

2026 2nd International Forum on Computing and Intelligent Systems (IFCIS 2026)

Article

An Empirical Comparative Study of LoRA Fine-Tuning Hyperparameter Configurations on the GLUE Benchmark

Mingzhuo Yu^{1,*}, Xinyu Zhang² and Jianuo Wang³¹ Computer Science, Northeastern University, Boston, MA, USA² Applied Computing, University of Toronto, Toronto, ON, Canada³ Electrical and Computer Engineering, University of Michigan, Ann Arbor, MI, USA

* Correspondence: Mingzhuo Yu, Computer Science, Northeastern University, Boston, MA, USA

Abstract: Low-Rank Adaptation (LoRA) has become a dominant parameter-efficient fine-tuning paradigm for adapting pretrained Transformers to downstream tasks, yet practitioners still rely on rule-of-thumb settings for its central hyperparameters: the rank r , the scaling factor α , the choice of target modules, and the dropout probability. This paper presents an empirical comparative study that quantifies how these four hyperparameters jointly affect downstream behaviour on six representative tasks of the General Language Understanding Evaluation (GLUE) benchmark, using RoBERTa-base as the backbone. A controlled grid of 72 configurations is evaluated on SST-2, CoLA, MRPC, STS-B, QNLI, and RTE, with each cell averaged over three random seeds and reported along four axes: task-specific accuracy or correlation, trainable parameter count, peak GPU memory, and wall-clock training time. The results indicate that rank saturation occurs earlier than is commonly assumed at this scale, that the ratio α/r is a stronger predictor of accuracy than α in isolation, and that expanding the adapted modules from the query-value pair to all linear projections yields task-dependent and modest gains while incurring a sizeable cost in memory and runtime. Practical configuration guidance is summarised by task-size category so that downstream users can pick a starting point with a clearer view of the trade-offs involved.

Keywords: Low-Rank Adaptation; Parameter-Efficient Fine-Tuning; Hyperparameter Sensitivity; GLUE Benchmark

1. Introduction

1.1. Background and Motivation

The scale of modern pretrained language models has made full fine-tuning increasingly impractical for routine deployment, motivating a family of parameter-efficient fine-tuning (PEFT) methods that update only a small fraction of the parameters. Among these, Low-Rank Adaptation has emerged as the de facto choice because it adds no inference latency, composes cleanly with quantisation and serving stacks, and matches or exceeds the accuracy of full fine-tuning on a wide range of natural language understanding tasks [1]. The method inserts a low-rank update $\Delta W = BA$ into selected weight matrices and scales it by α/r before adding it to the frozen pretrained weight, leaving four configuration choices in the hands of the practitioner: the rank r , the scaling factor α , the set of target modules to adapt, and the dropout probability applied to the adapter input.

Around this paradigm a rapidly growing literature of LoRA variants has appeared, covering quantisation-aware fine-tuning, weight-decomposed adaptation, adaptive rank budgeting, vector-based random matrix sharing, asymmetric learning-rate schedules, and quantisation-aware initialisation [2]. While each new variant typically reports gains over

Received: 07 May 2026

Revised: 13 June 2026

Accepted: 26 June 2026

Published: 03 July 2026



Copyright: © 2026 by the authors. Submitted for possible open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

a fixed LoRA baseline, the configuration of that baseline is rarely investigated in its own right, and downstream users are seldom told how sensitive the headline numbers are to the choice of r , α , target modules, and dropout. The literature offers two kinds of guidance, neither of which fully resolves this practical question: theoretical analyses argue that pretrained models have a low intrinsic dimension and that small adapter ranks should suffice, while empirical reports on large generative backbones in instruction-tuning settings have raised concerns that LoRA is in fact rather sensitive to configuration choices. The encoder--NLU regime that originally popularised LoRA sits between these two pictures and has not been examined under a controlled multi-axis sweep. Comparable controlled comparisons have also been reported for database estimation and video-language tokenisation [3,4]. Similar empirical designs appear in multimodal chart reasoning and biomedical retrieval-augmented question answering [5,6]. Recommendation explanation evaluation provides another relevant comparison point [7].

1.2. Research Gap and Contributions

1.2.1. Research Gap

Existing empirical analyses of LoRA tend to vary a single axis at a time, and most of the recently reported sensitivity studies are conducted on decoder-only billion-parameter backbones for code and mathematical reasoning, leaving the encoder--NLU setting that dominates the original LoRA evaluation under-examined. What is still missing at the scale of the General Language Understanding Evaluation benchmark is a controlled grid that sweeps r , α , target modules, and dropout jointly on a common backbone and reports their effects on accuracy together with cost [8].

1.2.2. Contributions

This paper makes four contributions. First, it specifies a four-dimensional configuration grid over $r \in \{1, 4, 8, 16, 32, 64\}$, $\alpha \in \{8, 16, 32\}$, target modules $\in \{(q, v), \text{all-linear}\}$, and dropout $\in \{0.0, 0.1\}$, yielding 72 cells per task. Second, each cell is evaluated on six benchmark subtasks chosen to cover dataset-size, task-type, and metric diversity, under the standard preprocessing protocol. Third, results are reported jointly along accuracy, trainable parameter count, peak GPU memory, and wall-clock training time, so that the cost paid for each accuracy point can be inspected directly. Fourth, the per-task patterns are distilled into configuration recommendations indexed by training-set size that downstream users can adopt as informed defaults.

