zip`.

Архив должен содержать:

1. Файл `submission_B.csv` с двумя колонками:
 - `ID` — номер наблюдения в `data_B.npy`;
 - `target` — предсказанное значение целевой переменной.
2. Файл `solution_B.ipynb` — Jupyter Notebook с вашим решением.

Допускается добавлять в архив дополнительные файлы, необходимые для работы решения. При этом архив должен содержать ровно один файл с расширением `.csv` и ровно один файл с расширением `.ipynb`.

Система оценивания

За эту задачу можно получить до 60 баллов.

Данные разбиты на публичную и приватную части. Когда вы отправляете `submission_B.csv`, вам показывается результат на **публичной** части. После завершения этапа результат будет пересчитан на **приватной** части. Публичная и приватная часть не пересекаются.

После окончания этапа ваша метрика будет приведена к 60-балльной шкале по следующему правилу:

- результат **baseline-решения** со значением $ARI \leq X$ оценивается в **0 баллов**;
- результат со значением $ARI \geq Y$ оценивается в **60 баллов**;
- если значение ARI находится между X и Y , количество баллов вычисляется по формуле линейной интерполяции:

$$\text{Score} = 60 \cdot \frac{ARI - X}{Y - X}.$$

Значения метрик X и Y будут доступны в тестирующей системе.

Итоговые баллы за задачу выставляются по **последней** посылке.

Метрика оценивания точности ответа

В этой задаче используется метрика **ARI (Adjusted Rand Index)**. Чем большее количество пар объектов участник правильно распределяет по кластерам (например, если оба объекта находятся в разных кластерах и участник также относит их к разным кластерам, ИЛИ оба объекта находятся в одном кластере и участник тоже относит их к одному кластеру), тем выше эта метрика. *ARI* принимает значение 0 для случайного разбиения на кластеры, значение 1 для идеально правильного разбиения и может принимать отрицательные значения в случае разбиения хуже случайного.

Пример расчета метрики **ARI** на ‘Python’:

```
from sklearn.metrics import adjusted_rand_score

labels_true = [0, 0, 1, 1, 2, 2]
labels_pred = [1, 1, 0, 0, 2, 2]

ari = adjusted_rand_score(labels_true, labels_pred)
print("ARI =", ari)
```

С. Ошибка новичка

Условие

Машу позвали на летнюю стажировку в научный институт, который занимается исследованием ареалов обитания животных. В качестве примера интересного задания, с которым Маше предстоит работать, исследователи поделились наблюдениями за местами обитания нового вида ленивцев (*Bradypus procrastinator shkolnikus*). Совсем недавно этот вид был выделен из уже известного ранее. Теперь учёные хотят обнаружить и другие потенциальные зоны его обитания.

Для удобства Маши данные были разбиты на `train` и `test`. Поскольку у исследователей есть только собственные данные о локациях, где уже живут ленивцы, они попросили статистическое агентство Южной Америки предоставить данные и по другим локациям, чтобы расширить тестовый датасет. Каждая строка в данных описывает некоторую локацию и содержит измеренные параметры, характеризующие климат, погоду, растительный и животный мир. Работники института пожаловались, что у них никак не получается построить хорошую модель для этой задачи. Маша посмотрела на данные и сразу заметила две типичные ошибки новичка, допущенные сотрудниками института.

Помогите Маше построить модель, которая по имеющимся данным предсказывает, относится ли локация к потенциальной зоне обитания нового вида ленивцев.

Формат ввода

К задаче прикреплены файлы:

- `map_C.json` — файл следующей структуры:

```
{
  "coastline": [...]
}
```

`coastline` — массив координат береговой линии, используемый для визуализации.

- `train_C.csv` — обучающая выборка. Каждый объект содержит признаки локации и целевую переменную `target` (бинарная метка 0 или 1);
- `test_C.csv` — тестовая выборка. Каждый объект содержит идентификатор `id` и признаки локации;
- `baseline_C.ipynb` — ноутбук с базовым решением задачи;
- `submission_C.csv` — пример решения, которое нужно отправить в тестирующую систему.

Формат вывода

Для проверки необходимо загрузить архив `solution_C.zip`.

Архив должен содержать:

1. Файл `submission_C.csv` с двумя колонками:
 - `ID` — идентификатор объекта из тестовой выборки `test`;
 - `target` — предсказанная бинарная метка класса (0 или 1).
2. Файл `solution_C.ipynb` — Jupyter Notebook с вашим решением.

Допускается добавлять в архив дополнительные файлы, необходимые для работы решения. При этом архив должен содержать ровно один файл с расширением `.csv` и ровно один файл с расширением `.ipynb`.

Система оценивания

За эту задачу можно получить до 60 баллов.

