

Research and Analysis of Dataset and Model Evaluation Limitations and Optimization Solutions Based on the Anchor-Point Evaluation Method

Jinjin Han

*College of Arts and Information Engineering, Dalian Polytechnic University, Dalian, China
hanjinjinresearch@outlook.com*

Abstract. From the perspective of the current development trend of large language models, full evaluation on complete test sets entails high computational, time, and financial costs and cannot meet the need for frequent and rapid testing. Anchor-point evaluation relies on a small number of representative anchor samples to approximately assess the overall performance of a model, making it a feasible solution for low-cost and efficient evaluation. However, this method depends heavily on the quality of the anchor samples and the design of the evaluation framework, and it has evident shortcomings in terms of data distribution shift, sample noise, experimental conditions, and error control. This paper summarizes various problems in anchor-point evaluation from the two levels of datasets and model evaluation, compares optimization approaches such as active sampling, domain adaptation, data denoising, weighted evaluation, and multidimensional robust evaluation, and discusses the scope of application and constraints of general optimization solutions. Active sampling and lightweight data correction can improve the stability and accuracy of anchor-sample evaluation; in contrast, complex domain adaptation, multilevel adversarial experiments, and large-scale error analysis have excessively high computational costs, weakening the core lightweight advantage of anchor-point evaluation. Future research needs to balance evaluation accuracy, computational efficiency, and actual business needs, with an emphasis on developing lightweight and reliable anchor-point evaluation solutions suited to few-sample scenarios.

Keywords: large language models, anchor-point evaluation, active sampling, domain adaptation, robust evaluation

1. Introduction

The number of parameters in currently used large language models is becoming increasingly large, and the computing resources required for training and deployment are growing accordingly. Under traditional methods, full-set testing generally involves submitting all test data to the model for processing and then measuring the model's performance in various respects by accuracy, generation quality, or recall capability. In terms of the completeness of the results, this type of evaluation is relatively comprehensive, but it entails substantial computational consumption and a long evaluation

time and therefore cannot effectively meet the practical needs arising from continuous model improvement, rapid version iteration, or real-time business evaluation of large numbers of instances.

To address the pain point of the excessive cost of full-set evaluation, lightweight few-sample evaluation has become a mainstream research direction, and few-sample evaluation solutions such as anchor-point evaluation and tinyBenchmarks have been proposed successively [1, 2]. Reviews in related fields have also indicated that traditional evaluation systems have evident shortcomings [3], while few-shot learning theory provides support for optimizing anchor effectiveness [4].

What is now referred to as anchor-point evaluation has begun to attract attention. It uses a small number of representative samples from a complete test set as anchors. Anchors are small-scale characteristic samples used to rapidly infer a model's overall performance. Researchers rely on the evaluation results of anchor samples to approximately infer the model's comprehensive performance on all test data. The analysis of anchor samples is treated as a clue, and its results are used to roughly determine how the model processes the complete test set. Compared with full-set evaluation, anchor-point evaluation uses fewer test instances and can therefore save the time and computing resources expended on evaluation. However, the effectiveness involved in anchor-point evaluation cannot be judged solely by the number of samples; the quality and distribution of the anchor samples and the evaluation method must also be considered together. With respect to the anchor samples themselves, they are often designed for a particular task or situation and may therefore be affected by data distribution shift. If old anchor samples are used to analyze a new task or new business scenario, the relevant characteristics may not be described completely, and the evaluation results may consequently be biased. At the same time, because the number of anchor samples is generally small, individual atypical, interfering, or marginal instances contribute heavily to the overall result, and the resulting evaluation statistics, such as variance, can easily become very large.

In addition to the discussion concerning datasets, the model evaluation method on which anchor-point evaluation relies remains incomplete. Much current work focuses only on the positive performance of methods under ideal conditions, while giving insufficient consideration to multidomain applications, special cases, and failure situations. It also fails to clearly distinguish whether the model itself is good or whether the anchor-point estimate is erroneous, and it lacks a unified method for measuring errors. These shortcomings constitute a fundamental obstacle that makes the conclusions of anchor-point evaluation weak in explanatory power or unreliable.

