

Knowledge Distillation for Efficient Zen Models: Semantic-Preserving Multi-Teacher Distillation

Technical Report v2025.07

Antje Worring, Zach Kelling
Zen LM Research Team
research@zenlm.org

July 2025

Abstract

Knowledge distillation compresses large “teacher” models into smaller, faster “student” models while retaining as much task performance as possible. We present a comprehensive distillation framework for the Qwen3-based Zen models (dense 0.6B/4B/8B/32B and the Qwen3-30B-A3B mixture-of-experts variant, all Apache-2.0), covering offline distillation, online distillation, progressive distillation, and task-specific distillation. Our central innovation is *semantic-preserving distillation (SPD)*, which augments the standard KL-divergence objective with a semantic alignment loss that matches student and teacher representations in a shared anchor embedding space. We further introduce *multi-teacher distillation*, where multiple teacher models of different specializations contribute to a single student. SPD closes a substantially larger fraction of the teacher–student performance gap than standard KL-divergence distillation, while producing students that are several times faster at inference. We present scaling laws for distillation efficiency as a function of student size and teacher size.

Contents

1	Introduction	3
1.1	Distillation Taxonomy	3
2	Semantic-Preserving Distillation	3
2.1	Motivation	3
2.2	SPD Loss	3
2.3	Layer Matching Strategy	4
3	Multi-Teacher Distillation	4
3.1	Motivation	4
3.2	Multi-Teacher Loss	4
4	Distillation Variants	4
4.1	Offline Distillation	4
4.2	Online Distillation	4
4.3	Progressive Distillation	5
4.4	Task-Specific Distillation	5
5	Experiments	5
5.1	Teacher-Student Performance Gaps	5
5.2	Effect of Distillation Method on Student Quality	5

5.3	Distillation Efficiency	5
5.4	Scaling Laws for Distillation	5
5.5	Layer Matching Ablation	6
6	Task-Specific Distillation Results	6
7	Conclusion	6

1 Introduction

Deploying 32B-parameter models at scale demands significant compute: a single GPU serves far fewer tokens/second for a 32B model in FP16 than for a 4B or 8B model. For cost-sensitive applications (consumer products, edge devices, batch processing pipelines), a well-distilled 4B or 8B student can provide near-teacher performance at a fraction of the inference cost.

Knowledge distillation [1] trains a student model S_θ to mimic a teacher model T by minimizing the KL divergence between their output distributions:

$$\mathcal{L}_{\text{KD}} = \text{KL}(p_T(\cdot | x) \| p_S(\cdot | x)) \quad (1)$$

Standard KD operates at the token probability level, transferring the teacher’s soft output distributions. This captures distributional knowledge but ignores the rich internal representations that underlie the teacher’s competence. SPD augments KD with an intermediate representation alignment loss in the semantic anchor space, closing an additional portion of the teacher-student gap beyond what distributional KD alone achieves.

1.1 Distillation Taxonomy

We distinguish four distillation regimes:

- **Offline distillation:** teacher outputs precomputed and stored; student trains on static soft labels. Highest throughput, lowest cost.
- **Online distillation:** teacher runs synchronously during student training, providing dynamic soft labels. Higher quality, higher compute cost.
- **Progressive distillation:** iterative chain $T \rightarrow S_1 \rightarrow S_2 \rightarrow \dots$, each stage compressing by a fixed factor. Enables extreme compression ratios.
- **Task-specific distillation:** distillation on domain-specific data with task-adapted teacher soft labels. Maximizes student performance on target tasks.

2 Semantic-Preserving Distillation

2.1 Motivation

Standard KD transfers distributional knowledge from teacher to student but does not explicitly align their internal representations. A student may produce the correct output distribution through a qualitatively different computational path than the teacher, which can lead to brittleness on out-of-distribution inputs. SPD adds a representation alignment objective that encourages the student to develop semantically similar intermediate representations.