Данные разбиты на публичную и приватную части. Когда вы отправляете `submission_C.csv`, вам показывается результат на **публичной** части. После завершения этапа результат будет пересчитан на **приватной** части. Публичная и приватная часть не пересекаются.

После окончания этапа ваша метрика будет приведена к 60-балльной шкале по следующему правилу:

- результат **baseline-решения** со значением $F1 \leq X$ оценивается в **0 баллов**;
- результат со значением $F1 \geq Y$ оценивается в **60 баллов**;
- если значение $F1$ находится между X и Y , количество баллов вычисляется по формуле линейной интерполяции:

$$\text{Score} = 60 \cdot \frac{F1 - X}{Y - X}.$$

Значения метрик X и Y будут доступны в тестирующей системе.

Итоговые баллы за задачу выставляются по **последней** посылке.

Метрика оценивания точности ответа

В этой задаче используется метрика **F1-score**. Участник отправляет бинарные метки класса (0 или 1).

Обозначим:

- TP — количество истинно положительных предсказаний;
- FP — количество ложно положительных предсказаний;
- FN — количество ложно отрицательных предсказаний.

Тогда

$$\text{Precision} = \frac{TP}{TP + FP},$$

$$\text{Recall} = \frac{TP}{TP + FN},$$

а

$$F1 = \frac{2 \cdot \text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}.$$

Значение **F1-score** лежит в диапазоне от 0 до 1. Чем выше значение, тем лучше качество классификации.

Пример расчёта метрики **F1-score** на Python:

```
from sklearn.metrics import f1_score

score = f1_score(y_true, y_pred)
print("F1-score =", score)
```

Д. Наследие Человечества

Условие

Извержение вулкана Везувий в 79 году нашей эры засыпало пеплом и уничтожило город Помпеи и многие другие города, в том числе Геркуланум, где в 1750 году была обнаружена древнеримская вилла с большим количеством античных папирусов. Свитки, хранившиеся в библиотеке, обуглились и превратились в хрупкие чёрные цилиндры. В наши дни учёные стремятся прочесть их, используя методы рентгеновской томографии, компьютерной реконструкции и машинного обучения.

Когда Вадим узнал об этом, он подумал, что тоже хотел бы поучаствовать в процессе расшифровки. Узнав, что учёные уже решают задачу восстановления текста внутри каждого отдельного папируса, он понял, что следующий шаг — восстановить правильный порядок самих папирусов, ведь античные трактаты обычно состояли из нескольких частей.

Поскольку Вадим не знает древнегреческого, он решил сначала рассмотреть аналогичную задачу на русском языке. Для этого он взял сборник стихотворений, часть стихотворений оставил неизменной (`train_D.json`). Каждое из оставшихся стихотворений (`test_D.json`) он разделил на две части, а затем перемешал между собой все левые и правые части. Теперь задача Вадима — восстановить исходные пары и собрать стихотворения обратно.

Формат ввода

К задаче прикреплены следующие файлы:

- `model_D.zip` — веса для большой языковой модели.
- `baseline_D.ipynb` — ноутбук с базовым решением задачи и примером использования большой языковой модели `model_D.zip` для получения эмбедингов текста
- `train_D.json` — дополнительные стихотворения целиком;
- `test_D.json` — перемешанные левые и правые части стихотворений;
- `submission_D.csv` — пример решения, которое необходимо отправить в тестирующую систему.

Формат вывода

Для проверки необходимо загрузить архив `solution_D.zip`.

Архив должен содержать:

1. Файл `submission_D.csv` с двумя колонками:
 - `left_id` — ID левой половины из `test_D.json`;
 - `right_id` — ID правой половины из `test_D.json`.
2. Файл `solution_D.ipynb` — Jupyter Notebook с вашим решением.

Допускается добавлять в архив дополнительные файлы, необходимые для работы решения. При этом архив должен содержать ровно один файл с расширением `.csv` и ровно один файл с расширением `.ipynb`.

Система оценивания

За эту задачу можно получить до 60 баллов.

Данные разбиты на публичную и приватную части. Когда вы отправляете `submission_D.csv`, вам показывается результат на **публичной** части. После завершения этапа результат будет пересчитан на **приватной** части. Публичная и приватная часть не пересекаются.

После окончания этапа ваша метрика будет приведена к 60-балльной шкале по следующему правилу:

- результат **baseline-решения** со значением $\text{Accurasy} \leq X$ оценивается в **0 баллов**;
- результат со значением $\text{Accurasy} \geq Y$ оценивается в **60 баллов**;
- если значение Accurasy находится между X и Y , количество баллов вычисляется по формуле линейной интерполяции:

$$\text{Score} = 60 \cdot \frac{\text{Accurasy} - X}{Y - X}.$$

Значения метрик X и Y будут доступны в тестирующей системе.

Итоговые баллы за задачу выставляются по **последней** посылке.

