

МИНИСТЕРСТВО НАУКИ И ВЫСШЕГО ОБРАЗОВАНИЯ
РОССИЙСКОЙ ФЕДЕРАЦИИ

ФЕДЕРАЛЬНОЕ ГОСУДАРСТВЕННОЕ БЮДЖЕТНОЕ ОБРАЗОВАТЕЛЬНОЕ
УЧРЕЖДЕНИЕ ВЫСШЕГО ОБРАЗОВАНИЯ
«Рязанский государственный радиотехнический университет им. В.Ф. Уткина»

КАФЕДРА ЭЛЕКТРОННЫЕ ВЫЧИСЛИТЕЛЬНЫЕ МАШИНЫ

ОЦЕНОЧНЫЕ МАТЕРИАЛЫ

по дисциплине

«Основы систем ИИ»

Направление подготовки

09.03.01 – «Информатика и вычислительная техника»

Направленность (профиль) подготовки

«Программно-аппаратное обеспечение вычислительных комплексов и систем искусственного интеллекта»

Уровень подготовки - бакалавриат

Квалификация выпускника – бакалавр

Форма обучения – очная

1 ОБЩИЕ ПОЛОЖЕНИЯ

Оценочные материалы – это совокупность учебно-методических материалов (практических заданий, описаний форм и процедур проверки), предназначенных для оценки качества освоения обучающимися данной дисциплины как части ОПОП.

Цель – оценить соответствие знаний, умений и владений, приобретенных обучающимися в процессе изучения дисциплины, целям и требованиям ОПОП в ходе проведения промежуточной аттестации.

Основная задача – обеспечить оценку уровня сформированности компетенций. Контроль знаний обучающихся проводится в форме промежуточной аттестации. Промежуточная аттестация проводится в форме экзамена. Форма проведения экзамена – тестирование, письменный опрос по теоретическим вопросам и выполнение практического задания.

2 ОПИСАНИЕ ПОКАЗАТЕЛЕЙ И КРИТЕРИЕВ ОЦЕНИВАНИЯ КОМПЕТЕНЦИЙ

Сформированность каждой компетенции (или ее части) в рамках освоения данной дисциплины оценивается по трехуровневой шкале:

- 1) пороговый уровень является обязательным для всех обучающихся по завершении освоения дисциплины;
- 2) продвинутый уровень характеризуется превышением минимальных характеристик сформированности компетенций по завершении освоения дисциплины;
- 3) эталонный уровень характеризуется максимально возможной выраженностью компетенций и является важным качественным ориентиром для самосовершенствования.

Уровень освоения компетенций, формируемых дисциплиной:

Описание критериев и шкалы оценивания тестирования:

Шкала оценивания	Критерий
3 балла (эталонный уровень)	уровень усвоения материала, предусмотренного программой: процент верных ответов на тестовые вопросы от 85 до 100%
2 балла (продвинутый уровень)	уровень усвоения материала, предусмотренного программой: процент верных ответов на тестовые вопросы от 70 до 84%
1 балл (пороговый уровень)	уровень усвоения материала, предусмотренного программой: процент верных ответов на тестовые вопросы от 50 до 69%
0 баллов	уровень усвоения материала, предусмотренного программой: процент верных ответов на тестовые вопросы от 0 до 49%

Описание критериев и шкалы оценивания теоретического вопроса:

Шкала оценивания	Критерий
3 балла (эталонный уровень)	выставляется студенту, который дал полный ответ на вопрос, показал глубокие систематизированные знания, смог привести примеры, ответил на дополнительные вопросы преподавателя
2 балла (продвинутый уровень)	выставляется студенту, который дал полный ответ на вопрос, но на некоторые дополнительные вопросы преподавателя ответил только с помощью наводящих вопросов
1 балл (пороговый уровень)	выставляется студенту, который дал неполный ответ на вопрос в билете и смог ответить на дополнительные вопросы только с помощью преподавателя
0 баллов	выставляется студенту, который не смог ответить на вопрос

Описание критериев и шкалы оценивания практического задания:

<i>Шкала оценивания</i>	Критерий
<i>3 балла</i> (эталонный уровень)	Задача решена верно
<i>2 балла</i> (продвинутый уровень)	Задача решена верно, но имеются неточности в логике решения
<i>1 балл</i> (пороговый уровень)	Задача решена верно, с дополнительными наводящими вопросами преподавателя
<i>0 баллов</i>	Задача не решена

На промежуточную аттестацию выносятся тест, два теоретических вопроса и задача. Максимально студент может набрать 12 баллов. Итоговый суммарный балл студента, полученный при прохождении промежуточной аттестации, переводится в традиционную форму по системе «отлично», «хорошо», «удовлетворительно» и «неудовлетворительно».

Оценка «отлично» выставляется студенту, который набрал в сумме 12 баллов (выполнил все задания на эталонном уровне). Обязательным условием является выполнение всех предусмотренных в течение семестра практических заданий.

Оценка «хорошо» выставляется студенту, который набрал в сумме от 8 до 11 баллов при условии выполнения всех заданий на уровне не ниже продвинутого. Обязательным условием является выполнение всех предусмотренных в течение семестра практических заданий.

Оценка «удовлетворительно» выставляется студенту, который набрал в сумме от 4 до 7 баллов при условии выполнения всех заданий на уровне не ниже порогового. Обязательным условием является выполнение всех предусмотренных в течение семестра практических заданий.

Оценка «неудовлетворительно» выставляется студенту, который набрал в сумме менее 4 баллов или не выполнил всех предусмотренных в течение семестра практических заданий.

3 ПАСПОРТ ОЦЕНОЧНЫХ МАТЕРИАЛОВ ПО ДИСЦИПЛИНЕ

<i>Контролируемые разделы (темы) дисциплины</i>	<i>Код контролируемой компетенции (или её части)</i>	Вид, метод, форма оценочного мероприятия
Тема 1. Задача классификации в ИИ. Метрическая классификация.	УК-11.1 УК-11.2	Зачет
Тема 2. Линейная классификация.	ПК-2.1 ПК-2.2	Зачет
Тема 3. Линейная регрессия.	ПК-2.1 ПК-2.2	Зачет
Тема 4. Байесовские методы классификации.	ПК-18.2 ПК-18.3	Зачет
Тема 5. Кластеризация.	ПК-6.1 ПК-6.2	Зачет
Тема 6. Основы нейронных сетей.	ПК-6.1 ПК-6.2	Зачет
Тема 7. Принципы обучения нейронных сетей.	ПК-18.1	Зачет
Тема 8. Основные архитектуры нейронных сетей.	ПК-18.3	

4 ТИПОВЫЕ КОНТРОЛЬНЫЕ ЗАДАНИЯ ИЛИ ИНЫЕ МАТЕРИАЛЫ

4.1. Промежуточная аттестация в форме экзамена

Код компетенции	Результаты освоения ОПОП Содержание компетенций
УК-11	Способен осуществлять свою трудовую деятельность с учетом определения корректной роли ИИ в различных процессах, критического анализа последствий применения ИИ-технологий, этических принципов (SS-1)

УК-11.1. Определяет ценностные предпосылки, когнитивные искажения, культурно-обусловленные предвзятости в данных, алгоритмах, постановке задач для ИИ

Типовые тестовые вопросы

1. Что является обязательным свойством для функции, чтобы она считалась метрикой?

- a) Неотрицательность
- b) Симметричность
- c) Тождество неразличимых и неравенство треугольника
- d) Все вышеперечисленное

Правильный ответ d)

2. В методе k ближайших соседей (kNN) при каком значении k модель становится более устойчивой к шуму в данных, но при этом границы между классами размываются?

- a) $k = 1$
- b) Малое k (например, 3 или 5)
- c) Большое k (например, 15 или 20)
- d) Значение k не влияет на устойчивость

Правильный ответ c)

3. Какой из перечисленных методов НЕ является вариацией метода ближайших соседей?

- a) Взвешенный kNN (где голос ближайшего соседа имеет больший вес)
- b) Радиусные соседи (RadiusNeighbors)
- c) Метод опорных векторов (SupportVectorMachines)
- d) kNN с различными метриками расстояния (например, Манхэттенское)

Правильный ответ c)

4. Какова основная цель нормализации (например, Min-MaxScaling) данных перед применением метрических методов?

- a) Увеличить скорость вычислений
- b) Удалить все выбросы из набора данных
- c) Привести все признаки к единому масштабу, чтобы признаки с большим размахом не доминировали при расчете расстояния
- d) Увеличить точность модели на 100%

Правильный ответ c)

5. Какой метод быстрого поиска ближайших соседей основан на рекурсивном разбиении пространства данных гиперплоскостями?

- a) KD-дерево (KD-Tree)
- b) Локально-чувствительное хеширование (LSH)
- c) Кластеризация
- d) Полный перебор (Brute-Force)

Правильный ответа)

6. Что является главной целью процедуры кросс-валидации (например, k-foldcross-validation)?

- a) Увеличить объем обучающих данных
- b) Более надежно оценить способность модели к обобщению на новых, невидимых данных
- c) Ускорить процесс обучения модели
- d) Автоматически выбрать лучший алгоритм классификации

Правильный ответb)

Правильный ответ c)

7. Система распознавания лиц, показывающая значительно более высокую точность для мужчин со светлой кожей, чем для женщин с темной кожей, скорее всего, страдает от:

- a) Недостаточной вычислительной мощности
- b) Смещения в обучающих данных (datasetbias), где были недостаточно представлены определенные демографические группы
- c) Слишком маленького значения k в алгоритме kNN
- d) Отсутствия нормализации данных

Правильный ответ b)

8. Что из перечисленного относится к "когнитивным искажениям" на этапе постановки задачи для ИИ?

- a) Разработчик невольно формулирует задачу так, что ее решение будет благоприятствовать одной социальной группе над другой, исходя из своих стереотипов.
- b) Алгоритм kNN всегда выбирает не тот класс.
- c) Евклидова метрика дает неверные результаты для всех категориальных признаков.
- d) Данные не были нормализованы.

Правильный ответ a)

9. Какая из этих метрик расстояния может быть более уместной для данных с преимущественно категориальными признаками?

- a) Евклидово расстояние
- b) Манхэттенское расстояние
- c) Расстояние Хэмминга
- d) Косинусное расстояние

Правильный ответc)

10. Выбор одной метрики расстояния (например, Евклидовой) вместо другой (например, Косинусной) для сравнения объектов может быть примером:

- a) Стохастической градиентной оптимизации
- b) Ценностной предпосылки, которая скрыто влияет на то, что модель будет считать "похожими" объектами

- c) Обязательного шага предобработки данных
- d) Признака переобучения модели

Правильный ответ b)

11. Культурно-обусловленная предвзятость в алгоритме рекомендации новостей может проявляться в том, что:

- a) Алгоритм использует кэширование для ускорения работы.
- b) Алгоритм стабильно рекомендует контент, который укрепляет точку зрения, доминирующую в культуре, где собирались данные, игнорируя альтернативные точки зрения.
- c) Алгоритм требует много оперативной памяти.
- d) Для поиска соседей используется метод полного перебора.