Discussions of active sampling, domain adaptation, distribution-shift correction, and robustness evaluation are currently continuing to develop, and the proposed ideas offer new inspiration for anchor-point evaluation [5-7]. However, it should be acknowledged that some conventional optimization methods are computationally burdensome; if they are applied to anchor-point evaluation without restraint, they will undermine its original lightweight purpose.

The fundamental issues to be discussed in this paper are to examine the inherent defects of anchor-point evaluation from the perspective of datasets and analyze whether known data optimization methods can be expected to resolve such defects, and, with respect to model evaluation, to determine what mechanistic difficulties anchor-point evaluation has and how traditional robust evaluation can be applied in coordination with few-sample evaluation. After summarizing and comparing various existing studies, this paper attempts to propose ideas that can serve as references, namely, establishing a rapid evaluation system for large models in a lower-cost, more stable, and more reliable manner.

2. Dataset limitations and optimization solutions in anchor-point evaluation tasks

2.1. Main problems at the dataset level

Most current discussions of anchor-point evaluation methods concentrate on improving efficiency or reducing evaluation costs and give little consideration to the problems involved in data distribution shift, samples, and cross-domain transfer. Some studies treat test data as high-quality objects and assume that anchor samples can fully describe the overall distribution. In reality, however, data may be unclean, classes may not be balanced, and the domains to which the data belong often differ. Therefore, experiments conducted under idealized assumptions may not necessarily yield reliable conclusions applicable to actual business scenarios.

The accuracy of anchor-point evaluation is highly dependent on the representativeness of the anchor samples. Compared with full-set testing, a small number of anchor samples can reduce evaluation costs, but it can also amplify the effects of uneven data distributions and sample-selection bias. At present, the following three main problems exist.

2.1.1. Limited cross-domain generalization capability

Anchor samples are generally selected according to a fixed task or fixed domain. If a model is to be applied to actual situations in different business scenarios, the data encountered may differ substantially from the original data in task type, language characteristics, relevant knowledge, and difficulty. In such cases, the samples previously used as anchors may not be appropriate, and the resulting anchor-point evaluation may not correctly describe the model's actual processing capability in the target situation.

2.1.2. Significant influence of sampling noise

Because the number of anchor samples is small, individual anomalous samples or samples with low-quality annotations may exert a large influence on the overall evaluation results. If sample selection mainly relies on random sampling, substantial differences in results may occur among different sets of anchor samples, thereby reducing the stability and reproducibility of the evaluation results.

2.1.3. Relatively limited experimental validation scenarios

Some studies of anchor-point evaluation are validated mainly on a single dataset or a single task and lack systematic testing across tasks, across domains, and in out-of-distribution scenarios. Such experimental designs can demonstrate the effectiveness of a method under specific conditions, but they are insufficient to comprehensively prove its general applicability in complex real-world environments.

2.2. Dataset optimization methods

To address the above problems, existing research has mainly developed the following categories of optimization approaches.

2.2.1. Active sampling and representative sample selection

Active sampling does not treat randomness as the sole criterion; instead, it uses sample uncertainty, confidence level, within-class distance, or gradients as clues to deliberately select representative samples with high information content [5, 8]. Compared with random sampling, this method can allocate limited evaluation quotas more reasonably. Sample-stratification methods such as dataset cartography can also assist in identifying easy-to-learn, ambiguous, and hard-to-learn samples, thereby improving the coverage of the overall distribution by the anchors [9].

2.2.2. Domain adaptation and distribution-shift correction

Domain adaptation usually reduces differences between the source domain and the target domain through feature alignment, distribution correction, or statistical normalization. When the source and target domains mainly exhibit covariate shift, importance weighting can also be used to estimate risk under the target distribution [10]. From the perspective of cross-domain anchor-point evaluation, the above methods help improve the applicability of existing anchors in target scenarios, but they increase the costs of data organization, weight estimation, or model computation.

2.2.3. Data denoising and weighted evaluation

To handle anomalous samples or low-quality annotations, anomaly detection, sample screening, or weight adjustment can be used to reduce their influence. Confident learning can identify suspected mislabeled samples by estimating the joint distribution between labels and latent true classes [11]. In anchor-point evaluation, the weights of low-confidence samples can be reduced and those of more representative samples can be increased to prevent individual erroneous samples from excessively influencing the overall conclusion.