Метрика оценивания точности ответа

В этой задаче оценивается корректность сопоставления левых и правых частей стихотворений.

Пусть:

- N — общее количество левых частей в `test.json`;
- K — количество правильно восстановленных пар, то есть таких пар `(left_id, right_id)`, которые совпадают с истинным соответствием.

Тогда метрика вычисляется по формуле

$$\text{Accurasy} = \frac{K}{N}.$$

Иными словами, метрика показывает долю правильно восстановленных стихотворений.

Значение **Accurasy** лежит в диапазоне от 0 до 1. Чем выше значение, тем лучше качество решения.

Пример расчёта **Accurasy** на Python:

```
import pandas as pd

merged = true_pairs.merge(pred_pairs, on="left_id", suffixes=("_true",
    "_pred"))
correct = (merged["right_id_true"] == merged["right_id_pred"]).sum()
accuracy = correct / len(merged)

print("Accurasy =", accuracy)
```

Е. Подозрительные пирожные

Условие

Саша часто болеет и пытается понять причину. Он подозревает, что проблема может быть связана с пирожными, которые он регулярно ест.

Недавно Саша купил новую партию пирожных, и среди них могли оказаться необычные, которые для него вредны. Саша очень хочет съесть все пирожные, но не хочет снова заболеть. Для этого Саше нужно научиться определять, какие пирожные отличаются от остальных.

Саша уверен, что подозрительные пирожные можно отличить визуально. Чтобы научиться их находить, он решил обучить свёрточную нейросеть анализировать изображения пирожных. Однако для обучения у Саши есть только «нормальные» пирожные, про которые известно, что они ему точно не вредны. Всего Саша ест 10 видов пирожных, поэтому он обучил нейросеть классифицировать «нормальные» пирожные на 10 классов. Теперь Саше нужно придумать, как с помощью этой модели находить «вредные» пирожные. Саша просит вашей помощи, чтобы найти подозрительные пирожные среди новой партии.

Вам даны изображения пирожных из новой партии и веса предобученной нейросети. Нейросеть обучена только на «нормальных» пирожных и решает задачу классификации на 10 классов. Вам нужно определить, какие пирожные из новой партии отличаются от «нормальных». Известно, что подозрительных пирожных ровно 1000.

Формат ввода

К задаче прикреплены файлы:

- `public_test_package_E.npz` — тестовые изображения размера 32×32 ;
- `model_weights_E.pt` — веса предобученной CNN;
- `baseline_E.ipynb` — базовый ноутбук с примером решения;
- `submission_E.csv` — пример файла ответа.

Формат вывода

Для проверки необходимо загрузить архив `solution_E.zip`.

Архив должен содержать:

1. Файл `submission_E.csv` с двумя колонками:

- `id` — индекс объекта от 0 до $N - 1$;
- `is_outlier` — ваш прогноз:
 - 1, если объект является подозрительным;
 - 0, если объект считается нормальным.

2. Файл `solution_E.ipynb` — Jupyter Notebook с вашим решением.

В файле `submission_E.csv` должно быть ровно 1000 строк со значением `is_outlier = 1`.

Допускается добавлять в архив дополнительные файлы, необходимые для работы решения. При этом архив должен содержать ровно один файл с расширением `.csv` и ровно один файл с расширением `.ipynb`.

Система оценивания

За эту задачу можно получить до 60 баллов.

Если в отправленном файле количество объектов, помеченных как подозрительные, **отличается от 1000, то за задачу выставляется 0 баллов.**

Данные разбиты на публичную и приватную части. Когда вы отправляете `submission_E.csv`, вам показывается результат на **публичной** части. После завершения этапа результат будет пересчитан на **приватной** части. Публичная и приватная часть не пересекаются.

После окончания этапа значение метрики будет приведено к 60-балльной шкале по следующему правилу:

- результат **baseline-решения** со значением $\text{Recall@K} \leq X$ оценивается в **0 баллов**;
- результат со значением $\text{Recall@K} \geq Y$ оценивается в **60 баллов**;
- если значение Recall@K находится между X и Y , количество баллов вычисляется по формуле линейной интерполяции:

$$\text{Score} = 60 \cdot \frac{\text{Recall@K} - X}{Y - X}.$$

Значения метрик X и Y будут доступны в тестирующей системе.

Итоговые баллы за задачу выставляются по **последней** посылке.

Метрика оценивания точности ответа

В этой задаче используется метрика **Recall@K**.

Участник должен отметить ровно $K = 1000$ объектов как подозрительные. Пусть hits@K — число правильно найденных подозрительных объектов среди отмеченных. Тогда

$$\text{Recall@K} = \frac{\text{hits@K}}{K}.$$

Иными словами, чем больше настоящих подозрительных объектов вы найдёте среди выбранных 1000, тем выше результат.