Правильный ответ b)

12. Какая из перечисленных практик НЕ помогает смягчить предвзятость в метрических классификаторах?

- a) Аудит и балансировка обучающего набора данных.
- b) Осознанный выбор и проверка метрики расстояния на предмет fairness (справедливости).
- c) Использование как можно большего значения k , чтобы модель игнорировала данные.
- d) Тестирование модели на разнообразных контрольных выборках, представляющих разные подгруппы.

Правильный ответ c)

Типовые вопросы открытого типа:

1. Как выбор конкретной функции расстояния может отражать ценностные предпосылки разработчика при построении метрического классификатора?
2. Каким образом евклидова метрика может усиливать культурно-обусловленные предвзятости при работе с разнородными данными?
3. Как метрика расстояния скрыто влияет на то, какие объекты алгоритм будет считать «похожими» в контексте социальных данных?
4. В чем могут заключаться когнитивные искажения при выборе метрики для задач классификации людей?
5. Как метод k -ближайших соседей может воспроизводить и усиливать исторические предвзятости, присутствующие в данных?
6. Каким образом выбор параметра k в алгоритме k -NN связан с компромиссом между точностью и воспроизводством системных предубеждений?
7. Как взвешенный метод ближайших соседей может непреднамеренно усиливать дискриминацию против определенных групп?
8. Какие культурные стереотипы могут быть закодированы в результатах работы метрического классификатора?
9. Каким образом нормализация данных может маскировать существующие в них системные предвзятости?
10. Как методы предобработки данных могут непреднамеренно устранять важные культурные особенности, значимые для классификации?
11. В чем заключаются этические риски автоматической предобработки данных без анализа их культурного контекста?
12. Как выбор стратегии обработки пропущенных данных может отражать когнитивные искажения разработчика?
13. Каким образом стандартная кросс-валидация может скрывать низкое качество работы модели для отдельных культурных групп?

14. Как следует модифицировать процедуру кросс-валидации для выявления культурно-обусловленных предвзятостей?
15. Какие метрики оценки качества, помимо точности, следует использовать для обнаружения дискриминационных эффектов?
16. Как алгоритмы быстрого поиска ближайших соседей могут структурно закреплять существующие социальные неравенства?
17. Каким образом индексирование данных для ускорения поиска может делать «невидимыми» определенные группы объектов?
18. При сравнении метрических методов классификации, какие аспекты следует анализировать для выявления скрытых предубеждений?
19. Как ценностные установки разработчика влияют на выбор «лучшего» алгоритма при сравнении метрических методов?
20. Каким образом тестовые данные для сравнения алгоритмов могут содержать культурные стереотипы?

УК-11.2. Применяет методики работы с этическими и социальными рисками, возникающими на разных стадиях жизненного цикла ИИ

Типовые тестовые вопросы

1. Разработчик выбирает Евклидову метрику для системы подбора кандидатов, так как она "естественна" для него. Какая проблема может быть скрыта в этом выборе?

- a) Это технически неверный выбор, так как Евклидова метрика не является метрикой.
- b) Это **ценностная предпосылка**: метрика скрыто определяет, какие профили считаются "похожими", потенциально усиливая сходство с исторически успешными (и, возможно, нерепрезентативными) шаблонами.
- c) Это приведет к неизбежному переобучению модели.
- d) Евклидова метрика несовместима с категориальными данными по своей природе.

Правильный ответ b)

2. Алгоритм kNN для кредитного скоринга, обученный на данных из регионов с преобладанием одного этнического класса, показывает низкую точность для заемщиков из других регионов. В чем коренная причина проблемы?

- a) Слишком маленькое значение k.
- b) **Культурно-обусловленная предвзятость в данных**: модель выучила специфические, локальные паттерны, которые не обобщаются на другие культурные контексты.
- c) Отсутствие нормализации данных о доходе.
- d) Необходимо использовать взвешенный kNN.

Правильный ответ b)

3. При постановке задачи для ИИ-системы отбора в университет формулировка "найти абитуриентов, наиболее похожих на наших лучших выпускников" содержит:

- a) Оптимальную постановку для метода kNN.
- b) **Когнитивное искажение "якорения"** и потенциальную **предвзятость воспроизведения**: система будет увековечивать существующий демографический состав, не оставляя места для "алмазов в грубой оправе" с нестандартными профилями.
- c) Четкое техническое требование.
- d) Необходимость использования косинусной меры.

Правильный ответ b)

4. Нормализация данных (например, Min-MaxScaling) перед применением kNN может непреднамеренно скрыть предвзятость, потому что:

- a) Она всегда полностью ее устраняет.
- b) Она делает данные нечитаемыми для человека.
- c) Она **маскирует систематические различия** между группами, "приводя к общему знаменателю" признаки, по которым группы изначально различались, делая эти различия невидимыми на этапе анализа.
- d) Она замедляет работу алгоритма.

Правильный ответ c)

5. Какая из вариаций kNN может быть более уязвима к унаследованной предвзятости из-за "проклятия размерности" в данных о пользователях?

- a) kNN с большим k.
- b) kNN с малым k.
- c) **Взвешенный kNN**, где ближайшие соседи имеют наибольшее влияние, и в разреженном пространстве "близость" может определяться шумом или смещенными признаками.
- d) Метод не имеет к этому отношения.

Правильный ответ c)

6. При использовании KD-дерева для быстрого поиска соседей в системе рекомендаций книг, сама структура дерева:

- a) Гарантирует справедливые рекомендации для всех пользователей.
- b) **Может неявно кодировать предвзятость**, так как порядок разбиения пространства признаков (например, по жанру, популярности) приоритизирует одни критерии похожести над другими с самого начала.
- c) Не влияет на результат, а только на скорость.
- d) Автоматически балансирует данные.

Правильный ответ b)

7. Кросс-валидация, проведенная на сегментированных по демографическому признаку данных, помогает выявить:

- a) Оптимальное значение k.
- b) Самый быстрый алгоритм поиска соседей.
- c) **"Слепые зоны" модели**: если точность на фолдах, соответствующих меньшинствам, стабильно ниже, это сигнал о предвзятости алгоритма или данных.
- d) Необходимость использования евклидовой метрики.

Правильный ответ c)

8. Сравнивая два метрических классификатора, один показывает чуть более высокую общую точность, но второй – значительно более сбалансированные метрики (precision, recall) для всех защищаемых групп. Какой вывод корректен?

- a) Всегда следует выбирать модель с максимальной общей точностью.
- b) **Ценностный выбор в пользу "справедливости" (fairness)** может обосновать выбор второй модели, даже в ущерб незначительному проценту общей точности.
- c) Модели идентичны, разница вызвана случайностью.
- d) Вторая модель обязательно переобучена.

Правильный ответ b)

9. Исторические данные для обучения модели прогнозирования рецидивов преступлений собраны в условиях системных предубеждений правовой системы. Использование kNN на таких данных приведет к:

- a) Справедливым прогнозам, так как данные объективны.
- b) **Воспроизведению и усилению системной предвзятости:** модель будет считать "похожими" на ранее осужденных людей из тех же групп, на которые была настроена правовая система.
- c) Случайным ошибкам, не связанным с предвзятостью.
- d) Автоматической коррекции исторических несправедливостей.

Правильный ответ b)

10. Выбор в качестве метрики "косинусного сходства" для сравнения резюме кандидатов вместо "евклидова расстояния" является:

- a) Чисто техническим решением, не несущим ценностной нагрузки.
- b) **Ценностной предпосылкой:** косинусное сходство фокусируется на направлении (наборе навыков), а не на абсолютной величине (например, общем количестве мест работы), что может быть более справедливым для кандидатов с нетрадиционным карьерным путем.
- c) Ошибкой, так как для резюме подходит только расстояние Хэмминга.
- d) Признаком переобучения модели.

Правильный ответ b)

Типовые вопросы открытого типа:

1. Опишите методику анализа этических рисков, связанных с выбором функции расстояния для социально-значимых данных.
2. Какие процедуры следует применять для оценки социальных последствий использования различных метрик в классификаторе?
3. Как выявить и документировать потенциальные этические риски, заложенные в самой функции расстояния?
4. Разработайте план этической оценки метрики перед ее применением к персональным данным.
5. Опишите методику проведения этического аудита вариаций метода k-NN на предмет дискриминационных эффектов.
6. Какие социальные риски возникают при использовании взвешенных методов ближайших соседей в социальных приложениях?
7. Разработайте процедуру оценки воздействия параметра k на различные демографические группы.
8. Как следует модифицировать метод ближайших соседей для минимизации социальных рисков?
9. Опишите методику выявления этических проблем при нормализации разнородных социальных данных.
10. Какие социальные риски могут возникнуть на этапе предобработки и как их mitigate?
11. Разработайте руководство по этической предобработке данных для муниципальных сервисов.
12. Как оценить, не приводит ли нормализация к стиранию культурных особенностей данных?
13. Предложите методику интеграции этических метрик в процедуру кросс-валидации.
14. Как модифицировать стратегию кросс-валидации для выявления социальных рисков?
15. Опишите процесс создания этического чек-листа для валидации метрических моделей.
16. Какие социальные показатели следует отслеживать при кросс-валидации моделей, работающих с социальными данными?

17. Разработайте методику оценки социальных последствий использования алгоритмов быстрого поиска в публичных сервисах.
18. Какие этические риски возникают при индексировании данных для ускорения поиска?
19. Опишите процедуру тестирования алгоритмов быстрого поиска на справедливость.
20. Предложите методику сравнительного анализа метрических методов с точки зрения этических последствий.
21. Как интегрировать оценку социальных рисков в процесс выбора между различными метрическими классификаторами?
22. Разработайте framework для сравнения алгоритмов по критериям этичности и социальной приемлемости.

Типовые вопросы открытого типа:

20.

Код компетенции	Результаты освоения ОПОП Содержание компетенций
ПК-2	Способен проектировать и разрабатывать программное обеспечение

ПК-2.1. Проектирует и разрабатывает программное обеспечение

Типовые тестовые вопросы

1. Какой этап жизненного цикла ПО напрямую связан с выбором и обучением «модели классификации» для задачи вроде определения спам-комментариев?

- a) Написание пользовательской документации
- b) Составление технического задания
- c) **Проектирование архитектуры системы и реализация алгоритмов**
- d) Нагрузочное тестирование

Правильный ответ c)

2. При проектировании системы для автоматической категоризации багрепортов (ошибок/новых функций) разработчик предполагает, что их можно разделить по типам простым линейным правилом. Какое понятие он использует?

- a) Объектно-ориентированное программирование
- b) **Принцип линейной делимости**
- c) Паттерн проектирования "Стратегия"
- d) Инкапсуляция

Правильный ответ b)

3. В контексте разработки ПО, «разделяющая гиперплоскость» — это:

- a) Граница между модулями в архитектуре приложения.
- b) **Решающее правило, реализованное в коде, которое на основе признаков (например, ключевых слов в тексте) относит объект к одному из классов.**
- c) Интерфейс между клиентской и серверной частью.
- d) Уровень абстракции в многослойной архитектуре.