2.2.4. Multidomain mixed benchmark datasets

By integrating data from different domains, tasks, and levels of difficulty, more diverse evaluation benchmarks can be constructed. MMLU and C-Eval provide multidisciplinary knowledge and reasoning tasks in English and Chinese, respectively [12, 13], while HELM and BIG-bench further emphasize comprehensive evaluation across multiple scenarios, metrics, and capability dimensions [7-14]. Such mixed benchmarks can alleviate evaluation bias caused by a single dataset, but their data construction, maintenance, and continuous updating costs are high.

2.3. Adaptability of general optimization methods to anchor-point evaluation

Active sampling, sample weighting, and relatively mild denoising treatments are all highly beneficial to anchor-point evaluation. These methods generally do not increase the total number of samples by much, yet they can improve the representativeness of the anchors and make the evaluation results more reliable; therefore, they are likely to be feasible solutions.

In contrast, although methods involving complex domain adaptation, multidomain data construction, and multiple rounds of distribution correction can improve cross-domain adaptation to some extent, they are themselves difficult to implement and computationally demanding. If all of the proposed methods are applied to anchor-point evaluation, evaluation time may increase, resource use may grow, and the lightweight nature of anchor-point evaluation may consequently be partly

offset. Therefore, for specific tasks and intended objectives, appropriately simplifying or combining the various methods is a feasible approach.

3. Model evaluation limitations and optimization solutions in anchor-point evaluation tasks

3.1. Main problems at the model evaluation level

In addition to dataset factors, anchor-point evaluation also has certain limitations in its model-evaluation logic and experimental design.

However, most existing studies focus on final performance metrics and pay insufficient attention to the analysis of error sources. Some articles consider only whether the results corresponding to anchor-point evaluation and full-set evaluation are highly consistent, without conducting much analysis of whether this consistency holds across various models, tasks, and data conditions.

The experimental scenarios used in existing related studies may be subject to selection bias. Some studies tend to select tasks or datasets on which good results are easy to obtain, while paying insufficient attention to extreme scenarios in which models are prone to failure. Such experimental designs may produce overly optimistic evaluation results and make it difficult to comprehensively reflect the boundaries of applicability of anchor-point evaluation methods. At the same time, evaluation results obtained through anchors make it difficult to distinguish model errors from anchor-estimation errors. In general, the final evaluation must take into account the quality of both the model and the anchor samples. If the model fails to process the anchor samples successfully, it is difficult to determine whether the model lacks the required capability or whether the anchor samples fail to cover all of the model's weaknesses. It can therefore be concluded that analysis should be used to distinguish between model errors and sampling errors. In addition, a unified error-quantification metric has not yet been established. Different studies currently use somewhat different methods to calculate anchor-point evaluation errors. Common metrics include absolute error, relative error, and correlation coefficients relative to full-set evaluation results, but no unified standard has yet been formed. This affects performance comparisons among different methods to a certain extent.

3.2. Robust model evaluation methods

To address the above problems, this paper constructs a lightweight robust evaluation solution at three levels: statistical stability, multidimensional stress testing, and experimental standardization. All three types of methods use comparisons between anchor subsets and full-set test results as the benchmark. Rather than pursuing an unbounded increase in test items, they use reproducible procedures to determine under what conditions anchor-point evaluation is credible and where its errors originate.

3.2.1. Generalization error and stability analysis

The first type of method uses repeated sampling to quantify the generalization error of anchor-point evaluation. Data may be selected from Chinese-domain samples in education, medicine, law, and other fields, with tasks from similar disciplines in C-Eval and MMLU used as general benchmark comparisons [12, 13]. During implementation, the model version, prompt template, and decoding parameters are fixed, and anchors of the same size are repeatedly drawn from the candidate set under different random seeds. Bootstrap resampling is performed on statistics such as accuracy,

average score, or model ranking, and 95% confidence intervals are reported [15]. The Pearson or Spearman correlation coefficient, mean absolute error, and variance across repeated samples between the anchor results and the full-set results are then calculated, and appropriate significance tests are used to compare different solutions [16]. If the correlation is high and the interval is narrow, the anchor estimate is relatively stable; if the model is stable on the full dataset while the anchor results fluctuate greatly, the error is more likely to originate from the sampling process. This method adds only offline resampling and statistical computation, does not require the model to accept new adversarial tasks, and is suitable for priority use in low-cost error tracing.