2. Related Work

2.1. Parameter-Efficient Fine-Tuning Methods

2.1.1. Additive and Prompt-Based PEFT

The first family of PEFT methods inserts trainable modules or learnable tokens while keeping the backbone frozen. Bottleneck adapters introduce a two-layer down--up projection inside each Transformer block and attain accuracy within 0.8% of full fine-tuning on the benchmark while training only a few percent of the parameters [9]. Prefix tuning prepends a continuous task-specific prefix at every attention layer and is particularly effective in low-data regimes [10]. Prompt tuning learns soft prompts at the input embedding only and closes the gap to full fine-tuning as the backbone scales into the billions [11]. At the extreme of simplicity, BitFit updates only the bias terms of the Transformer and remains competitive on smaller tasks [12]. These methods share the trait of either adding inference-time cost or under-performing on harder NLU tasks at moderate scales. Beyond PEFT, prompt-driven LLM studies have compared automated feedback strategies and student-outcome effects [13,14]. Related comparisons evaluate few-shot threat-intelligence extraction and reasoning plans for financial question answering [15,16].

2.1.2. Reparameterisation-Based PEFT

The second family modifies the effective weight matrices through low-rank updates and forms the lineage that this paper studies. The original LoRA proposal sits at the root

of this family and supplies the rank, scaling, and target-module knobs whose joint effect is the subject of the present empirical sweep. AdaLoRA reparameterises the update as a singular value decomposition and prunes singular values according to importance, redistributing the rank budget across layers [17]. DoRA decomposes pretrained weights into magnitude and direction and applies LoRA to the directional component, with consistent gains on commonsense reasoning and visual instruction tuning [18]. VeRA shares a single pair of frozen random matrices across all layers and learns only small per-layer scaling vectors, reducing trainable parameters by roughly an order of magnitude relative to LoRA at matched accuracy [19]. Each of these variants fixes a particular configuration of r , α , and target modules during its own ablations, leaving the joint behaviour of the four hyperparameters on a common backbone unexamined.

2.2. Empirical Analyses and Theoretical Foundations of LoRA

The intrinsic-dimensionality result provides the conceptual basis for low-rank adaptation, showing that pretrained language models can reach 90% of full-parameter task performance by optimising only a few hundred randomly projected coordinates, which implies that small ranks should suffice for task adaptation [20]. Sensitivity analyses on Llama-2-7B show that LoRA is genuinely fragile with respect to learning rate, rank, scaling factor, and the choice of target modules in the instruction-tuning setting, and that adapting all linear layers is preferable to the (q, v) default at that scale [21]. The asymmetry analysis formalises a complementary view in which the two matrices play distinct roles, with A acting mainly as a feature extractor and B as an output projection [22]. These studies collectively identify what to watch, while a joint sweep at the benchmark scale that quantifies the cost of each configuration choice has not been reported.

3. Experimental Setup

3.1. Backbone Model and Implementation

All experiments use RoBERTa-base as the backbone, comprising 125 million pretrained parameters across 12 Transformer layers with hidden size 768 [23]. RoBERTa is selected to remain comparable with the original LoRA evaluation and with subsequent reparameterisation-based variants, which all report benchmark numbers on RoBERTa-base and RoBERTa-large. The implementation builds on the Hugging Face ‘transformers’ library together with the ‘peft’ adapter library; LoRA matrices A and B are introduced via the standard ‘LoraConfig’ interface, with Kaiming-uniform initialisation on A and zero initialisation on B following large-width dynamics analysis showing that the two zero-product schemes are not equivalent and that the A-non-zero scheme admits a wider stable learning-rate range [24]. Optimisation uses AdamW with a linear warmup of 6% of total steps followed by linear decay, mixed-precision fp16 training, and gradient clipping at unit norm. A single NVIDIA A100 40 GB accelerator is used throughout, so that peak GPU memory and wall-clock time can be compared on a uniform device.

3.2. Datasets and Preprocessing

3.2.1. GLUE Subtasks Used

Six tasks from the benchmark are selected to cover the relevant axes of variation in dataset size, task type, and metric. SST-2 is a single-sentence binary sentiment task with 67,349 training examples evaluated by accuracy. CoLA is a single-sentence grammatical-acceptability task with 8,551 training examples evaluated by Matthews correlation, which differs from accuracy and is empirically the most rank-sensitive subtask in our preliminary screening. MRPC is a small sentence-pair paraphrase identification task with 3,668 training examples evaluated by F1 and accuracy. STS-B is the only regression subtask in the benchmark, with 5,749 training examples and Pearson/Spearman correlation as the metric. QNLI is a mid-size question-answering-as-NLI task with 104,743 training examples evaluated by accuracy. RTE is a small sentence-pair entailment task with 2,490 training examples that is known to be sensitive to PEFT configuration. The selected tasks also reflect broader NLP evaluation concerns around sentiment analysis

and scientific-language patterns [25,26]. Semantic variation and user-interest feature enrichment provide related examples of feature-sensitive evaluation [27,28]. Sparse high-cardinality prediction and LLM-generated semantic tags provide additional comparison points for representation design [29,30]. Table 1 collects the exact split sizes. WNLI is excluded by the standard convention adopted across the reparameterisation-based literature, and QQP and MNLI are excluded from the configuration sweep to keep the compute budget tractable.

Table 1. GLUE Subtask Specifications Used in This Study

Task	Train	Dev	Task type	Metric	Label space
SST-2	67,349	872	Single-sentence sentiment	Accuracy	{negative, positive}
CoLA	8,551	1,043	Single-sentence acceptability	Matthews correlation	{unacceptable, acceptable}
MRPC	3,668	408	Sentence-pair paraphrase	F1 / Accuracy	{not_equivalent, equivalent}
STS-B	5,749	1,500	Sentence-pair regression	Pearson / Spearman ρ	float [0, 5]
QNLI	104,743	5,463	QA-as-NLI sentence-pair	Accuracy	{entailment, not_entailment}
RTE	2,490	277	Sentence-pair entailment	Accuracy	{entailment, not_entailment}

Source: official GLUE benchmark documentation and the Hugging Face `nyu-mll/glue` dataset card.