Правильный ответ b)

4. Как разработчик может использовать «функцию потерь» при создании алгоритма классификации отзывов на положительные и отрицательные?

- a) Для хранения логов приложения.
- b) Для измерения пропускной способности API.
- c) **В качестве критерия, который алгоритм оптимизации будет минимизировать в процессе обучения, чтобы улучшить точность прогнозов.**
- d) Для расчета нагрузки на базу данных.

Правильный ответ c)

5. Процесс «минимизации эмпирического риска» при разработке ПО — это:

- a) Снижение количества строк кода.
- b) **Настройка параметров модели на обучающих данных так, чтобы она допускала как можно меньше ошибок.**
- c) Минимизация числа внешних зависимостей проекта.
- d) Упрощение пользовательского интерфейса.

Правильный ответ b)

6. Какой тип функции потерь наиболее уместно использовать при разработке классификатора, выдающего вероятности принадлежности к классу (например, вероятность того, что транзакция мошенническая)?

- a) Квадратичная ошибка
- b) **Логистическая потеря (LogLoss)**
- c) Абсолютная ошибка
- d) HingeLoss

Правильный ответ b)

7. Если данные для спринтной планировки задач (срочно/несрочно, сложно/легко) не являются линейно разделимыми, что должен предусмотреть разработчик в архитектуре системы?

- a) Упростить модель, убрав несколько признаков.
- b) **Использовать более сложную модель (например, с ядром) или нейронную сеть, способную уловить нелинейные зависимости.**
- c) Увеличить объем данных, не меняя модель.
- d) Отказаться от автоматической классификации.

Правильный ответ b)

8. При реализации алгоритма линейного классификатора для фильтрации контента, «веса» признаков в модели — это:

- a) **Параметры, которые настраиваются в процессе обучения и определяют важность каждого признака для итогового решения.**
- b) Нагрузка на процессор от вычисления каждого признака.
- c) Статические константы, заданные разработчиком.
- d) Размеры данных, занимаемые каждым признаком в памяти.

Правильный ответ a)

9. Процесс «минимизации эмпирического риска» в коде реализуется с помощью:

- a) Ручного тестирования.
- b) Алгоритмов оптимизации (например, градиентного спуска), которые итеративно обновляют параметры модели.**
- c) Компилятора языка программирования.
- d) Системы контроля версий.

Правильный ответ b)

10. Почему принцип линейной разделимости является важным концептом для разработчика?

- a) Он гарантирует, что код будет компилироваться без ошибок.
- b) Он помогает выбрать подходящий алгоритм для задачи: если данные линейно разделимы, можно начать с простых и эффективных линейных моделей.**
- c) Он определяет структуру базы данных.
- d) Он заменяет необходимость в модульном тестировании.

Правильный ответ b)

11. Какая концепция машинного обучения наиболее точно соответствует задаче контроля сетевого трафика, где система должна решить, является ли пакет данных «разрешенным» или «запрещенным»?

- a) Кластеризация
- b) Модель классификации**
- c) Регрессия
- d) Ассоциативные правила

Правильный ответ b)

12. Если система контроля доступа к ПО логически разделяет попытки входа на «легитимные» и «несанкционированные» с помощью линейного правила, это означает, что данные являются:

- a) Нормально распределенными
- b) Линейно разделимыми**
- c) Автокоррелированными
- d) Гетероскедастичными

Правильный ответ b)

13. В контексте контроля доступа к сетевым устройствам, «разделяющая гиперплоскость» — это:

- a) Физический межсетевой экран
- b) Абстрактная граница в пространстве признаков (например, время доступа, тип запроса), отделяющая легитимные операции от нарушений**
- c) Маршрутизатор, разделяющий сеть на сегменты
- d) Протокол шифрования данных

Правильный ответ b)

14. Функция потерь в модели, классифицирующей сетевые атаки, необходима для:

- a) Количественной оценки ошибки модели, когда она помечает легитимный трафик как атаку или пропускает реальную угрозу**

- b) Ускорения передачи сетевых пакетов
- c) Потери пакетов данных при перегрузке сети
- d) Снижения пропускной способности канала

Правильный ответ > а)

15. Процесс «минимизации эмпирического риска» при обучении системы контроля доступа направлен на:

- a) Максимизацию числа пользователей
- b) **Снижение среднего количества ошибок классификации на исторических данных о попытках доступа**
- c) Минимизацию использования оперативной памяти
- d) Упрощение сетевой топологии

Правильный ответ b)

16. Какая функция потерь наиболее подходит для бинарной классификации сетевых событий на «нормальные» и «вторжение», где важно одинаково штрафовать ложные пропуски и ложные срабатывания?

- a) Функция потерь для регрессии
- b) **Логистическая потеря (LogLoss)**
- c) Функция потерь Хубера
- d) Функция потерь для ранжирования

Правильный ответ b)

17. Если данные о использовании программного обеспечения (время работы, потребляемая память, активность сети) не являются линейно разделимыми на классы «разрешенное использование» и «нарушение», то:

- a) Контроль использования невозможен
- b) **Следует использовать нелинейную модель классификации или преобразовать признаки**
- c) Необходимо увеличить объем данных
- d) Следует отказаться от системы контроля

Правильный ответ b)

18. Какой параметр модели классификации напрямую влияет на положение разделяющей гиперплоскости в системе мониторинга сетевой активности?

- a) Скорость передачи данных
- b) **Веса (weights) признаков в модели**
- c) Размер пакета данных
- d) Время задержки в сети

Правильный ответ b)

19. Процесс минимизации эмпирического риска в системе контроля доступа к устройствам предполагает:

- a) **Поиск таких параметров модели, при которых совокупные потери на примерах из журналов доступа минимальны**
- b) Снижение стоимости оборудования
- c) Минимизацию количества правил фильтрации
- d) Уменьшение времени отклика системы

Правильный ответ а)

20. Какое утверждение о линейной разделимости в контексте контроля ПО является верным?

- а) Это идеализированное понятие; в реальных данных сессии легитимного и запрещенного использования редко разделяются строго линейно**
- б) Всегда достигается при достаточном количестве признаков
- с) Гарантирует 100% точность системы контроля
- д) Зависит только от типа используемого шифрования

Правильный ответ а)

Типовые вопросы открытого типа:

1. Как выбрать тип модели классификации при проектировании архитектуры ПО для задачи прогнозирования отказов оборудования?
2. Какие архитектурные паттерны наиболее подходят для интеграции моделей классификации в микросервисную архитектуру?
3. Как спроектировать интерфейсы для взаимодействия между модулем классификации и другими компонентами системы?
4. Какие требования к данным следует учитывать при проектировании API для модели классификации?
5. Как проверить гипотезу о линейной разделимости данных на этапе проектирования системы машинного обучения?
6. Какие архитектурные решения следует принять, если данные не являются линейно разделимыми?
7. Как проектировать систему, которая автоматически определяет применимость линейных моделей к решаемой задаче?
8. Как визуализировать разделяющую гиперплоскость в пользовательском интерфейсе системы принятия решений?
9. Какие алгоритмы оптимизации следует реализовать для нахождения разделяющей гиперплоскости в реальном времени?
10. Как спроектировать модуль для интерактивного изменения параметров разделяющей гиперплоскости?
11. Как спроектировать систему выбора функции потерь в зависимости от специфики бизнес-задачи?
12. Какие шаблоны проектирования применять для реализации различных функций потерь с возможностью расширения?
13. Как организовать модульное тестирование реализации функций потерь?
14. Как спроектировать систему мониторинга процесса минимизации эмпирического риска в production-среде?
15. Какие механизмы отката следует предусмотреть при ухудшении показателей эмпирического риска?
16. Как архитектурно разделить процесс обучения и инференса для минимизации влияния на производительность?
17. Как интегрировать линейную регрессию в качестве выходного слоя нейронной сети при проектировании архитектуры?
18. Какие преимущества дает использование линейной регрессии в комбинации с нелинейными активационными функциями?
19. Как спроектировать систему для автоматического выбора между простой линейной регрессией и нейросетевой архитектурой?

20. Как реализовать множественную линейную регрессию с поддержкой динамического добавления признаков?
21. Какие методы валидации данных следует заложить в систему при работе с множественными признаками?
22. Как проектировать систему для обработки пропущенных значений в множественной регрессии?
23. Как модель классификации может быть применена для разграничения прав доступа к сетевым ресурсам?
24. Опишите архитектуру системы контроля доступа на основе классификации сетевых событий.
25. Какие признаки сетевой активности наиболее информативны для классификации легитимного и несанкционированного доступа?
26. Как спроектировать модель классификации для обнаружения аномального использования программного обеспечения?
27. При каких условиях данные о использовании сетевых устройств можно считать линейно разделимыми на классы "норма" и "аномалия"?
28. Как проверить гипотезу о линейной разделимости для журналов доступа к сетевым сервисам?
29. Какие методы преобразования признаков можно применить для достижения линейной разделимости в задачах контроля ПО?
30. Как интерпретировать положение разделяющей гиперплоскости в контексте политик контроля доступа?
31. Каким образом изменения в правилах информационной безопасности влияют на положение разделяющей гиперплоскости?
32. Как визуализировать разделяющую гиперплоскость для многомерных данных сетевой активности?
33. Как выбрать функцию потерь для задачи классификации сетевых атак с учетом асимметрии последствий ошибок?
34. Какова должна быть функция потерь при классификации попыток несанкционированного доступа, если ложные positives критичны?
35. Как модифицировать функцию потерь для учета различных типов нарушений политик использования ПО?
36. Как оценить эмпирический риск модели, классифицирующей нарушения политик использования сетевых устройств?
37. Какие методы минимизации эмпирического риска наиболее эффективны для задач обнаружения сетевых аномалий?
38. Как объем и качество данных журналов доступа влияют на процесс минимизации эмпирического риска?
39. Как интегрировать модель классификации в существующую систему контроля сетевых устройств?
40. Какие архитектурные решения необходимо предусмотреть для реального времени классификации сетевых событий?
41. Как обеспечить масштабируемость системы классификации для крупной сетевой инфраструктуры?
42. Какие метрики качества классификации наиболее значимы для задач контроля использования ПО?
43. Как оценивать надежность классификатора в условиях постоянно меняющихся сетевых угроз?
44. Какие процедуры валидации необходимы для модели, классифицирующей доступ к критичным сетевым ресурсам?

Типовые тестовые вопросы

1. Какой современный инструмент наиболее эффективен для быстрого прототипирования модели классификации с использованием принципа линейной разделимости?

- a) Microsoft Visio
- b) Jupyter Notebook с библиотеками Scikit-learn/PyTorch
- c) Adobe Photoshop
- d) MySQL Workbench

Правильный ответ b)

2. Для визуализации разделяющей гиперплоскости обученной модели в пространстве признаков разработчик должен использовать:

- a) Библиотеки визуализации (Matplotlib/Seaborn) в сочетании с методами снижения размерности like PCA
- b) Систему контроля версий Git
- c) Инструмент для создания UML-диаграмм
- d) Средство профилирования кода PyCharm Profiler

Правильный ответ a)

3. Какой инструмент позволяет автоматически отслеживать эксперименты с различными функциями потерь в процессе минимизации эмпирического риска?

