

REPL: Pseudo-label Refinement for Semi-supervised LiDAR Semantic Segmentation

Donghyeon Kwon^{1*}, Taegyu Park^{2*}, and Suha Kwak^{2,3}

¹GENGENAI, Republic of Korea ²Dept. of CSE, POSTECH, Republic of Korea

³Graduate School of AI, POSTECH, Republic of Korea

¹kinux98@gmail.com, ²{taegyu.park,suha.kwak}@postech.ac.kr

Abstract. Semi-supervised learning for LiDAR semantic segmentation often suffers from error propagation and confirmation bias caused by noisy pseudo-labels. To tackle this chronic issue, we introduce REPL, a novel framework that enhances pseudo-label quality by identifying and correcting potential errors in pseudo-labels through masked reconstruction, along with a dedicated training strategy. We also provide a theoretical analysis demonstrating the condition under which the pseudo-label refinement is beneficial, and empirically confirm that the condition is mild and clearly met by REPL. Extensive evaluations on the nuScenes-lidarseg and SemanticKITTI datasets show that REPL improves pseudo-label quality substantially, and in consequence, achieves the state of the art in semi-supervised LiDAR semantic segmentation.

Keywords: Semi-supervised learning · LiDAR semantic segmentation

1 Introduction

Outdoor LiDAR semantic segmentation, the task of assigning semantic labels to every point in outdoor 3D scenes, plays a crucial role in diverse applications such as autonomous driving [7, 29, 31] and robotics [33, 37]. Recent progress in this field has been largely driven by supervised learning on large-scale point cloud datasets [1, 6]. However, collecting dense annotations for 3D point clouds is prohibitively costly and time-consuming, which limits the scale and class diversity of training data. To alleviate this bottleneck, a large body of research has explored data-efficient training paradigms such as semi-supervised learning [10, 13, 19–21, 24], weakly supervised learning [22, 35], and unsupervised learning [27, 43]. In this work, we study semi-supervised learning for LiDAR semantic segmentation, where only a subset of 3D scenes for training is manually labeled and the remainder is unlabeled.

At the core of semi-supervised learning lies the challenge of leveraging unlabeled data effectively for training. To this end, existing methods for LiDAR semantic segmentation commonly leverage consistency regularization [13, 21, 24] and contrastive learning [10, 19, 24]. Consistency regularization encourages stable and invariant predictions by enforcing similar outputs under different perturbations of the same input, while contrastive learning promotes feature representations that bring samples of the same pseudo-labels closer and push those of

* Equal contribution.

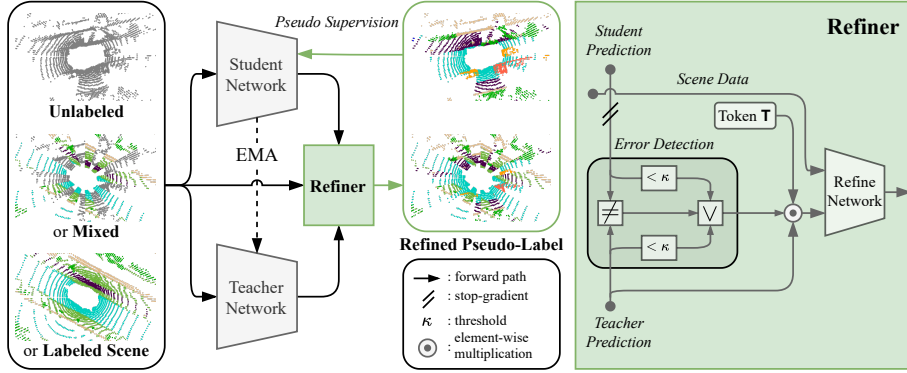


Fig. 1: Overview of REPL. The teacher generates predictions for unlabeled LiDAR scenes, which are used as pseudo-labels for the student, and is updated via exponential moving average (EMA) of the student. The pseudo-label refiner detects erroneous pseudo-labels by confidence-based agreement between the teacher and student, and then corrects them through masked reconstruction with learnable tokens. The final refined pseudo-labels combine reliable teacher predictions with reconstructed ones, yielding improved supervision for semi-supervised learning of the student.

different pseudo-labels farther apart. Although these approaches often yield substantial gains, they have a fundamental vulnerability, confirmation bias towards erroneous pseudo-labels [15, 41], since they blindly exploit the model’s predictions as pseudo-labels for training the model with unlabeled data; this could cause performance to deteriorate as training progresses.

To handle noisy pseudo-labels, recent studies have proposed confidence-based filtering [13, 20, 21], which discards predictions with low confidence under the assumption that such labels are more error-prone. Loss reweighting methods [19, 24] take a complementary approach by retaining all pseudo-labeled samples but adjusting their contribution through confidence weighted loss scaling. While both strategies assuage the impact of unreliable labels, they remain post-hoc in nature as they adjust sample utilization only after pseudo-labels have been assigned, rather than improving their quality at the point of generation.