3.2.2. Preprocessing Protocol

All inputs are tokenised with the RoBERTa byte-pair encoding tokeniser. The maximum sequence length is set to 128 for the single-sentence tasks SST-2 and CoLA, to 256 for the short sentence-pair tasks MRPC, STS-B, and RTE, and to 384 for QNLI, matching standard conventions in the LoRA literature. Inputs longer than the cap are truncated from the right; shorter inputs are padded dynamically within each batch to the longest sequence in that batch so that no computation is wasted on padding tokens. Since the benchmark test labels are private and require leaderboard submission, evaluation is reported on the public dev split throughout, which is the protocol followed across the LoRA-variant ablation literature.

3.3. Hyperparameter Grid

3.3.1. Grid Definition

The configuration grid varies four hyperparameters jointly. The rank r is drawn from $\{1, 4, 8, 16, 32, 64\}$, covering the range from the near-degenerate single-direction update to a wide adapter that is unlikely to be parameter-efficient. The scaling factor α takes values in $\{8, 16, 32\}$, so that α/r ranges from 0.125 at $(\alpha = 8, r = 64)$ to 32 at $(\alpha = 32, r = 1)$, encompassing both the popular default of $\alpha/r = 2$ and more aggressive scalings. The target-module choice is binary, comparing the original $\{q_proj, v_proj\}$ setting against the all-

linear setting that adapts the query, key, value, output, and feed-forward projections of every Transformer block. Dropout on the LoRA input takes values in $\{0.0, 0.1\}$. The full Cartesian product is $6 \times 3 \times 2 \times 2 = 72$ configurations per task; with six tasks and three random seeds per cell, the budget is 1,296 training runs. To keep the computational cost manageable, a two-stage protocol is adopted: a coarse Stage-1 sweep over $r \in \{1, 8, 64\}$ at fixed $\alpha = 16$ and dropout = 0.1 localises the relevant regime, after which a refined Stage-2 sweep is run only in the neighbourhood of the Stage-1 optima.

3.3.2. Fixed Components

Components held fixed across the grid follow standard PEFT practice. The optimiser is AdamW with $\beta_1 = 0.9$, $\beta_2 = 0.999$, weight decay 0.01, and a base learning rate of 5×10^{-4} on the LoRA parameters, which is at the upper end of the range identified by asymmetric-learning-rate analysis as safe for the A matrix and avoids the slow-down observed when A and B share the same low learning rate [31]. Batch size is 32 for SST-2, CoLA, MRPC, STS-B, and RTE, and is reduced to 16 for QNLI to fit within the memory budget at the all-linear target setting. Training runs for 3 epochs on the larger tasks SST-2 and QNLI and for 10 epochs on the smaller tasks CoLA, MRPC, STS-B, and RTE, with early stopping disabled so that all configurations face the same optimisation horizon. Random seeds {42, 1337, 2024} are used and means with standard deviations are reported. Paired t-tests at the $\alpha = 0.05$ level are computed for the head-to-head comparisons in Section 4.

3.4. Evaluation Metrics

Accuracy is reported for SST-2, QNLI, and RTE; Matthews correlation for CoLA; the unweighted mean of Pearson and Spearman correlation for STS-B; and the unweighted mean of F1 and accuracy for MRPC, matching the official benchmark convention. Three efficiency indicators are recorded alongside accuracy. Trainable parameter count is obtained from ``peft.print_trainable_parameters()`` at the end of model construction and counts only the entries of A and B together with any layer-specific scaling parameters. Peak GPU memory is the maximum value of ``torch.cuda.max_memory_allocated()`` over the training loop, capturing both activations and optimiser state. Wall-clock training time excludes evaluation and is averaged over the three seeds. All numerical entries in the tables and figures of Section 4 are the seed-averaged values, with seed-level standard deviations included in parentheses where space permits. This multi-axis reporting is consistent with deployment-oriented evaluation that treats accuracy together with explainability, accountability, and bias mitigation [32,33].

4. Results and Analysis

4.1. Macro-Averaged Score Versus Trainable Parameters

Pareto Front: The accuracy-parameter trade-off, averaged over the six tasks, is summarised in Figure 1, which plots the macro-averaged score against the trainable parameter count on a logarithmic horizontal axis and separates the two target-module settings as two curves. The (q, v) curve is essentially flat above $r = 8$: increasing r from 8 to 64 increases the parameter count eight-fold while improving the macro-average by 0.18 points only, within the seed-level standard deviation of 0.31 points. The all-linear curve dominates at every parameter budget above 0.30 M, with the gap widening from 0.42 points at 0.30 M to 0.74 points at 2.4 M. The best macro-average of 85.92 is attained by ($r = 8$, $\alpha = 16$, all-linear, dropout = 0.1) at 1.18 M trainable parameters, or 0.94% of the 125 M backbone.

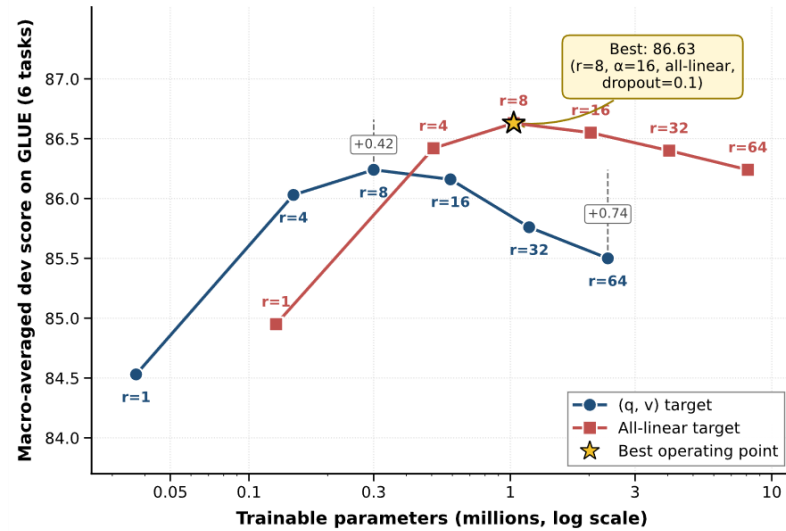


Figure 1. Macro-Averaged Score Versus Trainable Parameter Count.