- a) MLflow или Weights&Biases
- b) Docker
- c) Postman
- d) Jenkins

Правильный ответ a)

4. Для эффективной минимизации эмпирического риска на больших datasets современные разработчики используют:

- a) Ручной перебор параметров
- b) Фреймворки с автоматическим дифференцированием (TensorFlow/PyTorch) и GPU-ускорение
- c) Таблицы Excel
- d) Текстовые редакторы like Notepad++

Правильный ответ b)

5. Какой инструмент обеспечивает воспроизводимость процесса обучения модели классификации?

- a) Docker-контейнеры с зафиксированными версиями библиотек
- b) Microsoft Word
- c) Google Docs
- d) Файловый менеджер Windows Explorer

Правильный ответ a)

6. Для проверки линейной разделимости данных перед созданием модели разработчик должен применить:

- a) Инструменты exploratory data analysis (Pandas Profiling) и визуализации
- b) Статический анализатор кода Pylint
- c) Инструмент для нагрузочного тестирования JMeter
- d) Средство моделирования баз данных ERwin

Правильный ответ a)

7. Какой современный подход помогает автоматизировать подбор гиперпараметров при минимизации эмпирического риска?

- a) Optuna или Hyperopt
- b) Ручное документирование в Confluence
- c) Регулярные выражения
- d) Модульное тестирование с unittest

Правильный ответ a)

8. Для промышленной реализации модели классификации с регулярным переобучением и мониторингом функции потерь используется:

- a) Airflow или Prefect для оркестрации ML пайплайнов
- b) PowerPoint для презентаций
- c) Виртуальная машина без автоматизации
- d) Ручные скрипты без планировщика

Правильный ответ a)

9. Какой инструмент позволяет эффективно управлять различными версиями датасетов в процессе минимизации эмпирического риска?

- a) DVC (Data Version Control)
- b) Git LFS
- c) Google Drive
- d) Файловая система NTFS

Правильный ответ a)

10. Для отладки процесса обучения и анализа поведения функции потерь разработчик использует:

- a) TensorBoard или Weights & Biases дашборды
- b) Логирование в текстовый файл
- c) Вывод print-операторов в консоль
- d) Статический анализ кода

Правильный ответ a)

Типовые вопросы открытого типа:

1. Какие современные фреймворки машинного обучения наиболее эффективны для быстрого прототипирования моделей классификации и почему?
2. Как использовать Jupyter Notebook в сочетании с библиотеками scikit-learn для исследования принципа линейной разделимости данных?
3. Какие инструменты визуализации позволяют эффективно анализировать положение разделяющей гиперплоскости в многомерном пространстве признаков?
4. Как средства автоматического машинного обучения (AutoML) могут помочь в подборе оптимальной функции потерь для конкретной задачи?

5. Какие библиотеки оптимизации (SciPy, CVXPY) наиболее подходят для реализации различных стратегий минимизации эмпирического риска?
6. Как использовать инструменты профилирования кода для оптимизации вычислений в процессе минимизации эмпирического риска?
7. Какие системы управления экспериментами (MLflow, Weights&Biases) эффективны для отслеживания процесса минимизации риска?
8. Как современные фреймворки глубокого обучения (PyTorch, TensorFlow) реализуют линейную регрессию в качестве базового строительного блока?
9. Какие инструменты позволяют эффективно реализовать множественную линейную регрессию с автоматическим дифференцированием?
10. Как использовать библиотеки автоматического дифференцирования для реализации метода наименьших квадратов с различными функциями потерь?
11. Какие численные библиотеки (NumPy, SciPy) предоставляют наиболее устойчивые реализации сингулярного разложения для больших матриц?
12. Как использовать специализированные библиотеки линейной алгебры (CuPy, JAX) для ускорения вычислений МНК на GPU?
13. Какие инструменты диагностики мультиколлинеарности доступны в современных библиотеках анализа данных?
14. Какие встроенные механизмы регуляризации предоставляют современные фреймворки для предотвращения переобучения в линейных моделях?
15. Как использовать инструменты кросс-валидации из scikit-learn для оценки устойчивости линейной регрессии к переобучению?
16. Какие библиотеки предоставляют готовые реализации линейной монотонной регрессии с настраиваемыми ограничениями?
17. Как использовать Docker для создания воспроизводимых окружений при разработке моделей линейной регрессии?
18. Какие инструменты CI/CD наиболее эффективны для автоматизации тестирования и развертывания линейных моделей?
19. Как системы управления версиями моделей (DVC) интегрируются с пайплайнами линейной регрессии?
20. Какие распределенные вычислительные **фреймворки** (ApacheSpark, Dask) эффективны для обучения линейных моделей на больших данных?
21. Как использовать инструменты оптимизации памяти при работе с большими матрицами в задачах линейной регрессии?
22. Какие библиотеки предоставляют GPU-ускоренные реализации алгоритмов линейной регрессии?

Код компетенции	Результаты освоения ОПОП Содержание компетенций
ПК-6	Способен организовывать хранение данных, выбирая адекватные технологические решения (BD-3)

ПК-6.1. Разрабатывает, отлаживает и тестирует прикладные решения с элементами ИИ с применением различных технологий хранения структурированных данных, оценивает качество построенных прикладных решений

Типовые тестовые вопросы

1. При разработке системы сегментации пользователей для интернет-магазина наиболее подходящим алгоритмом будет:

а) Линейная регрессия

- b) Алгоритмы кластеризации (например, K-means)
- c) Дерево решений
- d) Сверточная нейронная сеть

Правильный ответ b)

2. Для оценки качества кластеризации при отладке системы рекомендаций товаров используется:

- a) Accuracy
- b) Функционалы качества кластеризации (SilhouetteScore, Calinski-HarabaszIndex)
- c) F1-score
- d) ROC-AUC

Правильный ответ b)

3. При тестировании системы распознавания образов разработчик выбрал сигмоиду в качестве функции активации потому, что:

- a) Она преобразует выход в диапазон (0,1), что удобно для интерпретации как вероятности
- b) Она не имеет проблемы исчезающего градиента
- c) Она всегда работает быстрее других функций
- d) Она линейна

Правильный ответ a)

4. Для хранения структурированных данных о параметрах нейронной сети наиболее эффективно использовать:

- a) Текстовые файлы
- b) Реляционные БД (PostgreSQL) или специализированные форматы (HDF5)
- c) XML-файлы
- d) Кэш в оперативной памяти

Правильный ответ b)

5. При разработке системы классификации изображений архитектура нейронной сети должна включать:

- a) Сверточные слои для извлечения пространственных 特征
- b) Только полносвязные слои
- c) Рекуррентные слои
- d) Один перцептрон Розенблатта

Правильный ответ a)

6. Для отладки работы иерархического алгоритма кластеризации разработчик использует:

- a) Дендрограммы для визуализации процесса объединения кластеров
- b) Матрицу корреляции
- c) Гистограммы распределения
- d) Линейные графики

Правильный ответ a)

7. При оценке вычислительных возможностей нейронной сети для обработки текстов учитывают:

- a) **Способность сети изучать сложные нелинейные зависимости в последовательностях**
- b) Только количество нейронов
- c) Только глубину сети
- d) Скорость работы на CPU

Правильный ответ а)

8. Для тестирования работы перцептрона Розенблатта на различных типах данных используют:

- a) **Наборы данных с линейной разделимостью**
- b) Только изображения высокого разрешения
- c) Только текстовые данные
- d) Случайные шумовые данные

Правильный ответ а)

9. При разработке системы анализа временных рядов применяют статистические алгоритмы кластеризации, потому что они:

- a) **Учитывают вероятностные распределения данных**
- b) Всегда работают быстрее эвристических
- c) Не требуют настройки параметров
- d) Гарантируют оптимальное решение

Правильный ответ а)

10. Для хранения и управления версиями обученных моделей нейронных сетей используют:

- a) **MLflow или Weights&Biases**
- b) Только локальные папки
- c) Реляционные БД без схемы версионирования
- d) Файловые системы без метаданных

Правильный ответа)

11. При отладке эвристических алгоритмов кластеризации важно тестировать на данных:

- a) **С различной плотностью и формой кластеров**
- b) Только с нормальным распределением
- c) Только с малым объемом
- d) Только с большим количеством признаков

Правильный ответ а)

12. Архитектура нейронной сети для обработки естественного языка должна включать:

- a) **Рекуррентные или трансформерные слои**
- b) Только сверточные слои
- c) Только входной и выходной слои
- d) Линейные активационные функции

Правильный ответ а)

Типовые вопросы открытого типа:

1. Как выбрать подходящий алгоритм кластеризации для сегментации пользователей в веб-приложении и интегрировать его в существующую архитектуру?
2. Какие метрики качества кластеризации следует использовать при тестировании системы рекомендаций на основе кластеров?
3. Как спроектировать пайплайн обработки данных для эвристических алгоритмов кластеризации в реальном времени?
4. Каким образом иерархические алгоритмы кластеризации можно применить для анализа логов структурированных данных в корпоративной системе?
5. Как оценить эффективность статистических алгоритмов кластеризации при работе с транзакционными данными из реляционной БД?
6. Как архитектура нейронной сети должна учитывать структуру хранения обучающих данных в SQL-базе?
7. Каким образом функции активации влияют на способность сети извлекать паттерны из структурированных табличных данных?
8. Как провести отладку перцептрона при работе с неполными или зашумленными данными из корпоративных источников?
9. Какие вычислительные возможности нейронных сетей наиболее важны при обработке больших объемов структурированных данных в реальном времени?
10. Как организовать взаимодействие между кластеризационным алгоритмом и системой управления базами данных для минимизации передачи данных?
11. Какие архитектурные решения позволяют эффективно сочетать нейросетевую обработку с транзакционными операциями в БД?
12. Как спроектировать систему кэширования промежуточных результатов кластеризации для ускорения работы приложения?
13. Как разработать стратегию тестирования кластеризационных алгоритмов на реалистичных структурированных данных?
14. Какие методы оценки качества нейросетевой модели применить при работе с импутированными данными из БД?
15. Как провести А/В тестирование различных архитектур нейронных сетей на продакшн-данных без нарушения работы системы?
16. Как оптимизировать работу иерархических алгоритмов кластеризации при больших объемах структурированных данных?
17. Какие подходы к индексации данных в БД ускоряют процесс кластеризации без потери качества?
18. Как распределить вычисления в нейросетевой модели между приложением и СУБД для максимизации производительности?
19. Как организовать предобработку структурированных данных для статистических алгоритмов кластеризации в промышленном пайплайне?
20. Какие методы обработки пропущенных значений в табличных данных наиболее совместимы с перцептронами?
21. Как спроектировать систему featureengineering для нейросетей, работающую напрямую с реляционными данными?
22. Какие метрики мониторинга качества кластеризации следует встроить в продакшн-систему?
23. Как обеспечить масштабируемость нейросетевого решения при росте объема структурированных данных?
24. Какие инструменты профайлинга использовать для оптимизации работы алгоритмов кластеризации на больших датасетах?