To tackle this issue, we introduce a novel pseudo-label refinement framework, named REPL (**R**efinement of **P**seudo-**L**abels), which is illustrated in Fig. 1. REPL integrates two key components: teacher-student networks for LiDAR semantic segmentation [34] and a pseudo-label refiner that identifies errors on pseudo-labels and corrects them. The teacher network offers initial predictions for LiDAR scenes, which are in turn used as pseudo-labels for training the student with unlabeled data; the teacher is then updated by exponential moving average (EMA) of the student. The pseudo-label refiner identifies potentially unreliable regions in the pseudo-labels and reconstructs them into cleaner supervisory signals for the student. Specifically, the refiner operates through two stages. First, it estimates potentially erroneous pseudo-labels using a simple confidence-based agreement between teacher and student predictions. Second, it reconstructs these pseudo-labels through a process inspired by masked autoen-

coders [8]; in this step, unreliable areas are replaced with mask tokens, and the refiner reconstructs them to produce refined predictions. Even with a simple error estimation strategy, REPL achieves significant performance improvements, while remaining extensible to more sophisticated error detection methods.

The success of our method heavily depends on the ability of the pseudo-label refiner. However, training the refiner in the semi-supervised learning setting is challenging due to the scarce supervision: the supervision derived from the discrepancies between predicted and ground-truth labels is available only for a small subset of training data. We address this issue in two ways. First, during training the refiner, we apply additional random masking to make its reconstruction more challenging. This forces the refiner to develop a better contextual understanding rather than simply memorizing patterns. Second, we mix 3D scenes, augmenting labeled scenes with unlabeled ones before feeding them to the segmentation networks. This produces strongly augmented segmentation predictions with higher and diverse error rates, and enables the refiner to partially experience prediction errors for unlabeled scenes. We also provide a theoretical analysis demonstrating the condition under which the refinement is beneficial, and empirically confirm that the condition is mild and clearly met by REPL.

Following the established practice [13, 21, 24], we evaluated our method on the nuScenes-lidarseg [6] and SemanticKITTI [1] benchmarks while varying the ratio of labeled data, where it demonstrated significant performance improvements over the supervised learning baseline and outperformed latest methods for semi-supervised learning. In summary, our contribution is three-fold:

- We propose a semi-supervised LiDAR semantic segmentation framework, dubbed REPL, that refines pseudo-labels via error estimation and masked reconstruction.
- We provide a theoretical analysis establishing the condition under which pseudo-label refinement improves upon the teacher-only baseline.
- REPL achieved the state of the art on the two public benchmarks.

2 Related Work

LiDAR Semantic Segmentation. LiDAR semantic segmentation assigns semantic labels to each point in large-scale point clouds [26, 44]. Early work transformed raw 3D scans into 2D representations such as range images or bird’s-eye-view maps to reuse established 2D convolutional backbones. RangeNet++ [26] and SalsaNext [5] demonstrated that simple projection can yield competitive results. Another research stream discretized 3D space into regular or cylindrical grids, as seen in PolarNet [42], Cylinder3D [44], and sparse convolutional frameworks like MinkowskiNet [4]. Meanwhile, methods directly processing raw points gained traction, from PointNet [30] to RandLA-Net [9] and stratified architectures [11, 18], capturing both local details and long-range context. Despite performance gains, these methods require dense manual annotations, which are costly and not scalable.

Semi-supervised LiDAR Semantic Segmentation. To address this annotation bottleneck, semi-supervised learning leverages limited labeled data along with abundant unlabeled point clouds [13, 24]. Earlier work improved pseudo-label reliability: GPC [10] used confidence thresholds to reduce error propagation, LaserMix [13] exploited spatial priors of LiDAR beams, enforcing prediction consistency through beam mixing across scans, and Lim3D [20] employed a memory bank for contrastive learning to alleviate class imbalance. Recent methods introduced richer constraints: DDSemi [19] employed density-guided contrastive learning with dual-space hardness sampling for sparse regions, AIScene [21] addressed intra-scene inconsistency through patch-based mixing at scene and instance levels, and IT2 [24] introduced consistency learning across peer LiDAR representations, treating representation differences as perturbations. These approaches remain post-hoc, adjusting pseudo-label usage rather than improving their intrinsic quality. REPL directly enhances pseudo-labels by correcting erroneous pseudo-labels, providing improved supervision.

3 Method

A chronic issue in semi-supervised learning is the confirmation bias towards errors of pseudo-labels, which has often been handled by post-hoc strategies such as confidence filtering or loss reweighting. In this work, we explore a different direction: improving the quality of unreliable pseudo-labels by refinement instead of discarding them.

To this end, we propose REPL, a refinement-based semi-supervised learning framework for LiDAR semantic segmentation. REPL consists of two modules: teacher-student segmentation networks and a pseudo-label refiner. The teacher network generates pseudo-labels for unlabeled data and is updated by exponential moving average of the student parameters, while the student is trained with both labeled and pseudo-labeled data. To handle unreliable pseudo-labels, the refiner identifies uncertain voxels and corrects their pseudo-labels through masked reconstruction. The final pseudo-labels combine reliable teacher predictions with the refined output on uncertain regions.