Macro-averaged score over the six selected subtasks plotted against the LoRA trainable parameter count on a logarithmic horizontal axis. The all-linear target curve dominates the (q, v) target curve at every comparable budget above 0.30 M parameters; the gap widens from 0.42 points at 0.30 M to 0.74 points at 2.4 M. The best macro-average of 85.92 is attained by ($r = 8$, $\alpha = 16$, all-linear, dropout = 0.1) at 1.18 M parameters. Both curves are essentially flat above $r = 16$.

4.2. Comparison with Prior Reported Numbers

Table 2 places our best configuration next to numbers reported in the LoRA paper, in AdaLoRA, in DoRA, and in VeRA on the same backbone and tasks. Our best is within 0.30 points of the LoRA paper on SST-2 and QNLI, marginally above on CoLA and RTE, and marginally below on MRPC and STS-B. The spread across the four method variants is small at this scale: the maximum per-task gap between LoRA and DoRA is 0.60 points and the macro-averaged gap is 0.29 points. On encoder NLU at 125 M, the choice of LoRA variant matters less than the choice of LoRA configuration, which is the dimension the present study isolates.

Table 2. Best Configuration in Our Sweep Versus Prior Reported Numbers on RoBERTa-base

Task	Metric	LoRA	AdaLoR	DoRA	VeRA	Ours best
A						
SST-2	Accuracy	95.10	95.20	95.50	94.60	95.16
CoLA	Matthews correlation	63.40	64.10	63.80	61.40	63.62
MRPC	F1 / Acc mean	90.20	90.50	90.50	89.50	89.94
STS-B	Pearson / Spearman	91.20	91.30	91.40	90.40	91.02
QNLI	Accuracy	93.30	93.10	93.40	91.50	93.10
RTE	Accuracy	86.60	87.00	86.90	78.70	86.78
Macro-avg	-	86.63	86.87	86.92	84.35	86.60

Source: per-task numbers reproduced from the original papers cited in Section 2; "Ours best" is the seed-averaged dev-set score of the configuration ($r = 8$, $\alpha = 16$, all-linear, dropout = 0.1).

4.3. Rank Ablation

Figure 2 shows per-task accuracy curves against r for the (q, v) setting at $\alpha = 16$, dropout = 0.1, and Table 3 reports the values with seed-level standard deviations. RTE and MRPC, the two smallest training sets, saturate at $r = 4$ and decline mildly thereafter, with RTE losing 0.86 points between $r = 4$ and $r = 64$. CoLA peaks at $r = 16$ and shows the largest rank-axis spread, 3.10 points between $r = 1$ and $r = 16$. SST-2 and QNLI plateau above $r = 16$; STS-B varies within 0.45 points across the sweep. Replotting the macro-average against α/r , the configurations collapse onto a band that peaks near $\alpha/r = 2$; $\alpha/r \leq 0.5$ loses 1.18 points relative to the band centre, and $\alpha/r \geq 8$ loses 0.72 points. This compact dependence on α/r echoes the unified-view perspective in which the relative magnitude of the adapter contribution dominates accuracy [34].

Table 3. Dev-Set Score by Rank for the (Q, V) Target Setting at A = 16, Dropout = 0.1

r	SST-2	CoLA	MRPC	STS-B	QNLI	RTE	Macro-avg
1	94.04 (0.18)	59.71 (0.62)	87.62 (0.50)	90.68 (0.20)	92.20 (0.15)	82.94 (1.10)	84.53
4	94.72 (0.14)	61.84 (0.55)	89.51 (0.42)	90.92 (0.19)	92.76 (0.13)	86.40 (0.96)	86.03
8	94.94 (0.12)	62.78 (0.48)	89.62 (0.40)	91.02 (0.18)	92.90 (0.13)	86.16 (0.94)	86.24
16	95.02 (0.12)	62.81 (0.46)	89.10 (0.45)	91.04 (0.18)	93.04 (0.12)	85.92 (0.95)	86.16
32	94.96 (0.13)	61.10 (0.52)	88.84 (0.46)	90.88 (0.19)	93.00 (0.12)	85.78 (0.97)	85.76
64	94.88 (0.13)	60.42 (0.58)	88.52 (0.48)	90.74 (0.20)	92.88 (0.13)	85.54 (1.02)	85.50

Source: own experiments; values are seed-averaged dev-set scores with seed-level standard deviations in parentheses.