ПК-6.2. Разрабатывает, отлаживает и тестирует прикладные решения с элементами ИИ с применением различных технологий хранения неструктурированных данных, оценивает качество

Типовые тестовые вопросы

1. Для кластеризации текстовых документов (неструктурированные данные) наиболее подходящим алгоритмом будет:

- a) K-means с евклидовой метрикой
- b) Иерархические алгоритмы или алгоритмы на основе косинусного сходства**
- c) Линейная регрессия
- d) DBSCAN с Манхэттенским расстоянием

Правильный ответ b)

2. При разработке системы хранения эмбедингов изображений для нейронной сети наиболее эффективно использовать:

- a) Реляционные БД
- b) Векторные базы данных (Pinecone, Weaviate) или специализированные форматы (HDF5)**
- c) JSON-файлы
- d) XML-базы данных

Правильный ответ b)

3. Для оценки качества кластеризации неструктурированных данных при отсутствии размеченных эталонов используют:

- a) Accuracy
- b) Функционалы внутренней качества (SilhouetteScore, Davies-BouldinIndex)**
- c) F1-score
- d) ROC-AUC

Правильный ответ b)

4. При обработке изображений в сверточной нейронной сети функция активации ReLu предпочтительнее сигмоиды, потому что:

- a) Она помогает бороться с затуханием градиентов в глубоких сетях**
- b) Она ограничивает выходные значения диапазоном (0,1)
- c) Она более биологически правдоподобна
- d) Она требует меньше вычислительных ресурсов

Правильный ответ a)

5. Для хранения и обработки потоковых неструктурированных данных (аудио, видео) в реальном времени используют:

- a) ApacheKafka + специализированные форматы (Avro, Parquet)**
- b) Только локальные текстовые файлы
- c) Реляционные БД с BLOB-полями
- d) Excel-таблицы

Правильный ответ a)

6. Архитектура нейронной сети для обработки неструктурированных данных изображений должна включать:

- a) Сверточные слои для автоматического извлечения признаков**
- b) Только полносвязные слои

- c) Рекуррентные слои LSTM
- d) Один перцептрон Розенблатта

Правильный ответ а)

7. При отладке кластеризации неструктурированных данных эвристическими алгоритмами важно:

- a) Тестировать различные метрики расстояния и параметры алгоритма
- b) Использовать только евклидову метрику
- c) Игнорировать шумовые точки
- d) Работать только с малыми объемами данных

Правильный ответ а)

8. Для оценки вычислительных возможностей нейронной сети при работе с видео данными учитывают:

- a) Способность обрабатывать пространственно-временные зависимости
- b) Только количество параметров модели
- c) Только объем тренировочных данных
- d) Скорость работы на CPU

Правильный ответ а)

9. При тестировании статистических алгоритмов кластеризации для текстовых данных проверяют:

- a) Сходимость ЕМ-алгоритма и устойчивость к шуму
- b) Только скорость работы
- c) Только использование памяти
- d) Совместимость с различными ОС

Правильный ответ а)

10. Для хранения моделей нейронных сетей, обученных на неструктурированных данных, используют:

- a) ONNX-формат для кроссплатформенной совместимости
- b) Только бинарные форматы фреймворков
- c) Текстовые описания архитектуры
- d) SQL-дампы весов

Правильный ответ а)

Типовые вопросы открытого типа:

1. Как выбрать алгоритм кластеризации для группировки текстовых документов и интегрировать его в систему обработки неструктурированных данных?
2. Какие метрики качества кластеризации наиболее информативны при работе с векторными представлениями изображений?
3. Как спроектировать пайплайн обработки потоковых неструктурированных данных для эвристических алгоритмов кластеризации?
4. Каким образом иерархические алгоритмы кластеризации можно применить для анализа графов социальных сетей?
5. Как оценить устойчивость статистических алгоритмов кластеризации к шуму в аудиоданных?

6. Как архитектура нейронной сети должна учитывать особенности хранения неструктурированных данных в объектных хранилищах?
7. Каким образом функции активации влияют на способность сети извлекать features из изображений и видео?
8. Как провести отладку перцептрона при работе с частично размеченными текстовыми данными?
9. Какие вычислительные возможности нейронных сетей наиболее важны для обработки потоковых неструктурированных данных?
10. Как организовать взаимодействие между алгоритмом кластеризации и распределенной файловой системой для больших объемов медиаданных?
11. Какие архитектурные решения позволяют эффективно сочетать нейросетевую обработку с NoSQL базами для JSON-документов?
12. Как спроектировать систему кэширования эмбедингов для ускорения кластеризации текстовых данных?
13. Как разработать стратегию тестирования кластеризационных алгоритмов на разнородных неструктурированных данных?
14. Какие методы оценки качества нейросетевой модели применить при работе с частично аннотированными изображениями?
15. Как провести А/В тестирование различных архитектур нейронных сетей на потоковых видео данных?
16. Как организовать предобработку разнородных неструктурированных данных для статистических алгоритмов кластеризации?
17. Какие методы векторизации текстовых данных наиболее совместимы с перцептронами различной архитектуры?
18. Как спроектировать систему извлечения признаков для нейросетей, работающую с потоковыми аудиоданными?
19. Как оптимизировать работу иерархических алгоритмов кластеризации для больших объемов видеоархивов?
20. Какие подходы к индексированию данных в векторных базах ускоряют процесс кластеризации эмбедингов?
21. Как распределить вычисления в нейросетевой модели между CPU и GPU при обработке потоковых данных?

Код компетенции	Результаты освоения ОПОП Содержание компетенций
ПК-18	Способен применять современную теоретическую математику для разработки новых алгоритмов и формулирования перспективных задач ИИ (MF-1)

ПК-18.1. Обосновывает способы и варианты применения методов и моделей в задачах искусственного интеллекта, включая их модификацию и адаптацию к специфике задачи

Типовые тестовые вопросы

1. При решении задачи классификации с несбалансированными классами обоснованным выбором функции потерь будет:

- a) MSE (MeanSquaredError)
- b) **FocalLoss или взвешенная BinaryCross-Entropy**
- c) Абсолютная ошибка (MAE)
- d) HuberLoss

Правильный ответ b)

2. Для задачи регрессии с наличием выбросов в данных наиболее обоснованным выбором функции потерь является:

- a) MSE (MeanSquaredError)
- b) **HuberLoss, которая менее чувствительна к выбросам**
- c) BinaryCross-Entropy
- d) LogCoshLoss

Правильный ответ b)

3. При адаптации метода обратного распространения ошибки для глубоких сетей с проблемой затухающих градиентов обоснованным решением будет:

- a) Увеличение learningrate
- b) **Использование skip-connections (ResNet) или LSTM-ячеек**
- c) Уменьшение количества слоев
- d) Увеличение размера батча

Правильный ответ b)

4. Для задачи с шумными данными и множеством локальных минимумов обоснованным выбором модификации градиентного спуска является:

- a) **SGD с моментом для прохождения плоских областей и локальных минимумов**
- b) Уменьшение learningrate
- c) Использование только полного градиента
- d) Увеличение размера датасета

Правильный ответ a)

5. При адаптации обучения для задачи с разреженными градиентами обоснованным выбором оптимизатора будет:

- a) **Adam, который адаптирует learningrate для каждого параметра**
- b) Vanilla SGD
- c) SGD без момента
- d) RProp

Правильный ответ a)

6. Для улучшения сходимости при обучении GAN (GenerativeAdversarialNetworks) обоснованной эвристикой является:

- a) **Использование разных функций потерь для генератора и дискриминатора**
- b) Установка одинакового learningrate для обеих сетей
- c) Одновременное обучение обеих сетей
- d) Использование только MSE loss

Правильный ответ a)

7. При адаптации метода обратного распространения для рекуррентных сетей необходимо модифицировать:

- a) **Алгоритм через BPTT (BackpropagationThroughTime)**
- b) Функцию активации на выходном слое
- c) Размер входного слоя
- d) Количество скрытых слоев

Правильный ответ a)

8.

Для задачи с большим объемом данных и ограниченными вычислительными ресурсами обоснованным выбором является:

- a) **Mini-batch gradient descent** с оптимальным размером батча
- b) Полный batch gradient descent
- c) Stochastic gradient descent (батч=1)
- d) Использование только CPU для вычислений

Правильный ответ а)

9. При адаптации обучения для transfer learning обоснованной стратегией изменения градиентного спуска является:

- a) **Использование разных learning rates для разных слоев сети**
- b) Установка одинакового learning rate для всех слоев
- c) Заморозка всех слоев предобученной сети
- d) Использование только предобученной сети без дообучения

Правильный ответ а)

10. Для улучшения сходимости при обучении сверточных сетей обоснованной эвристикой является:

- a) **Использование Batch Normalization для стабилизации распределения активаций**
- b) Увеличение количества фильтров в каждом слое
- c) Использование только линейных функций активации
- d) Уменьшение глубины сети

Правильный ответ а)

Типовые вопросы открытого типа:

1. Как обосновать выбор функции потерь для задачи классификации с сильно несбалансированными классами?
2. Какие модификации функции потерь следует рассмотреть для задачи детекции объектов на изображениях с пересекающимися bounding boxes?
3. Как адаптировать функцию потерь для задачи многозадачного обучения с разнородными типами выходных данных?
4. В каких случаях следует предпочесть Huber loss вместо MSE или MAE и как это обосновать?
5. Как модифицировать функцию потерь для учета доменных знаний и бизнес-требований в задаче прогнозирования?
6. Как обосновать применение различных вариантов обратного распространения для рекуррентных нейронных сетей?
7. Какие модификации обратного распространения следует рассмотреть для работы с разреженными градиентами?
8. Как адаптировать метод обратного распространения для сетей с недифференцируемыми операциями?
9. В каких случаях целесообразно использовать методы приближенного обратного распространения и как это обосновать?
10. Как выбрать стратегию распространения ошибки в ансамблевых архитектурах нейронных сетей?
11. Как обосновать выбор оптимизатора для задачи с шумными и нестационарными градиентами?
12. Какие эвристики улучшения сходимости следует применить для задачи с множеством локальных минимумов?

13. Как адаптировать стратегию обучения для модели с сильно различающимися по масштабу градиентами в разных слоях?
14. В каких случаях следует предпочесть адаптивные методы оптимизации классическому SGD с моментом?
15. Как модифицировать алгоритм градиентного спуска для работы с очень большими батчами?
16. Как обосновать выбор и адаптацию методов оптимизации для задачи обучения с подкреплением?
17. Какие модификации градиентного спуска следует рассмотреть для GAN-архитектур с неустойчивым обучением?
18. Как адаптировать процесс обучения для transferlearning с заморозкой части слоев?
19. В каких случаях следует использовать curriculumlearning и как обосновать его эффективность?
20. Как модифицировать функцию потерь и процесс оптимизации для semi-supervised обучения?