3.2.2. Multidimensional robust evaluation

The second type of method uses multidimensional probe testing to examine cross-domain generalization, annotation noise, and input perturbations while controlling the sample size. The probe library can be stratified and sampled from the knowledge, reasoning, and language tasks covered by general comprehensive evaluation frameworks [7-14], and three groups of paired samples can be established: original samples and cross-domain rewritten samples, clean labels and a small number of noisy labels, and original inputs and inputs containing synonym substitutions or interfering information. CheckList's behavioral-testing concept, Robustness Gym's slice-based evaluation method, and TextFooler's text-perturbation strategy can provide references for probe construction [17-19]. During experiments, the number of questions, difficulty, and scoring rules in each dimension should be kept consistent. The degree of performance degradation, model-ranking retention rate, and evaluation cost per sample should be reported separately and then compared with random anchors and full-set robustness testing. If small-scale probes can reproduce the main model rankings and weak dimensions found in full-set testing, they can be considered diagnostically valuable at a lower cost; otherwise, the probe coverage must be expanded, and the effectiveness of the solution cannot be claimed solely on the basis of a single accuracy score.

3.2.3. Standardized experimental procedure

The third type of method reduces human selection bias through a standardized experimental procedure. The specific procedure includes predetermining the rules for dividing the training, validation, and test sets; fixing the random seed, model version, prompt template, temperature, and maximum generation length; establishing control groups such as random sampling, active sampling, and weighted sampling; conducting separate ablation studies on sampling strategies, sample weights, and robustness probes; and uniformly reporting model-ranking fidelity, average cross-domain error, confidence-interval width, running time, and number of inference calls. The reporting methods for significance tests and repeated trials should remain consistent [16], while data preprocessing, the hyperparameter search range, and failed results should also be recorded to improve experimental reproducibility [20]. This method itself adds almost no model-inference steps. Its value lies in establishing fair conditions for horizontal comparison so that the accuracy gains and cost tradeoffs of different lightweight solutions can be examined together.

3.3. Adaptation bottlenecks of robust evaluation methods

Multidimensional robustness testing and error analysis can supplement a single performance metric, but in anchor-point evaluation they are jointly constrained by sample size and cost budgets. When there are too few anchors, statistics by domain and difficulty produce wide confidence intervals. If

large-scale adversarial testing, multifactor ablation, and multiple rounds of cross-domain repeated experiments are conducted simultaneously, the number of inference calls also increases substantially, weakening the lightweight advantage. In addition, if probe samples are generated according to fixed rules, they may cover only a certain type of superficial perturbation and may not represent distribution changes in actual business scenarios. Therefore, adding multidimensional metrics does not necessarily yield more reliable conclusions; the key still lies in the representativeness of the probes and the interpretability of the statistical results.

Therefore, conventional robust evaluation methods should not be applied to anchor-point evaluation unchanged and in their entirety; instead, they should be compressed according to few-sample conditions. During implementation, the risk dimensions that must be covered can first be determined according to the actual purpose, after which a small number of discriminative paired probes can be selected from each dimension and Bootstrap confidence intervals can be used to describe the uncertainty of the few-sample results [15]. When the confidence intervals are too wide, model rankings change frequently, or the probe results are clearly inconsistent with the full-set results, this should be treated as a signal that the anchors are insufficient rather than as a basis for continuing to interpret a single score. The resulting combined solution, in which "a small number of multidimensional probes are responsible for discovering risks, repeated sampling is responsible for quantifying uncertainty, and a standardized procedure is responsible for ensuring comparability," can achieve a relatively reasonable balance between evaluation completeness and computational cost.