2.2 SPD Loss

Let $h_T^\ell(x) \in \mathbb{R}^{d_T}$ and $h_S^\ell(x) \in \mathbb{R}^{d_S}$ denote the hidden states of teacher and student at layer ℓ (matched by relative depth). SPD projects both to the shared anchor embedding space:

$$\tilde{h}_T^\ell = \mathbf{P}_T h_T^\ell(x), \quad \tilde{h}_S^\ell = \mathbf{P}_S h_S^\ell(x) \quad (2)$$

where $\mathbf{P}_T \in \mathbb{R}^{d_A \times d_T}$ and $\mathbf{P}_S \in \mathbb{R}^{d_A \times d_S}$ are alignment projection matrices learned jointly with the student. The SPD alignment loss is:

$$\mathcal{L}_{\text{align}} = \frac{1}{L} \sum_{\ell=1}^L (1 - \cos(\tilde{h}_T^\ell, \tilde{h}_S^\ell)) \quad (3)$$

The full SPD training objective is:

$$\mathcal{L}_{\text{SPD}} = (1 - \alpha)\mathcal{L}_{\text{CE}}(y, p_S) + \alpha\mathcal{L}_{\text{KD}}(p_T, p_S) + \beta\mathcal{L}_{\text{align}} \quad (4)$$

where α controls the blend of ground truth and teacher labels, and β controls semantic alignment strength. Default: $\alpha = 0.5$, $\beta = 0.1$.

2.3 Layer Matching Strategy

Teacher and student have different depths ($L_T > L_S$). We match layers by relative position: student layer ℓ corresponds to teacher layer $\lfloor \ell \cdot L_T / L_S \rfloor$. We also experiment with task-specific matching, where alignment is only applied at layers that correlate most strongly with task-relevant representations (selected by probing).

3 Multi-Teacher Distillation

3.1 Motivation

The Zen family includes specialized variants: a code-focused model, a mathematics-tuned model, and a general-purpose model (each a fine-tune of a Qwen3-based Zen checkpoint). A single student trained on the general teacher may underperform on code and mathematics. Multi-teacher distillation aggregates knowledge from multiple specialized teachers.

3.2 Multi-Teacher Loss

Given K teacher models $\{T_k\}_{k=1}^K$, the multi-teacher distillation loss is:

$$\mathcal{L}_{\text{MT}} = (1 - \alpha)\mathcal{L}_{\text{CE}} + \alpha \sum_{k=1}^K w_k(x) \cdot \mathcal{L}_{\text{KD}}(p_{T_k}, p_S) + \beta \sum_{k=1}^K w_k(x) \cdot \mathcal{L}_{\text{align}}^{(k)} \quad (5)$$

where $w_k(x) \in [0, 1]$ is a domain relevance weight for teacher k on input x :

$$w_k(x) = \text{softmax}\left(\frac{z_k(x)}{\tau}\right), \quad z_k(x) = \cos(e(x), c_k) \quad (6)$$

with $e(x)$ the input embedding and c_k the centroid of teacher k 's training domain in embedding space. This implements a soft routing: for a code prompt, the code teacher receives higher weight; for mathematics, the math teacher is weighted more.

4 Distillation Variants

4.1 Offline Distillation

In offline distillation, teacher outputs $\{p_T(y | x)\}$ are precomputed for the entire training corpus and stored as soft labels. The student trains on these stored distributions using $\mathcal{L}_{\text{SPD}}$ with the alignment term approximated from stored teacher hidden states. Offline distillation is $4\times$ cheaper per student training step than online, at a cost of 5–8% task performance.

4.2 Online Distillation

In online distillation, the teacher model runs synchronously during student training, providing fresh soft labels and hidden states for each training batch. This enables the student to learn from teacher outputs on the exact sequence it is training on, including any data augmentation applied online.

4.3 Progressive Distillation

Progressive distillation [1, 2] applies a chain: $32\text{B} \rightarrow 8\text{B} \rightarrow 4\text{B} \rightarrow 0.6\text{B}$, where each stage uses the previous stage’s output as the teacher. This is more effective than direct compression from 32B to 0.6B because each step makes a manageable capacity reduction.

4.4 Task-Specific Distillation

For deployment scenarios where only a specific task matters (e.g., code generation, SQL synthesis), we distill on a domain-specific corpus with the corresponding specialized teacher. Task-specific students achieve near-specialist performance in their domain while fitting in an 8B parameter budget.

5 Experiments

5.1 Teacher-Student Performance Gaps

Table 1: Performance gap closure (%) for different distillation methods. Teacher = Zen-32B (Qwen3-32B), student = Zen-8B (Qwen3-8B). Gap closure = (student - base) / (teacher - base) $\times$ 100, i.e. the fraction of the 8B $\rightarrow$ 32B benchmark gap recovered by distillation. Values are relative to the undistilled 8B base and compare distillation *methods*; they are not standalone capability claims.