ПК-18.2. Применяет аппарат теории вероятностей, математической статистики и теории информации для формулирования и анализа задач искусственного интеллекта

Типовые тестовые вопросы

1. При формулировке задачи классификации в вероятностной постановке цель состоит в нахождении:

- a) Минимального расстояния до центроидов классов
- b) **Апостериорной вероятности $P(y|x)$ принадлежности объекта к классу**
- c) Евклидовой метрики между объектами
- d) Детерминированной разделяющей поверхности

Правильный ответ b)

2. Теорема Байеса применяется в задачах ИИ для:

- a) **Обновления априорных убеждений на основе новых данных**
- b) Вычисления евклидова расстояния
- c) Оптимизации функций активации
- d) Настройки learningrate

Правильный ответ a)

3. Оптимальное байесовское решающее правило минимизирует:

- a) Количество признаков
- b) **Математическое ожидание потери (risk)**
- c) Время обучения модели
- d) Объем используемой памяти

Правильный ответ b)

4. Наивный байесовский классификатор называется "наивным", потому что:

- a) Он использует простые математические операции
- b) **Предполагает условную независимость признаков при заданном классе**
- c) Он всегда дает неточные результаты
- d) Он не использует теорему Байеса

Правильный ответ b)

5. Задача восстановления плотности распределения решается методами:

- a) **Параметрического и непараметрического оценивания**
- b) Линейной алгебры
- c) Теории графов
- d) Дифференциальных уравнений

Правильный ответ a)

6. Для анализа неопределенности в прогнозах нейронной сети применяют:

- a) **Байесовские нейронные сети, которые оценивают распределение весов**
- b) Увеличение количества слоев
- c) Уменьшение размера батча
- d) Изменение функции активации

Правильный ответ a)

7. При работе с малым количеством данных наиболее обоснованным является применение:

- a) **Байесовских методов, которые позволяют инкорпорировать априорные знания**
- b) Глубоких нейронных сетей
- c) Методов, требующих больших объемов данных
- d) Случайного угадывания

Правильный ответ a)

8. Для задачи обнаружения аномалий в вероятностной постановке используют:

- a) **Оценивание плотности распределения и пороговое значение**
- b) Линейную регрессию
- c) Метод k-ближайших соседей
- d) Деревья решений

Правильный ответ a)

9. При построении байесовской сети доверия (BayesianBeliefNetwork) анализируют:

- a) **Условные вероятностные зависимости между переменными**
- b) Евклидовы расстояния между объектами
- c) Матрицы корреляции
- d) Градиенты функций потерь

Правильный ответ a)

10. Для оценки качества вероятностной модели используют:

- a) **Логарифмическую правдоподобие или кросс-энтропию**
- b) Только энтропию
- c) Евклидову метрику
- d) Косинусное расстояние

Правильный ответ a)

Типовые вопросы открытого типа:

1. Как применить теорему Байеса для обновления априорных вероятностей в системе принятия решений в условиях неопределенности?
2. Каким образом теорема Байеса позволяет комбинировать различные источники информации в задачах машинного обучения?
3. Как использовать байесовский подход для оценки неопределенности предсказаний в сложных моделях ИИ?
4. В каких практических задачах ИИ применение теоремы Байеса дает существенные преимущества перед частотным подходом?
5. Как сформулировать задачу классификации изображений в вероятностной постановке с использованием аппарата теории вероятностей?
6. Каким образом вероятностная постановка задачи позволяет учитывать стоимость различных типов ошибок классификации?
7. Как применить вероятностный подход к задаче прогнозирования временных рядов с оценкой доверительных интервалов?
8. В чем преимущества вероятностной постановки задачи перед детерминированной при работе с зашумленными данными?
9. Как вывести оптимальное байесовское решающее правило для задачи многоклассовой классификации с асимметричной матрицей потерь?
10. Каким образом оптимальное байесовское правило минимизирует математическое ожидание потерь в условиях неопределенности?
11. Как учитывать априорные знания о распределении классов при построении байесовского решающего правила?
12. В каких случаях применение оптимального байесовского правила становится computationally intractable и как это обходить?
13. Как обосновать применение наивного байесовского классификатора для задач обработки естественного языка?
14. Каким образом можно модифицировать наивный байесовский классификатор для работы с непрерывными признаками?
15. Как оценить влияние нарушения предположения о независимости признаков на качество наивного байесовского классификатора?
16. В каких прикладных задачах наивный байесовский классификатор показывает competitive performance несмотря на упрощающие предположения?
17. Как выбрать метод восстановления плотности распределения для задачи обнаружения аномалий в многомерных данных?
18. Каким образом методы восстановления плотности позволяют решать задачу генерации новых образцов данных?
19. Как оценить качество восстановленной плотности распределения с точки зрения теории информации?
20. В каких задачах непараметрические методы восстановления плотности предпочтительнее параметрических?

ПК-18.3. Применяет аппарат теории вероятностей для исследования методов и моделей машинного обучения

Типовые тестовые вопросы

1. При исследовании неопределенности предсказаний полносвязной сети (FC) применяют байесовский подход через:

- a) Замену точечных оценок весов на их апостериорные распределения
- b) Увеличение количества слоев
- c) Использование только линейных активаций
- d) Добавление сверточных слоев

Правильный ответ а)

2. Вероятностная интерпретация функции активации softmax в FC-сетях основана на:

- а) Теореме Байеса и предположении о мультиномиальном распределении**
- b) Линейной алгебре
- c) Теории графов
- d) Метриках расстояния

Правильный ответ а)

3. При исследовании устойчивости сверточных сетей (CNN) к шуму используют вероятностные методы для:

- а) Моделирования распределения adversarial examples**
- b) Увеличения количества фильтров
- c) Изменения stride свертки
- d) Настройки learning rate

Правильный ответ а)

4. Байесовская оптимизация гиперпараметров CNN применяется потому, что она:

- а) Эффективно исследует пространство параметров, используя априорные знания**
- b) Требуется полного перебора всех комбинаций
- c) Работает только с линейными моделями
- d) Игнорирует историю предыдущих экспериментов

Правильный ответ а)

5. При анализе рекуррентных сетей (RNN) вероятностная постановка задачи последовательности соответствует:

- а) Цепям Маркова и условным вероятностям $P(x_t|x_{t-1})$**
- b) Линейной регрессии
- c) Кластеризации k-means
- d) Методу главных компонент

Правильный ответ а)

6. LSTM-сети можно исследовать с позиции теории вероятностей как модели, которые:

- а) Оценивают вероятности долгосрочных зависимостей в последовательностях**
- b) Минимизируют евклидово расстояние
- c) Максимизируют корреляцию Пирсона
- d) Оптимизируют косинусное сходство

Правильный ответ а)

7. В байесовских GRU-сетях для текстовых данных анализируют:

- а) Распределение весов и их uncertainty для оценки надежности предсказаний**
- b) Только точность классификации
- c) Количество параметров модели
- d) Скорость обучения

Правильный ответ а)

8. Вероятностная постановка задачи обучения с подкреплением для RNN включает:

- a) **Policygradient методы с байесовской оценкой valuefunctions**
- b) Только Q-learning с табличным представлением
- c) Линейное программирование
- d) Метод наименьших квадратов

Правильный ответ а)

9. При исследовании CNN для сегментации изображений используют вероятностные методы для:

- a) **Моделирования пространственных зависимостей между пикселями**
- b) Увеличения разрешения изображений
- c) Настройки аугментации данных
- d) Оптимизации использования памяти

Правильный ответ а)

10. При анализе uncertainty в LSTM для медицинских данных важность байесовского подхода заключается в:

- a) **Оценке надежности предсказаний для принятия клинических решений**
- b) Увеличении скорости inference
- c) Уменьшении размера модели
- d) Упрощении архитектуры сети

Правильный ответ а)

Типовые вопросы открытого типа:

1. Как применить теорему Байеса для оценки неопределенности предсказаний полносвязной нейронной сети?
2. Каким образом байесовский подход позволяет интерпретировать веса сверточной нейронной сети как случайные величины?
3. Как использовать оптимальное байесовское решающее правило в сочетании с глубокими нейросетевыми архитектурами?
4. В чем преимущества байесовской трактовки рекуррентных нейронных сетей для задач обработки последовательностей?
5. Какова вероятностная постановка задачи обучения для сетей LSTM и GRU в рамках теории скрытых марковских моделей?
6. Каким образом можно представить сверточные нейронные сети как иерархические вероятностные модели?
7. Как задача восстановления плотности распределения связана с обучением автоэнкодеров на базе полносвязных сетей?
8. В чем заключается байесовская интерпретация регуляризации в глубоких нейронных сетях?
9. Как можно комбинировать наивный байесовский классификатор с глубокими нейросетевыми архитектурами?
10. Каким образом предположения наивного Байеса проявляются в архитектуре современных нейронных сетей?
11. Как использовать наивный байесовский подход для инициализации весов в полносвязных сетях?
12. Какие вероятностные programmingframeworks наиболее эффективны для реализации байесовских нейронных сетей?
13. Как интегрировать методы MCMC с обучением глубоких рекуррентных сетей?

14. Каким образом вариационный вывод применяется для байесовских сверточных нейронных сетей?
15. Как использовать байесовские методы для оценки epistemicuncertainty в предсказаниях CNN?
16. Каким образом теорема Байеса помогает в анализе aleatoricuncertainty в RNN?
17. Как задача восстановления плотности распределения связана с оценкой уверенности модели в предсказаниях?
18. Как применить аппарат математической статистики для сравнения различных архитектур нейронных сетей?
19. Каким образом методы теории информации позволяют анализировать информационный поток в полносвязных сетях?
20. Как использовать статистические критерии для выбора между LSTM и GRU архитектурами?

Код компетенции	Результаты освоения ОПОП Содержание компетенций
ПК-24	Способен организовывать взаимодействие с генеративными моделями через проектирование, анализ и применение промптов (LLM-5)

ПК-24.1. Использует базовые шаблоны промптов

Типовые тестовые вопросы

1. Какой тип промпта предполагает отсутствие примеров в запросе?

- a) Few-shot
- b) One-shot
- c) Zero-shot
- d) Chain-of-thought

Правильный ответ: c

2. Какой подход использует один пример решения задачи перед постановкой новой задачи?

- a) Zero-shot
- b) One-shot
- c) Retrieval
- d) Alignment

Правильный ответ: b

3. Для чего применяется few-shot prompting?

- a) Для ускорения работы API
- b) Для предоставления модели нескольких примеров выполнения задачи
- c) Для уменьшения размера модели
- d) Для хранения данных

Правильный ответ: b

4. Что является основной целью промпта?

- a) Обучение модели
- b) Формулирование инструкции модели
- c) Сжатие данных
- d) Шифрование информации

Правильный ответ: b

5. Какой шаблон лучше использовать для новой задачи без доступных примеров?

