

УДК 004.89

ПРИМЕНЕНИЕ МЕТОДОВ КЛАСТЕРИЗАЦИИ ДЛЯ ВЫЯВЛЕНИЯ НЕОДНОРОДНОЙ ПРОИЗВОДИТЕЛЬНОСТИ МОДЕЛЕЙ

Акулова Е.О., студент гр. ИТб-211, IV курс
Научный руководитель: Матисов А. В., старший преподаватель
Кузбасский государственный технический университет
имени Т.Ф. Горбачева, г. Кемерово

Применение методов кластеризации для выявления неоднородной производительности моделей является важной и актуальной задачей в области машинного обучения, особенно в контексте систем с многозадачностью и разными типами данных [1].

Неоднородная производительность может проявляться в том, что модель работает с высокой точностью для одних типов данных, но ошибается или демонстрирует низкую эффективность на других. Это явление может быть связано с несколькими факторами, такими как структурные особенности данных, особенности распределений признаков или даже скрытые зависимости, которые не учитываются в процессе обучения модели.

Методы кластеризации, в свою очередь, представляют собой мощный инструмент для анализа этих различий, так как они позволяют группировать данные в такие кластеры, которые имеют схожие характеристики, и затем исследовать производительность модели в этих различных группах.

Одним из ключевых этапов этого процесса является выделение подмножеств данных с похожими характеристиками и анализ производительности модели на каждом из этих подмножеств. В реальных приложениях данные часто имеют высокую сложность и многогранность, что делает традиционные методы анализа производительности модели менее эффективными. В таких случаях кластеризация может помочь разделить данные на группы, которые обладают определённой однородностью, и позволяет понять, на каких группах модель работает плохо, а на каких – хорошо. Это может выявить скрытые зависимости, такие как определённые комбинации признаков или значения, которые модель неправильно интерпретирует, но которые остаются игнорируемыми в стандартных подходах [2].

Для этого могут быть использованы такие методы кластеризации, как K-средние, иерархическая кластеризация или более сложные алгоритмы с применением плотностной кластеризации, такие как DBSCAN. Эти алгоритмы позволяют выявить подмножества данных, которые имеют схожие характеристики

и, возможно, сходную реакцию на модель. Например, модель может демонстрировать высокую точность на определённой группе данных, где характеристики признаков более сбалансированы или представляют собой несложные закономерности. В то же время она может ошибаться на других группах, где признаки имеют более сложные зависимости или распределены неравномерно. Применяя методы кластеризации, можно выделить такие группы и проанализировать причины, почему модель работает хуже на этих подмножествах.

Кластеризация также может помочь в выявлении случаев, когда модель имеет «непостоянную» производительность. Например, в некоторых случаях алгоритм может давать хорошие результаты для данных, близких по структуре, но плохие для других, которые имеют редкие или аномальные значения. Используя методы кластеризации, можно исследовать, почему модель плохо работает в этих случаях, а также найти специфические особенности таких данных, которые требуют дополнительной обработки, исправлений или улучшений в модели.

Помимо этого, методы кластеризации могут быть полезны для обнаружения системных ошибок в обучении модели. Например, если модель ошибается на всех данных одного кластера, это может указывать на наличие ошибки в процессе обучения или на недостаточную представительность этого типа данных в обучающем наборе. Анализ с использованием кластеризации может выявить, что модель неадекватно обрабатывает или не может обобщить какие-то типы данных, что приводит к ухудшению её общей производительности. В таких случаях можно предпринять шаги для улучшения обучения, такие как сбор дополнительных данных для недопредставленных кластеров, использование специальных техник аугментации данных или настройку гиперпараметров модели [3].

Для дальнейшего анализа неоднородности производительности модели можно применить методы мета-обучения, которые позволяют обучать модель на базе информации о производительности на разных кластерах данных. Это может помочь улучшить устойчивость модели и её способность обрабатывать более разнообразные данные.

Например, можно разработать схему, в которой модель сначала обучается для каждого кластера отдельно, а затем мета-обучающая модель принимает решения о том, как комбинировать эти результаты для достижения наилучшей общей производительности.

Кроме того, кластеризация может помочь выявить зоны неопределённости в данных, где модель проявляет нестабильность. Если в какой-то группе данных модель систематически ошибается или выдает результаты с низким уровнем уверенности, это может быть показателем того, что эта группа данных нуждается в более глубоком анализе или в использовании специфических

методов для повышения точности предсказаний. Например, если кластеризация показала, что одна группа данных имеет широкий диапазон значений для определённых признаков, и модель не способна адекватно обработать эти данные, можно применить техники, такие как нормализация или стандартизация признаков, чтобы уменьшить влияние этих вариаций.

Методы кластеризации также играют важную роль в оценке метрик производительности модели в разных группах данных. Важно не только отслеживать общую точность модели, но и понимать, как она работает для разных кластеров. Например, можно вычислять отдельные метрики для каждого кластера, такие как точность, полнота, F1-меру и другие, чтобы понять, в каких группах модель работает хорошо, а в каких – необходимо произвести доработки. Такие изолированные метрики могут дать более детальное представление о производительности модели, чем простая сводная метрика, и помочь выделить те области, которые требуют улучшений [4].

Сравнение производительности модели в разных кластерах также может быть полезно для обнаружения скрытых паттернов в данных, которые раньше могли быть не очевидны. Например, кластеризация может показать, что модель лучше работает с одними типами признаков (например, с текстовыми данными или с изображениями) и хуже – с другими (например, с числовыми значениями). Это может дать понимание, какие типы признаков необходимо улучшить или как адаптировать модель для лучшей работы с данными в определённых группах.

В заключение, использование методов кластеризации для выявления неоднородной производительности моделей предоставляет мощный инструмент для анализа и оптимизации алгоритмических систем. Эти методы помогают разделить данные на группы с похожими характеристиками и анализировать, как модель работает с различными типами данных. Это позволяет выявить скрытые закономерности и ошибки в обучении, улучшить точность предсказаний и сделать модели более устойчивыми и справедливыми.

Кластеризация также способствует улучшению общей интерпретируемости модели, что важно для принятия более обоснованных и прозрачных решений, особенно в таких областях, как медицина, финансы и право, где ошибки могут иметь серьёзные последствия.

Список литературы:

1. D'Amour, A. Underspecification Presents Challenges for Credibility in Modern Machine Learning / A. D'Amour, K. Heller, D. Moldovan [et al.] // Journal of Machine Learning Research. – 2022. – Vol. 23. – P. 1–61.
2. Lakkaraju, H. Interpretable Decision Sets: A Joint Framework for Description and Prediction / H. Lakkaraju, S. H. Bach, J. Leskovec // Proceedings of the

22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining. – New York : ACM, 2016. – P. 1675–1684. – DOI 10.1145/2939672.2939874.

3. Kriegel, H.-P. LOF: Identifying Density-Based Local Outliers / H.-P. Kriegel, M. M. Breunig, R. T. Ng, J. Sander // Proceedings of the ACM SIGMOD International Conference on Management of Data. – New York : ACM, 2000. – P. 93–104. – DOI 10.1145/342009.335388.

4. Zhao, Q. On the Explanatory Power of Clustering for Model Error Analysis / Q. Zhao, D. Adeli, A. H. Gee [et al.] // arXiv preprint. – 2021. – arXiv:2107.01235.