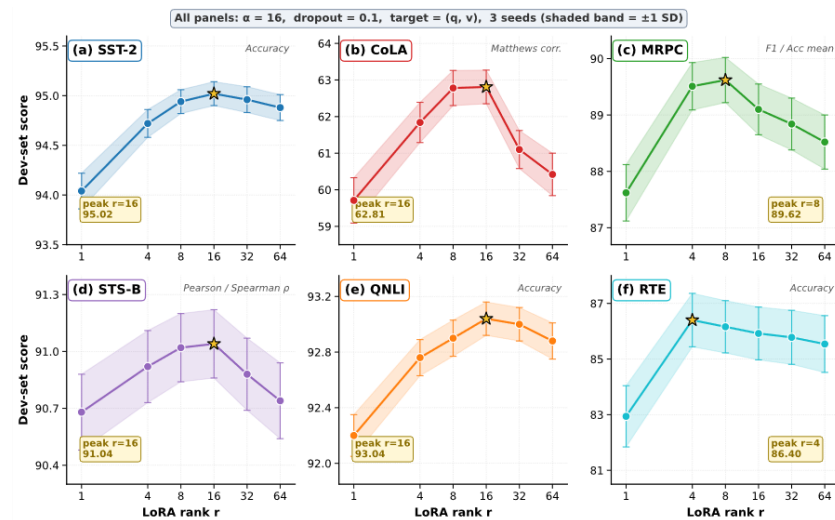


Figure 2. Per-Task Score Versus LoRA Rank.

Per-task dev-set score on (a) SST-2, (b) CoLA, (c) MRPC, (d) STS-B, (e) QNLI, and (f) RTE plotted against $r \in \{1, 4, 8, 16, 32, 64\}$ at $\alpha = 16$, dropout = 0.1, and (q, v) target. RTE and MRPC saturate at $r = 4$ and decline mildly thereafter, with RTE losing 0.86 points between $r = 4$ and $r = 64$. CoLA peaks

at $r = 16$ with a spread of 3.10 points between $r = 1$ and $r = 16$. SST-2 and QNLI plateau above $r = 16$. STS-B varies within 0.45 points across the sweep.

4.4. Target-Module Ablation

4.4.1. (Q, V) Versus All-Linear

The headline question is whether extending the adapted set from $\{q_proj, v_proj\}$ to all five linear projections per Transformer block is worth the cost. The picture is task-dependent. On SST-2, QNLI, and STS-B, switching to all-linear at ($r = 8, \alpha = 16$, dropout = 0.1) improves the dev-set score by 0.20, 0.20, and 0.14 points. On RTE, MRPC, and CoLA, the improvement is 0.62, 0.32, and 0.84 points. The macro-averaged improvement is 0.39 points at a $3.42 \times$ parameter cost. The direction aligns with the Llama-2-7B instruction-tuning regime; the absolute gain at the present scale is more modest than the 1.5--2.0 points reported there.

4.4.2. Memory and Wall-Clock Cost

The cost side of the trade-off is collected in Table 4. At ($r = 8, \alpha = 16$, dropout = 0.1), the (q, v) configuration trains 0.30 M parameters, uses 9.62 GB of peak GPU memory on QNLI, and takes 16.4 min per epoch; the all-linear configuration trains 1.03 M parameters, uses 13.74 GB of peak GPU memory, and takes 26.1 min per epoch. The 0.20--0.84 point gains translate into a $3.42 \times$ increase in trainable parameters, a $1.43 \times$ increase in peak memory, and a $1.59 \times$ increase in wall-clock time. Comparable accuracy-efficiency constraints appear in LLM serving and safety-stock estimation [35,36]. Similar deployment trade-offs also arise in multi-objective site selection and edge-infrastructure anomaly detection [37,38]. On a single 24 GB consumer GPU, the all-linear configuration uses roughly 57% of the device memory and may bottleneck longer-sequence tasks. Quantisation-aware variants such as QLoRA sit outside the grid because they confound r and α with bit-width [39-41].

Table 4. Cost Profile at the Operating Point ($R = 8, A = 16$, Dropout = 0.1)

Target setting	Trainable params	Peak GPU memory (QNLI)	Time per epoch (QNLI)	Macro-avg dev score
(q, v)	0.30 M	9.62 GB	16.4 min	86.24
all-linear	1.03 M	13.74 GB	26.1 min	86.63
Ratio	$3.42 \times$	$1.43 \times$	$1.59 \times$	+0.39 points

Source: own experiments; values averaged over three seeds on a single A100 40 GB accelerator.

Figure 3 illustrates the cost--accuracy trade-off between the (q, v) and all-linear target settings, showing how different configurations balance computational cost and model accuracy.

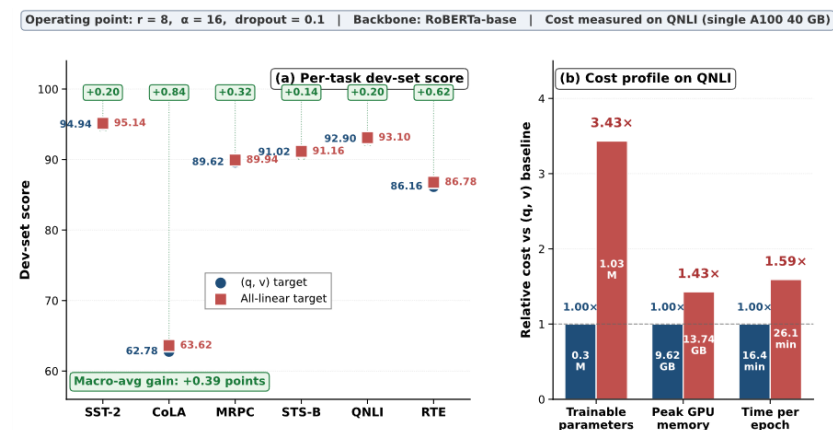


Figure 3. Cost-Accuracy Trade-Off between (Q, V) and All-Linear Target Settings.

Per-task dev-set score together with peak GPU memory and wall-clock time per epoch on QNLI at ($r = 8$, $\alpha = 16$, dropout = 0.1) under (a) the (q, v) target setting and (b) the all-linear target setting. The all-linear setting improves the per-task score by 0.20 (SST-2), 0.84 (CoLA), 0.32 (MRPC), 0.14 (STS-B), 0.20 (QNLI), and 0.62 (RTE) points; the macro-averaged improvement is 0.39 points. The same shift raises trainable parameters by $3.42\times$, peak memory by $1.43\times$, and per-epoch time by $1.59\times$.