- a) Zero-shot
- b) Few-shot
- c) Fine-tuning
- d) Embedding

Правильный ответ: a

6. Какой тип промпта предполагает использование нескольких примеров решения задачи перед постановкой нового запроса?

- a) Zero-shot
- b) One-shot
- c) Few-shot
- d) Reinforcement prompting

Правильный ответ: c

7. Что является преимуществом few-shot prompting по сравнению с zero-shot?

- a) Уменьшение размера модели
- b) Улучшение понимания формата ожидаемого ответа
- c) Снижение затрат на API
- d) Изменение архитектуры модели

Правильный ответ: b

8. Какой элемент наиболее важен при формировании простого промпта?

- a) Объем оперативной памяти
- b) Четкая постановка задачи
- c) Размер датасета
- d) Количество эпох обучения

Правильный ответ: b

9. Для какой задачи наиболее подходит one-shot prompting?

- a) Когда необходимо показать модели один образец правильного ответа
- b) Когда модель дообучается на новых данных
- c) Когда требуется уменьшить размер модели
- d) Когда используется API

Правильный ответ: a

10. Что следует сделать при получении некорректного результата от базового промпта?

- a) Изменить аппаратное обеспечение
- b) Добавить уточняющие инструкции или примеры
- c) Переобучить модель
- d) Удалить часть запроса

Правильный ответ: b

Типовые вопросы открытого типа:

1. Чем отличаются zero-shot, one-shot и few-shot промпты?
2. В каких случаях целесообразно использовать few-shot prompting?
3. Какие ошибки допускаются при формулировании простых промптов?
4. Как влияет полнота инструкции на качество ответа модели?
5. Какие преимущества дает использование шаблонов промптов?
6. Как влияет качество формулировки запроса на результат работы модели?
7. Какие критерии позволяют оценить корректность выбранного шаблона промпта?
8. Как использовать примеры для повышения качества генерации?

9. Какие ограничения имеют базовые шаблоны промптов?
10. Как организовать библиотеку типовых промптов для прикладного проекта?

ПК-24.2. Встраивает промпты в пайплайн взаимодействия

Типовые тестовые вопросы

1. Что понимается под пайплайном взаимодействия с LLM?

- a) Последовательность этапов обработки данных и взаимодействия с моделью
- b) Архитектура нейронной сети
- c) Метод обучения модели
- d) Формат хранения данных

Ответ: а

2. Какой этап обычно выполняется перед отправкой запроса модели?

- a) Предобработка данных
- b) Развертывание API
- c) Кластеризация
- d) Индексация БД

Ответ: а

3. Какой этап выполняется после получения ответа модели?

- a) Постобработка результатов
- b) Дообучение модели
- c) Регрессия
- d) Векторизация

Ответ: а

4. Что позволяет реализовать многошаговый пайплайн?

- a) Последовательное уточнение результата
- b) Уменьшение размера модели
- c) Изменение архитектуры LLM
- d) Исключение ошибок

Ответ: а

5. Какой компонент отвечает за передачу данных между этапами пайплайна?

- a) Оркестратор
- b) GPU
- c) Файловая система
- d) Компилятор

Ответ: а

6. Какой компонент пайплайна отвечает за подготовку данных перед отправкой запроса в LLM?

- a) Предобработка данных
- b) Токенизация модели
- c) Fine-tuning
- d) Кэширование

Правильный ответ: а

7. Для чего используется многошаговое взаимодействие с моделью?

- a) Для уменьшения объема памяти
- b) Для последовательного уточнения результата
- c) Для изменения архитектуры модели

d) Для ускорения работы сети

Правильный ответ: b

8. Какой этап выполняется после получения ответа модели?

a) Валидация и постобработка

b) Обучение модели

c) Индексация базы данных

d) Разметка данных

Правильный ответ: a

9. Что обеспечивает устойчивость пайплайна к ошибкам модели?

a) Контроль качества и обработка исключений

b) Увеличение числа параметров модели

c) Использование более длинных промптов

d) Сжатие данных

Правильный ответ: a

10. Для чего применяется оркестрация запросов?

a) Для управления последовательностью взаимодействий

b) Для обучения модели

c) Для хранения данных

d) Для визуализации результатов

Правильный ответ: a

Типовые вопросы открытого типа:

1. Что входит в типовой пайплайн взаимодействия с LLM?
2. Как организовать цепочку запросов к языковой модели?
3. Какие этапы предобработки данных наиболее востребованы?
4. Для чего используется постобработка ответов модели?
5. Как контролировать качество работы пайплайна?
6. Как организовать последовательность запросов к LLM для решения сложной задачи?
7. Какие способы передачи промежуточных результатов используются в пайплайнах?
8. Как обеспечить надежность работы пайплайна при ошибках модели?
9. Какие преимущества дают многошаговые сценарии взаимодействия?
10. Как оценить эффективность построенного пайплайна?

ПК-24.3. Настраивает API для работы с БЯМ

Типовые тестовые вопросы

1. Какой протокол чаще всего используется для доступа к API LLM?

a) HTTP/HTTPS

b) FTP

c) SMTP

d) SNMP

Ответ: a

2. Для чего используется API-ключ?

a) Для авторизации клиента

b) Для обучения модели

c) Для хранения данных

d) Для шифрования БД

Ответ: а

3. Что обычно передается в теле запроса к LLM API?

- a) Промпт и параметры генерации
- b) Архитектура модели
- c) Исходный код API
- d) Веса модели

Ответ: а

4. Какой параметр влияет на степень случайности генерации?

- a) Temperature
- b) Epochs
- c) Batch size
- d) Learning rate

Ответ: а

5. Почему нельзя публиковать API-ключ в открытом репозитории?

- a) Возможен несанкционированный доступ
- b) Снижается точность модели
- c) Увеличивается время ответа
- d) Изменяется архитектура сети

Ответ: а

6. Какой параметр API ограничивает длину ответа модели?

- a) temperature
- b) max_tokens
- c) embedding_size
- d) context_window

Правильный ответ: b

7. Для чего используется HTTPS при работе с API?

- a) Для защиты передаваемых данных
- b) Для увеличения скорости генерации
- c) Для обучения модели
- d) Для уменьшения размера ответа

Правильный ответ: а

8. Что означает код ответа HTTP 401?

- a) Ошибка авторизации
- b) Ошибка генерации
- c) Сервер недоступен
- d) Превышение лимита токенов

Правильный ответ: а

9. Что следует делать при превышении лимитов API?

- a) Реализовать повторные запросы с задержкой
- b) Игнорировать ошибку
- c) Изменить модель вручную
- d) Удалить API-ключ

Правильный ответ: а

20. Где рекомендуется хранить API-ключ?

- a) В переменных окружения

- b) В исходном коде
- c) В публичном репозитории
- d) В комментариях программы

Правильный ответ: а

Типовые вопросы открытого типа:

1. Как организовать подключение к API языковой модели?
2. Какие параметры генерации используются чаще всего?
3. Какие меры безопасности необходимо соблюдать при работе с API?
4. Как обрабатывать ошибки API-запросов?
5. Какие преимущества дает использование облачных LLM API?
6. Какие меры безопасности необходимо применять при работе с API?
7. Как организовать хранение ключей доступа?
8. Какие ошибки наиболее часто возникают при интеграции API?
9. Как реализовать журналирование запросов к API?
10. Как обеспечить отказоустойчивость взаимодействия с сервисом LLM?

ПК-24.4. Разрабатывает дизайн и структуру промптов

Типовые тестовые вопросы

1. Что понимается под ролью (role) в промпте?

- a) Контекст поведения модели
- b) Архитектура сети
- c) Тип базы данных
- d) Формат ответа API

Ответ: а

2. Для чего используется контекст в промпте?

- a) Для повышения релевантности ответа
- b) Для обучения модели
- c) Для сжатия данных
- d) Для ускорения API

Ответ: а

3. Что позволяет ограничить формат ответа модели?

- a) Явное указание структуры ответа
- b) Увеличение температуры
- c) Добавление эмбеддингов
- d) Fine-tuning

Ответ: а

4. Какой элемент повышает воспроизводимость ответа?

- a) Четкие инструкции
- b) Большая температура
- c) Случайный контекст
- d) Удаление ограничений

Ответ: а

5. Что обычно улучшает качество сложного промпта?

- a) Разделение задачи на этапы
- b) Уменьшение длины запроса до одного предложения
- c) Отсутствие контекста
- d) Исключение ограничений

Ответ: а

6. Что позволяет использование роли в промпте?

- a) Задать модель поведения генерации
- b) Ускорить работу API
- c) Снизить объем данных
- d) Изменить архитектуру модели

Правильный ответ: а

7. Для чего используются ограничения в промпте?

- a) Для управления содержанием ответа
- b) Для обучения модели
- c) Для сокращения контекста
- d) Для кэширования данных

Правильный ответ: а

8. Что повышает воспроизводимость ответа?

- a) Четкие инструкции и формат вывода
- b) Увеличение температуры
- c) Удаление контекста
- d) Случайные примеры

Правильный ответ: а

9. Что такое структурированный промпт?

- a) Промпт с явно выделенными инструкциями и контекстом
- b) Промпт, содержащий код программы
- c) Промпт для обучения модели
- d) Любой длинный запрос

Правильный ответ: а

10. Какой элемент помогает получить ответ в формате JSON?

- a) Явное указание структуры ответа
- b) Увеличение количества токенов
- c) Изменение модели
- d) Использование API-ключа

Правильный ответ: а

Типовые вопросы открытого типа:

1. Какие элементы включает хорошо структурированный промпт?
2. Как роль влияет на ответы модели?
3. Для чего задаются ограничения в промпте?
4. Как повысить качество ответов через изменение структуры промпта?
5. Какие принципы проектирования сложных промптов существуют?
6. Как использовать роли для повышения качества ответов?
7. Какие ограничения следует задавать в промптах?
8. Как проектировать промпты для структурированного вывода данных?
9. Какие элементы помогают повысить воспроизводимость ответов?

10. Как оценить качество разработанного промпта?

ПК-24.5. Использует методы контроля и выравнивания

Типовые тестовые вопросы

1. Что понимается под галлюцинацией LLM?

- a) Генерация недостоверной информации
- b) Ошибка API
- c) Отказ модели отвечать
- d) Переобучение модели

Ответ: а

2. Для чего применяется фактчекинг результатов?

- a) Для проверки достоверности ответа
- b) Для ускорения генерации
- c) Для уменьшения памяти
- d) Для обучения модели

Ответ: а

3. Что означает термин alignment?

- a) Согласование поведения модели с требованиями безопасности и полезности
- b) Выравнивание данных в таблице
- c) Сжатие параметров модели
- d) Кластеризация ответов

Ответ: а

4. Какой способ помогает снизить число галлюцинаций?

- a) Добавление проверяемого контекста
- b) Увеличение температуры
- c) Соккрытие инструкций
- d) Удаление ограничений

Ответ: а

5. Что следует делать с ответом модели в критически важных системах?