4. Discussion

Considering together the defects at the two major levels of datasets and model evaluation discussed above, the core contradiction in anchor-point evaluation lies in the conflict between the inherent shortcomings of few samples and the goal of lightweight evaluation. At the dataset level, the problems include insufficient cross-domain generalization, interference from sampling noise, and limited validation scenarios, and the corresponding optimization methods are active sampling, domain adaptation, data denoising, and multidomain mixed datasets. At the model-evaluation level, the limitations include experimental selection bias, difficulty in distinguishing errors, and the lack of unified quantitative metrics, and improvements can be achieved through stability analysis, multidimensional robustness probes, and standardized experimental procedures. The various optimization solutions each involve tradeoffs: lightweight sampling and weighted denoising have low costs and strong adaptability, whereas complex domain adaptation and large-scale adversarial testing, although capable of improving accuracy, substantially weaken the core advantage of anchor-point evaluation in saving computing power. Therefore, subsequent optimization should not blindly stack various technologies but should design lightweight adaptation solutions around a balance among evaluation accuracy, computational overhead, and actual business scenarios.

At the level of lightweight intelligent sampling, sample-screening methods with low computational costs should be designed to improve the coverage of the overall data distribution by anchor samples under a limited sample size and reduce result fluctuations caused by random sampling. On this basis, an optimization approach suited to few-sample denoising and weighting mechanisms can be established. Information such as sample quality, model confidence, and sample representativeness can be combined to perform lightweight correction of anchor samples, which can both weaken the interference of anomalous samples with the evaluation results and avoid the additional computational overhead produced by complex iterative calculations. In addition, a multidimensional evaluation protocol for few samples needs to be established. Concise and efficient robust evaluation methods should be designed for various anchor-point evaluation application

scenarios, incorporating the model's basic capabilities, cross-domain adaptation performance, resistance to input interference, and the stability of evaluation results into the dimensions examined with few samples. At the same time, confidence-interval calculation, repeated-sampling experiments, and systematic error-analysis methods should be used to improve the interpretability of anchor-point evaluation outputs.

Overall, for the future development of anchor-point evaluation, simply improving evaluation efficiency is not the only objective that should be pursued; reliability and scope of applicability must also be considered. For anchor-point evaluation to be useful for large language models, whether for preliminary screening or practical application, an appropriate balance must be achieved among cost, result accuracy, and stability across situations.

5. Conclusion

This paper discusses the limitations encountered by anchor-point evaluation in both data and model evaluation. The analysis presented in the paper shows that anchor-point evaluation can substantially reduce the number of samples and the amount of computation used to evaluate large language models, but the resulting results may be distorted because of the anchors themselves, sampling, or distribution bias.

From the perspective of datasets, active sampling, sample denoising, and weighted evaluation are all relatively appropriate methods and have positive effects on the quality of anchor samples and the stability of evaluation. Although domain adaptation or complex distribution-correction approaches can improve multidomain adaptation performance, the computational overhead involved can sometimes be relatively high.

To evaluate models, experiments should be handled in a standardized manner and multiple dimensions and errors should all be taken into account so that the results are more objective and interpretable. However, the number of samples used in anchor-point evaluation is limited, and conventional large-scale robust evaluation methods cannot simply be applied directly. Therefore, how to simultaneously achieve complete evaluation, reliable results, and computational economy when there are few samples is a fundamental problem that anchor-point evaluation must currently address.

The analysis throughout this paper shows that there is an inherent tradeoff between lightweight operation and high accuracy, and no general optimization solution is suitable for all business scenarios. The various optimization methods have their own advantages and disadvantages. Lightweight data-optimization solutions can balance evaluation accuracy and computing costs, whereas complex cross-domain correction and multiple rounds of adversarial testing, although capable of broadening the scope of evaluation applicability, impair the core advantages of high efficiency and low cost in anchor-point evaluation. This paper systematically reviews the defects at the two levels and the adaptation boundaries of various optimization paths, providing a complete reference basis for constructing a rapid few-sample evaluation system. Looking to the future, subsequent research can continue to explore three major specific directions. First, low-compute, lightweight intelligent sampling algorithms can be developed to achieve adaptive anchor selection on the basis of sample uncertainty and domain characteristics and reduce result fluctuations caused by random sampling. Second, a few-sample dynamic weighted-denoising mechanism integrating model confidence and sample-annotation quality can be constructed to alleviate evaluation bias caused by anomalous samples through lightweight iteration. Third, a standardized multidimensional robust evaluation protocol suited to few-sample scenarios can be formulated by integrating a small number of cross-domain perturbation probes with Bootstrap error-quantification methods and

unifying the standards for measuring anchor-point evaluation errors. Through continued exploration of the above directions, it is expected that evaluation accuracy, inference overhead, and cross-scenario generalization capability can be coordinated under a limited anchor-sample size, promoting the broad application of anchor-point evaluation in engineering-practice scenarios such as rapid iteration of large models, batch quality inspection for online business, and horizontal comparison of multiple model versions.