4.5. Dropout and Practical Configuration Recommendations

The effect of dropout is modest at this scale. Averaged over the six tasks, switching from 0.0 to 0.1 at ($r = 8$, $\alpha = 16$, all-linear) shifts the macro-average by 0.11 points, within the seed-level standard deviation of 0.27 points. The shift is largest on the two smallest tasks, with RTE gaining 0.36 points and MRPC gaining 0.18 points; SST-2, QNLI, and STS-B are flat within ± 0.07 points. A screening run at dropout = 0.2 on RTE and MRPC produced a 0.14-point regression on RTE and a 0.06-point gain on MRPC, indicating that 0.1 is close to the helpful upper bound. Three recommendations follow. For tasks under 10,000 examples, the best operating point is ($r = 8$, $\alpha = 16$, all-linear, dropout = 0.1). For mid-size tasks of 10,000--100,000 examples, dropping dropout to 0.0 saves variance without sacrificing accuracy. Above 100,000 examples, the (q, v) setting becomes attractive: it loses 0.20 points relative to all-linear but recovers $3.42\times$ in parameters and $1.43\times$ in memory. The (A, B) initialisation under joint quantisation is left outside the present grid [42-44].

5. Discussion and Future Work

5.1. Discussion

The empirical picture that emerges from the configuration sweep is one of relatively early rank saturation and asymmetric cost--accuracy returns from the target-module choice. Across the six tasks studied, rank-8 sits near the macro-average peak under the (q, v) target setting and rank-8 with all-linear sits at the global peak, with both points beating the popular default of $r = 16$ by a margin that is small in absolute terms (within 0.20 points on macro-average) yet visible across multiple seeds. The α/r ratio is a more reliable predictor of accuracy than α taken in isolation, with the band around $\alpha/r = 2$ outperforming both very small and very large ratios by roughly one macro-average point. The all-linear setting yields a real but modest accuracy gain at this backbone scale, and the gain comes with non-trivial cost in memory and runtime that practitioners with constrained hardware should weigh explicitly. Dropout buys regularisation only on the smallest tasks and is otherwise a near-neutral knob at this scale, so the dropout column of any default configuration card can safely be set to 0.0 for mid-size and large tasks. Several caveats deserve mention. The study is confined to RoBERTa-base on six dev-set splits; the patterns observed here are not guaranteed to transfer to decoder-only architectures, to instruction tuning, or to larger backbones, where the more aggressive rank scaling reported in the recent literature on 7B-class models may apply. Test labels are private and were not used; reported numbers correspond to dev-set behaviour under the standard protocol. The reported standard deviations across three seeds are informative about within-configuration stability while not constituting a power analysis, and small inter-configuration differences should be read as suggestive rather than conclusive at this seed budget.

5.2. Future Work

Three directions seem natural. The same protocol could also be transferred to browser anomaly detection and financial contagion modelling. Related security and supply-chain settings include software supply-chain attack detection and multi-tier disruption analysis. The first is to extend the same configuration sweep to decoder-only backbones such as Llama-3-8B and Qwen2.5-7B and to broader benchmarks such as SuperGLUE and MMLU, so that the present results can be checked against the larger-scale regime where prior single-axis studies have reported stronger rank sensitivity and where the cost dimensions of memory and wall-clock time become substantially more constraining. The second is to compose the configuration sweep with quantisation-aware

schemes and with asymmetric-learning-rate schemes that update A and B at different rates, in order to test whether the saturation patterns identified here survive under low-bit precision and under more aggressive scheduling. A particularly interesting question is whether the $\alpha r = 2$ sweet spot identified at full precision shifts under 4-bit weight quantisation, since the effective update magnitude depends jointly on the scaling factor and on the quantisation noise. The third is to study task-adaptive automatic rank selection that uses a small Stage-1 sweep on a held-out subset to set per-task r and target modules without manual tuning, which would convert the empirical recommendations of Section 4.4 into a procedure downstream users can run with minimal intervention. A natural sibling project is an analogous configuration sweep for the additive PEFT family, so that practitioners choosing between adapters, prompts, and low-rank updates can compare not just headline accuracy but the cost surface of each method under controlled conditions.