- a) Выполнять дополнительную проверку
- b) Использовать без проверки
- c) Игнорировать ограничения
- d) Всегда доверять модели

Ответ: а

6. Что является одним из способов снижения галлюцинаций?

- a) Предоставление достоверного контекста
- b) Увеличение температуры
- c) Сокращение запроса
- d) Исключение ограничений

Правильный ответ: а

7. Что понимается под фактчекингом?

- a) Проверка достоверности информации
- b) Обучение модели
- c) Настройка API
- d) Кластеризация данных

Правильный ответ: а

8. Для чего используются guardrails?

- a) Для ограничения нежелательного поведения модели
- b) Для ускорения генерации
- c) Для обучения модели
- d) Для хранения данных

Правильный ответ: а

9. Что означает alignment генеративной модели?

- a) Согласование поведения модели с требованиями безопасности и полезности
- b) Выравнивание данных в таблице
- c) Оптимизация скорости работы
- d) Уменьшение количества параметров

Правильный ответ: а

10. Почему критически важные ответы нельзя использовать без проверки?

- a) Возможны ошибки и недостоверные сведения
- b) Модель работает слишком быстро
- c) API ограничивает количество запросов
- d) Ответы слишком короткие

Правильный ответ: а

Типовые вопросы открытого типа:

1. Что такое галлюцинации и как их обнаружить?
2. Какие методы контроля качества ответов существуют?
3. Как организовать проверку достоверности ответа модели?
4. Что понимается под alignment генеративных моделей?
5. Какие риски возникают при использовании LLM без контроля качества?
6. Выполнить фактчекинг ответа модели по заданной теме.
7. Разработать набор правил контроля генерации.
8. Настроить ограничения для безопасной генерации.
9. Проанализировать риски использования LLM в прикладной задаче.
10. Разработать процедуру контроля качества ответов модели.

ПК-24.6. Анализирует и отлаживает промпты

Типовые тестовые вопросы

1. Что является целью отладки промпта?

- a) Повышение качества результата
- b) Уменьшение размера модели
- c) Изменение архитектуры сети
- d) Снижение числа параметров

Ответ: а

2. Какой подход чаще используется при улучшении промпта?

- a) Итеративное уточнение
- b) Полное удаление контекста
- c) Замена модели
- d) Удаление примеров

Ответ: а

3. Что следует анализировать при неудовлетворительном ответе?

- a) Формулировку инструкции
- b) Частоту процессора
- c) Размер монитора
- d) Версию ОС

Ответ: a

4. Что позволяет A/B-тестирование промптов?

- a) Сравнить эффективность разных вариантов
- b) Ускорить API
- c) Изменить модель
- d) Обучить модель

Ответ: a

5. Какая метрика может использоваться при оценке промпта?

- a) Полезность и точность ответа
- b) Размер файла
- c) Частота сети
- d) Количество потоков CPU

Ответ: a

6. Что позволяет A/B-тестирование промптов?

- a) Сравнить эффективность различных вариантов
- b) Уменьшить размер модели
- c) Изменить архитектуру модели
- d) Настроить API

Правильный ответ: a

7. Что является признаком неудачного промпта?

- a) Нестабильные и нерелевантные ответы
- b) Высокая скорость генерации
- c) Короткий текст ответа
- d) Использование API

Правильный ответ: a

8. Для чего проводится анализ ошибок генерации?

- a) Для улучшения качества промпта
- b) Для изменения модели
- c) Для хранения данных
- d) Для уменьшения объема контекста

Правильный ответ: a

9. Что такое итеративное улучшение промпта?

- a) Последовательная корректировка на основе результатов тестирования
- b) Дообучение модели
- c) Сжатие данных
- d) Удаление примеров

Правильный ответ: a

10. Какой результат свидетельствует об успешной отладке?

- a) Повышение качества и стабильности ответов
- b) Увеличение длины запроса
- c) Рост количества токенов

d) Уменьшение числа пользователей

Правильный ответ: а

Типовые вопросы открытого типа:

1. Какие этапы включает процесс отладки промпта?
2. Как выявить причину некорректного ответа модели?
3. Какие методы оценки эффективности промптов существуют?
4. Как организовать А/В-тестирование промптов?
5. Какие типовые ошибки встречаются при проектировании запросов к LLM?
6. Какие способы сравнения различных версий промптов существуют?
7. Как организовать процесс итеративного улучшения промптов?
8. Какие метрики использовать для оценки качества промптов?
9. Как документировать результаты тестирования?
10. Какие типовые причины приводят к некорректной генерации?

Типовые теоретические вопросы для экзамена по дисциплине

1. Объясните, как выбор метрики расстояния может внести скрытую предвзятость в модель на основе метода ближайших соседей.
2. Как методы кросс-валидации помогают выявить переобучение и оценить способность модели к обобщению на новых данных?
3. Опишите процесс и цели нормализации данных. К каким последствиям может привести ее отсутствие при использовании метрических алгоритмов?
4. В чем заключаются компромиссы при выборе значения k в методе k -NN и как это связано с понятием смещения-разложения (bias-variance tradeoff)?
5. Сравните эвристические и статистические алгоритмы кластеризации. В каких практических сценариях предпочтительнее каждый из подходов?
6. Как можно модифицировать метод ближайших соседей для работы с категориальными признаками и данными разного масштаба?
7. Объясните, как визуализация результатов кластеризации (например, с помощью дендрограмм или t-SNE) помогает в интерпретации и отладке моделей.
8. Дайте вероятностную интерпретацию логистической регрессии и объясните связь между функцией потерь и правдоподобием.
9. Что такое принцип линейной разделимости и как он влияет на выбор архитектуры модели при проектировании системы классификации?
10. Опишите геометрическую интерпретацию разделяющей гиперплоскости и как ее положение определяется в процессе оптимизации.
11. В чем разница между эмпирическим риском и ожидаемым риском, и почему минимизация первого не всегда гарантирует успех на новых данных?
12. Сравните функции потерь для задач регрессии (MSE, MAE, Huber). В каких ситуациях каждая из них является наиболее подходящей?
13. Как метод стохастического градиентного спуска помогает бороться с локальными минимумами и ускоряет обучение на больших наборах данных?
14. Объясните, как такие эвристики, как момент или адаптивная скорость обучения, улучшают сходимость градиентного спуска в глубоких сетях.
15. Сформулируйте задачу классификации в вероятностной постановке и объясните, как оптимальное байесовское решающее правило минимизирует ожидаемые потери.
16. В чем заключаются основные предположения наивного байесовского классификатора и как их нарушение влияет на качество модели?
17. Опишите задачу восстановления плотности распределения. Какие методы вы знаете и в каких прикладных задачах они используются?

18. Как с помощью теоремы Байеса можно обновлять уверенность модели в своих прогнозах по мере поступления новых данных?
19. Объясните, как байесовские методы позволяют оценивать неопределенность предсказаний и почему это важно для систем принятия решений.
20. В чем разница между байесовским и частотным подходами к машинному обучению? Приведите примеры задач, где байесовский подход предпочтительнее.
21. Проведите аналогию между биологическим нейроном и моделью искусственного нейрона. Какие компоненты являются наиболее важными для вычислительных возможностей?
22. Опишите архитектуру и принцип работы перцептрона Розенблатта. В чем его основные ограничения и как они были преодолены в современных нейросетях?
23. Объясните, как выбор функции активации влияет на способность сети обучаться и аппроксимировать сложные нелинейные функции.
24. В чем заключаются вычислительные преимущества и ограничения полносвязных сетей по сравнению с другими архитектурами?
25. Опишите индуктивное смещение, вносимое сверточными слоями, и почему оно так эффективно для задач компьютерного зрения.
26. Объясните принцип работы рекуррентных нейронных сетей и в чем их ключевое отличие от сетей прямого распространения.
27. Как архитектуры LSTM и GRU решают проблему исчезающего градиента и позволяют捕获 долгосрочные зависимости в последовательностях?
28. Опишите вычислительные возможности нейронных сетей с точки зрения теоремы универсальной аппроксимации. Каковы ее ограничения на практике?
29. Как метод обратного распространения ошибки использует цепное правило для эффективного вычисления градиентов в глубоких сетях?
30. Объясните, как такие техники, как Dropout и BatchNormalization, помогают бороться с переобучением и ускоряют сходимость обучения.
31. Сравните фреймворки для глубокого обучения (например, TensorFlow и PyTorch) с точки зрения гибкости, производительности и легкости отладки.
32. Как можно использовать перенос обучения (transfer learning) для адаптации предобученных моделей к специфическим задачам с ограниченными данными?
33. Опишите процесс и инструменты для отладки и мониторинга процесса обучения глубокой нейронной сети.
34. Как можно обнаружить и mitigate когнитивные искажения и культурно-обусловленные предвзятости в тренировочных данных?
35. Объясните, как ценностные предпосылки разработчика могут повлиять на постановку задачи, выбор метрик и, в конечном итоге, на поведение ИИ-системы.
36. Какие методологии и инструменты позволяют оценивать fairness (справедливость) и интерпретируемость моделей машинного обучения?
37. Опишите стратегии работы с несбалансированными данными на разных этапах разработки ML-модели.
38. Как подходы к обработке неструктурированных данных (текст, изображения, аудио) влияют на выбор архитектуры модели и методы ее обучения?
39. Обсудите проблему "черного ящика" в глубоком обучении и современные методы интерпретации решений нейронных сетей.
40. Какова роль различных технологий хранения данных (реляционные, векторные, NoSQL БД) в жизненном цикле ML-решения?
41. Опишите процесс проектирования и разработки сквозного ML-пайплайна — от сбора данных до развертывания и мониторинга модели.
42. Какие современные инструментальные средства (MLOps) используются для обеспечения воспроизводимости экспериментов и управления версиями моделей?
43. Как методы теории информации (такие как взаимная информация) могут быть использованы для отбора признаков и анализа моделей?

44. Обсудите компромиссы при выборе между сложностью модели, ее интерпретируемостью и вычислительными затратами в прикладных задачах.
45. Каковы современные тенденции в разработке гибридных моделей, сочетающих символьные подходы ИИ с методами глубокого обучения?
46. Как выбирать шаблон под задачу?
47. Как повысить качество ответа без изменения модели?
48. Какие ошибки возникают при составлении промптов?
49. Как влияет контекст на качество ответа?
50. Какие метрики использовать для оценки пайплайна?
51. Как интегрировать LLM в существующую информационную систему?
52. Как организовать безопасное хранение ключей?
53. Какие ограничения API следует учитывать?
54. Какие преимущества дают облачные API?
55. Как проектировать роль модели?
56. Как задавать формат ответа?
57. Как улучшать структуру промпта?
58. Как организовать многошаговый промпт?
59. Как выявлять галлюцинации?
60. Как организовать фактчекинг?
61. Что понимается под alignment?
62. Какие риски возникают при использовании LLM?
63. Опишите процесс отладки промпта.
64. Как проводить A/B-тестирование?
65. Как выявлять неоднозначные инструкции?
66. Как сравнивать версии промптов?