Method	MMLU	GSM8K	HumanEval	Avg.
No distillation (8B base)	0%	0%	0%	0%
Standard KD	58%	61%	64%	61%
SPD (ours)	72%	76%	73%	74%
SPD + multi-teacher	74%	81%	79%	78%

5.2 Effect of Distillation Method on Student Quality

Holding the student fixed at Zen-8B (Qwen3-8B) and varying the distillation method, we observe a consistent ordering on MMLU, GSM8K, and HumanEval: the undistilled base is weakest, standard KD improves it, SPD improves further, and SPD with multi-teacher aggregation is strongest, approaching the Zen-32B teacher. The same method ordering holds for the progressively-distilled 0.6B student, which recovers a smaller but still substantial fraction of teacher quality at a much smaller parameter budget.

5.3 Distillation Efficiency

The motivation for distillation is the inference-cost/quality trade-off: the smaller distilled students run several times faster than the 32B teacher while retaining most of its benchmark quality. The Zen-8B SPD+MT student offers a favorable operating point (near-teacher MMLU at a multiple of the teacher’s throughput), and the progressively distilled 0.6B student offers the highest throughput for the most cost-sensitive deployments, at a larger quality reduction. Exact throughput depends on hardware and serving configuration.

5.4 Scaling Laws for Distillation

We observe that distillation efficiency (performance per parameter) follows a power law in the student-to-teacher size ratio:

$$\text{Gap closure}(s) = 1 - C \cdot \left(\frac{s}{t}\right)^{-\delta} \quad (7)$$

where s is the student parameter count and t is the teacher parameter count. Fitting this form across the Qwen3-based Zen student sizes (0.6B, 4B, 8B) distilled from the Zen-32B teacher, gap closure increases smoothly and concavely with s/t : larger students recover a larger fraction of the teacher’s quality, with diminishing marginal returns as the student approaches the teacher’s size. The fitted curve lets practitioners estimate the gap closure achievable at a target student size before committing to a distillation run.

5.5 Layer Matching Ablation

Ablating the SPD layer-matching strategy on the Zen-8B student, adding any representation alignment improves over KD alone, and the choice of *which* layers to align matters: probing-based matching (aligning the few teacher layers most predictive of task accuracy) performs best, uniform matching (every few layers) is close behind, and aligning all layers performs slightly worse than the targeted strategies — indicating that a small number of well-chosen alignment points captures most of the benefit.

6 Task-Specific Distillation Results

Comparing task-specific Zen-8B (Qwen3-8B) students against the general SPD+MT Zen-8B student on four domain benchmarks (HumanEval, GSM8K, LegalBench, MedQA), each domain-specialized student is the strongest on its own target benchmark: the code-specific student leads on HumanEval, the math-specific student on GSM8K, the legal-specific student on LegalBench, and the medical-specific student on MedQA. Specialization trades some general breadth for a clear gain in the target domain, recovering close to specialist-level teacher performance within an 8B parameter budget.

7 Conclusion

Semantic-preserving distillation (SPD) with multi-teacher aggregation closes a substantially larger fraction of the 32B-to-8B performance gap than standard KD (78% vs. 61% of the gap in our experiments). Progressive distillation enables a large throughput increase at the 0.6B scale. The distillation scaling law (Equation 7) predicts gap closure as a function of s/t , enabling practitioners to select the optimal student size for their compute budget. Task-specific distillation recovers close to specialist-level teacher performance in target domains, providing an efficient deployment path for specialized applications. All teachers and students are Qwen3-based Zen models (Apache-2.0).

References

- [1] G. Hinton, O. Vinyals, J. Dean. *Distilling the Knowledge in a Neural Network*. arXiv:1503.02531, 2015.
- [2] T. Salimans, J. Ho. *Progressive Distillation for Fast Sampling of Diffusion Models*. ICLR, 2022.
- [3] W. Park, D. Kim, Y. Lu, M. Cho. *Relational Knowledge Distillation*. CVPR, 2019.
- [4] A. Romero, N. Ballas, S.E. Kahou, et al. *FitNets: Hints for Thin Deep Nets*. ICLR, 2015.