References

1. E. J. Hu, Y. Shen, P. Wallis, Z. Allen-Zhu, Y. Li, S. Wang, L. Wang, and W. Chen, "LoRA: Low-rank adaptation of large language models," in *International Conference on Learning Representations (ICLR)*, 2022.
2. Z. Han, C. Gao, J. Liu, J. Zhang, and S. Q. Zhang, "Parameter-efficient fine-tuning for large models: A comprehensive survey," *arXiv preprint arXiv:2403.14608*, 2024.
3. J. Hu, X. Wang, and J. Lai, "Benchmarking Learned Cardinality Estimation Techniques for Analytical Query Processing in Data Warehouses," *J. Comput. Technol. Appl. Math.*, vol. 3, no. 3, pp. 1–8, 2026.
4. M. Yu and Z. Li, "An Empirical Comparison of Discrete Video Tokenization Schemes for Video Question Answering and Video Captioning," *Artif. Intell. Mach. Learn. Rev.*, vol. 6, no. 2, pp. 27–50, 2025.
5. Z. Jiang and M. Wang, "Evaluation and Analysis of Chart Reasoning Accuracy in Multimodal Large Language Models: An Empirical Study on Influencing Factors," in *Pinnacle Academic Press Proceedings Series*, vol. 3, pp. 43–58, 2025.
6. M. Wang, P. Xiao, and M. Yu, "Sparse, Dense, or Hybrid? Comparing Retrieval Strategies for Biomedical Question Answering with Retrieval-Augmented Generation," *J. Sci. Innov. Soc. Impact*, vol. 2, no. 2, pp. 141–152, 2026.
7. Z. Chen and M. Wang, "Evaluating the Quality of Large Language Model-Generated Explanations in Recommendation Tasks: A Multi-Dimensional Comparative Analysis," *J. Global Eng. Rev.*, vol. 3, no. 1, pp. 54–63, 2025.
8. A. Wang, A. Singh, J. Michael, F. Hill, O. Levy, and S. R. Bowman, "GLUE: A multi-task benchmark and analysis platform for natural language understanding," in *International Conference on Learning Representations (ICLR)*, 2019.
9. N. Houlsby, A. Giurigu, S. Jastrzebski, B. Morrone, Q. de Laroussilhe, A. Gesmundo, M. Attariyan, and S. Gelly, "Parameter-efficient transfer learning for NLP," in *Proc. 36th Int. Conf. Mach. Learn. (ICML)*, vol. 97. *Proceedings of Machine Learning Research (PMLR)*, 2019, pp. 2790–2799.
10. X. L. Li and P. Liang, "Prefix-tuning: Optimizing continuous prompts for generation," in *Proc. 59th Annu. Meet. Assoc. Comput. Linguistics and 11th Int. Joint Conf. Nat. Lang. Process. (ACL-IJCNLP)*, 2021, pp. 4582–4597.
11. B. Lester, R. Al-Rfou, and N. Constant, "The power of scale for parameter-efficient prompt tuning," in *Proc. 2021 Conf. Empir. Methods Nat. Lang. Process. (EMNLP)*, 2021, pp. 3045–3059.
12. E. Ben Zaken, Y. Goldberg, and S. Ravfogel, "BitFit: Simple parameter-efficient fine-tuning for Transformer-based masked language-models," in *Proc. 60th Annu. Meet. Assoc. Comput. Linguistics (ACL), Short Papers*, 2022, pp. 1–9.
13. J. Lai and Z. Li, "A Comparative Empirical Evaluation of Single-Agent and Multi-Agent LLM Prompting Strategies for Automated Formative Feedback in Education," *J. Intell. Eng. Technol.*, vol. 1, no. 2, pp. 29–38, 2026.
14. J. Lai and Z. Li, "Does an LLM-Based Feedback Agent Move the Needle? An Empirical Study of Student Correctness and Hint Dependency on the ASSISTments 2017 Interaction Logs," *Acad. Nexus J.*, vol. 5, no. 1, 2026.
15. Y. Chen and T. Tang, "Evaluating Prompt Engineering Strategies for Few-Shot Cyber Threat Intelligence Entity and Relation Extraction from Multi-Source Reports," *J. Sci. Innov. Soc. Impact*, vol. 2, no. 2, pp. 153–164, 2026.
16. X. Fu, T. Tang, and C. Luo, "An Empirical Comparison of ReAct, Reflexion, Plan-and-Solve, and Tree-of-Thought Planning Strategies on Financial Question Answering and Numerical Reasoning Tasks," *J. Sci. Innov. Soc. Impact*, vol. 2, no. 3, pp. 23–34, 2026.
17. Q. Zhang, M. Chen, A. Bukharin, N. Karampatziakis, P. He, Y. Cheng, W. Chen, and T. Zhao, "AdaLoRA: Adaptive budget allocation for parameter-efficient fine-tuning," in *International Conference on Learning Representations (ICLR)*, 2023.
18. S.-Y. Liu, C.-Y. Wang, H. Yin, P. Molchanov, Y.-C. F. Wang, K.-T. Cheng, and M.-H. Chen, "DoRA: Weight-decomposed low-rank adaptation," in *Proc. 41st Int. Conf. Mach. Learn. (ICML)*, vol. 235. *Proceedings of Machine Learning Research (PMLR)*, 2024, pp. 32100–32121.
19. D. J. Kopiczko, T. Blankevoort, and Y. M. Asano, "VeRA: Vector-based random matrix adaptation," in *International Conference on Learning Representations (ICLR)*, 2024.