References

- [1] Vivek, R., Ethayarajh, K., et al. (2024). Anchor points: Benchmarking models with much fewer examples. *Proceedings of the European Chapter of the Association for Computational Linguistics*.
- [2] Polo, F. M., Weber, L., Choshen, L., et al. (2024). tinyBenchmarks: Evaluating LLMs with fewer examples. *Proceedings of the 41st International Conference on Machine Learning*.
- [3] Luo, W., & Wang, H. F. (2024). Evaluating large language models: A survey of research progress. *Journal of Chinese Information Processing*, 38, 1-23.
- [4] Zhao, K. L., Jin, X. L., & Wang, Y. Z. (2021). Survey on few-shot learning. *Journal of Software*, 32, 349-369.
- [5] Pan, S. J., & Yang, Q. (2010). A survey on transfer learning. *IEEE Transactions on Knowledge and Data Engineering*, 22, 1345-1359.
- [6] Liu, Z. Y., Sun, M. S., Lin, Y. K., et al. (2016). Knowledge representation learning: A survey. *Journal of Computer Research and Development*, 53, 247-261.
- [7] Liang, P., Bommasani, R., Lee, T., et al. (2023). Holistic evaluation of language models. *Transactions on Machine Learning Research*.
- [8] Settles, B. (2009). *Active learning literature survey*. University of Wisconsin-Madison.
- [9] Swayamdipta, S., Schwartz, R., Lourie, N., et al. (2020). Dataset cartography: Mapping and diagnosing datasets with training dynamics. *Proceedings of EMNLP*, 9275-9293.
- [10] Sugiyama, M., Krauledat, M., & Muller, K. R. (2007). Covariate shift adaptation by importance weighted cross validation. *Journal of Machine Learning Research*, 8, 985-1005.
- [11] Northcutt, C. G., Jiang, L., & Chuang, I. L. (2021). Confident learning: Estimating uncertainty in dataset labels. *Journal of Artificial Intelligence Research*, 70, 1373-1411.
- [12] Hendrycks, D., Burns, C., Basart, S., et al. (2021). Measuring massive multitask language understanding. *International Conference on Learning Representations*.
- [13] Huang, Y., Bai, Y., Zhu, Z., et al. (2023). C-Eval: A multi-level multi-discipline Chinese evaluation suite for foundation models. *Advances in Neural Information Processing Systems*, 36.
- [14] Srivastava, A., Rastogi, A., Rao, A., et al. (2023). Beyond the imitation game: Quantifying and extrapolating the capabilities of language models. *Transactions on Machine Learning Research*.
- [15] Efron, B., & Tibshirani, R. J. (1993). *An introduction to the bootstrap*. Chapman & Hall.
- [16] Dror, R., Baumer, G., Shlomov, S., et al. (2018). The hitchhiker's guide to testing statistical significance in natural language processing. *Proceedings of ACL*, 1383-1392.
- [17] Ribeiro, M. T., Wu, T., Guestrin, C., et al. (2020). Beyond accuracy: Behavioral testing of NLP models with CheckList. *Proceedings of ACL*, 4902-4912.
- [18] Goel, K., Rajani, N. F., Vig, J., et al. (2021). Robustness gym: Unifying the NLP evaluation landscape. *Proceedings of NAACL-HLT*, 42-55.
- [19] Jin, D., Jin, Z., Zhou, J. T., et al. (2020). Is BERT really robust? A strong baseline for natural language attack on text classification and entailment. *Proceedings of AAAI*, 34, 8018-8025.
- [20] Dodge, J., Ilharco, G., Schwartz, R., et al. (2019). Show your work: Improved reporting of experimental results. *Proceedings of EMNLP-IJCNLP*, 2185-2194.