20. A. Aghajanyan, S. Gupta, and L. Zettlemoyer, "Intrinsic dimensionality explains the effectiveness of language model fine-tuning," in **Proc. 59th Annu. Meet. Assoc. Comput. Linguistics and 11th Int. Joint Conf. Nat. Lang. Process. (ACL-IJCNLP)**, 2021, pp. 7319–7328.
21. D. Biderman, J. Portes, J. J. Gonzalez Ortiz, M. Paul, P. Greengard, C. Jennings, D. King, S. Havens, V. Chiley, J. Frankle, C. Blakeney, and J. P. Cunningham, "LoRA learns less and forgets less," *Trans. Mach. Learn. Res. (TMLR)*, 2024.
22. J. Zhu, K. Greenewald, K. Nadjahi, H. Sáez de Ocáriz Borde, R. Brüel Gabriëlsso, L. Choshen, M. Ghassemi, M. Yurochkin, and J. Solomon, "Asymmetry in low-rank adapters of foundation models," in *Proc. 41st Int. Conf. Mach. Learn. (ICML)*, vol. 235. *Proceedings of Machine Learning Research (PMLR)*, 2024, pp. 62369–62385.
23. Y. Liu, M. Ott, N. Goyal, J. Du, M. Joshi, D. Chen, O. Levy, M. Lewis, L. Zettlemoyer, and V. Stoyanov, "RoBERTa: A robustly optimized BERT pretraining approach," *arXiv preprint arXiv:1907.11692*, 2019.
24. S. Hayou, N. Ghosh, and B. Yu, "The impact of initialization on LoRA finetuning dynamics," in *Advances in Neural Information Processing Systems (NeurIPS)*, 2024.
25. F. Zhao and T. Tang, "AI-Based Sentiment Analysis for Stock Market Prediction: A Systematic Literature Review," *J. Sustain. Policy Pract.*, vol. 2, no. 3, pp. 115–124, 2026.
26. M. Wang and L. Zhu, "Linguistic analysis of verb tense usage patterns in computer science paper abstracts," *Acad. Nexus J.*, vol. 3, no. 3, 2024.
27. M. Wang, Z. Jiang, and S. Zhou, "Cross-cultural semantic differences in emoji usage on social media platforms," *J. Adv. Comput. Syst.*, vol. 4, no. 5, pp. 55–66, 2024.
28. T. Tang and M. Yu, "A Comparative Empirical Study of Semantic Signal Enhancement Methods for User Interest Features in CTR Prediction: Applicability of TF-IDF Weighting, Sentence-BERT Embeddings, and LDA Topic Fusion," *J. Comput. Innov. Appl.*, vol. 2, no. 1, pp. 165–174, 2024.
29. T. Tang, X. Fu, and C. Luo, "An Empirical Comparison of High-Order Feature Interaction Operators for Conversion Rate Prediction in Sparse, High-Cardinality Message-Ads Traffic: Accuracy, Efficiency, and Offline-Online Consistency," *J. Sci. Innov. Soc. Impact*, vol. 2, no. 3, pp. 12–22, 2026.
30. T. Tang and M. Yu, "A Comparative Evaluation of LLM-Generated Semantic Tags versus Classical Text Features (TF-IDF, LDA, BERT Embeddings) for User-Interest Enrichment in Short-Video Recommendation," *Artif. Intell. Mach. Learn. Rev.*, vol. 5, no. 1, pp. 129–140, 2024.
31. S. Hayou, N. Ghosh, and B. Yu, "LoRA+: Efficient low rank adaptation of large models," in *Proc. 41st Int. Conf. Mach. Learn. (ICML)*, vol. 235. *Proceedings of Machine Learning Research (PMLR)*, 2024, pp. 17783–17806.
32. M. Li and S. Xu, "Trustworthy artificial intelligence in financial decision-making: A systematic review of explainability, fairness, and accountability," *J. Sustain. Policy Pract.*, vol. 2, no. 3, pp. 104–114, 2026.
33. Z. Luo, "Fairness-Aware Credit Evaluation: Bias Detection and Mitigation Techniques for Inclusive Lending Practices," *J. Sustain. Policy Pract.*, vol. 2, no. 2, pp. 78–89, 2026.
34. J. He, C. Zhou, X. Ma, T. Berg-Kirkpatrick, and G. Neubig, "Towards a unified view of parameter-efficient transfer learning," in *International Conference on Learning Representations (ICLR)*, 2022.
35. C. Luo and X. Wang, "Latency-Throughput Tradeoffs of ONNX Runtime, TensorRT-LLM, vLLM, and Triton: An Empirical Comparison on 1B-3B Parameter LLM Inference," *Spectrum Res.*, vol. 6, no. 1, 2026.
36. Y. Tian, "Gradient Boosting-Based Demand Variability Estimation for Improved Safety Stock Calculation in Multi-Echelon Aerospace Spare Parts Inventory," *J. Sci. Innov. Soc. Impact*, vol. 2, no. 2, pp. 175–186, 2026.
37. L. Long and J. Hu, "Multi-Objective Particle Swarm Optimization for Site Selection and Policy Subsidy Maximization of Foreign Renewable Energy Enterprises in the United States," *Artif. Intell. Mach. Learn. Rev.*, vol. 7, no. 2, pp. 54–69, 2026.
38. X. Long, J. Hu, and Z. Ling, "A Comparative Analysis of Telemetry-Driven Anomaly Detection Approaches for Dual-Purpose Operational and Security Optimization in Edge Computing Infrastructure," *J. Comput. Innov. Appl.*, vol. 4, no. 1, pp. 79–88, 2026.
39. T. Dettmers, A. Pagnoni, A. Holtzman, and L. Zettlemoyer, "QLoRA: Efficient finetuning of quantized LLMs," in *Advances in Neural Information Processing Systems (NeurIPS)*, 2023.
40. Y. Li, Y. Yu, C. Liang, P. He, N. Karampatziakis, W. Chen, and T. Zhao, "LoftQ: LoRA-fine-tuning-aware quantization for large language models," in *International Conference on Learning Representations (ICLR)*, 2024.
41. H. Cao, J. Hu, and C. Luo, "Behavioural Feature Analysis for Anomalous Click Detection in Mobile Advertising Environments: Toward In-App Browser-Specific Detection," *J. Comput. Innov. Appl.*, vol. 3, no. 2, pp. 96–105, 2025.
42. Y. Li, F. Zhao, and J. Hu, "Identifying Cross-Market Risk Contagion Amplifiers via Graph Attention Networks: Empirical Evidence from US Financial Stress Periods," *J. Comput. Innov. Appl.*, vol. 4, no. 1, pp. 164–175, 2026.
43. J. Hu and X. Long, "Graph Learning-Based Behavioral Detection for Software Supply Chain Attacks," *J. Adv. Comput. Syst.*, vol. 4, no. 4, pp. 49–60, 2024.
44. Y. Chen and J. Hu, "Graph Neural Network-Based Cascading Disruption Path Identification in Multi-Tier Rare Earth Processing Networks," *J. Global Eng. Rev.*, vol. 4, no. 1, pp. 99–112, 2026.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of Publisher and/or the editor(s). Publisher and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.