The remainder of this section first elaborates on three training steps of REPL: (1) training the student network using the standard segmentation objectives on labeled data (Sec. 3.2), (2) training the pseudo-label refiner on both labeled and unlabeled data (Sec. 3.3), and (3) semi-supervised learning of the student network with the pseudo-labels improved by the refiner (Sec. 3.4). Then a theoretical analysis is presented to establish when the refinement is beneficial and to validate that REPL operates within this regime (Sec. 3.5).

3.1 Preliminaries

We consider training a segmentation network on a small labeled dataset and a large unlabeled dataset. Each point cloud is voxelized into a regular grid $X_i \in \mathbb{R}^{C \times H \times W \times L}$, where C denotes the number of feature channels and $H \times W \times L$

is the grid size. The labels are converted into voxel-wise one-hot tensors $Y_i \in \{0, 1\}^{K \times H \times W \times L}$, where K is the number of classes. The labeled and unlabeled datasets are denoted as $D_L = \{(X_i, Y_i)\}_{i=1}^{N_l}$ and $D_U = \{X_j\}_{j=1}^{N_u}$, respectively.

3.2 Supervised Learning with Labeled Data

Following prior work [13, 20, 24, 44], we train the student network $f(\cdot)$ on labeled data with two complementary segmentation losses: cross-entropy for voxel-wise classification and Lovász-Softmax [2] for direct IoU optimization. We denote the set of voxel grid indices as $\Omega = \{1, \dots, H\} \times \{1, \dots, W\} \times \{1, \dots, L\}$. For input X_i , the network outputs predictions $P_i = f(X_i) \in \mathbb{R}^{K \times H \times W \times L}$, where $[P_i]_{k,\omega}$ denotes the probability of class k at voxel $\omega \in \Omega$. The ground-truth label is denoted by $[Y_i]_{k,\omega} \in \{0, 1\}$.

The voxel-wise cross-entropy loss is defined as

$$\mathcal{L}_{ce}(P_i, Y_i) = -\frac{1}{|\Omega|} \sum_{\omega \in \Omega} \sum_{k=1}^K [Y_i]_{k,\omega} \log[P_i]_{k,\omega}. \quad (1)$$

For Lovász-Softmax, we first define the one-versus-rest error for each class k :

$$\mathbf{e}_i^{(k)}(\omega) = (1 - [P_i]_{k,\omega}) \cdot \mathbf{1}_{\{[Y_i]_{k,\omega}=1\}} + [P_i]_{k,\omega} \cdot \mathbf{1}_{\{[Y_i]_{k,\omega} \neq 1\}}. \quad (2)$$

The multiclass Lovász extension is then given by

$$\mathcal{L}_{ls}(P_i, Y_i) = \frac{1}{|C'_i|} \sum_{k \in C'_i} \bar{\Delta}_{\text{Jacc}}(\mathbf{e}_i^{(k)}), \quad (3)$$

where $C'_i = \{k \mid \sum_{\omega \in \Omega} [Y_i]_{k,\omega} > 0\}$ and $\bar{\Delta}_{\text{Jacc}}$ denotes the Lovász extension [2] of the Jaccard loss. The supervised learning objective combines both terms:

$$\mathcal{L}_{ssup} = \frac{1}{N_l} \sum_{i=1}^{N_l} \left\{ \mathcal{L}_{ce}(P_i, Y_i) + \lambda_{ls} \mathcal{L}_{ls}(P_i, Y_i) \right\}. \quad (4)$$

3.3 Training Pseudo-label Refiner

We employ the teacher-student framework where the teacher f^τ network outputs predictions for unlabeled data, $Q = f^\tau(X) \in \mathbb{R}^{K \times H \times W \times L}$, as their pseudo-labels and is updated by exponential moving average of the student network f [34]. The pseudo-label refiner network $g(\cdot)$ aims to identify unreliable teacher's predictions and refine them to improve the quality of the pseudo-labels.

Unreliable Voxel Identification. Voxels with unreliable pseudo-labels are identified by the agreement between the student's and teacher's predictions along with their confidence levels. For each voxel, we compute the confidence score as the maximum prediction probability across all classes for each of the two

networks. We also establish adaptive confidence thresholds using the $(100 - \kappa)$ -th percentile of all confidence scores within each scene, where κ controls the strictness of the confidence requirement. A voxel is considered reliable only if all the following conditions hold: (1) the student and teacher predict the same class, (2) the student’s confidence exceeds its adaptive threshold, and (3) the teacher’s confidence exceeds its adaptive threshold. All other voxels are treated as unreliable and marked for refinement. The error candidate mask is then given by $M = \mathbf{1} - \mathbf{1}_{\text{rel}} \in \{0, 1\}^{H \times W \times L}$, where $\mathbf{1}_{\text{rel}}$ denotes the binary mask indicating voxels that satisfy all the three conditions above. To prevent the refiner from overfitting to error-prone regions and develop a better contextual understanding rather than simply memorizing patterns, we also apply additional random masking $R \sim \text{Bernoulli}(\sigma)$ and define the final mask as $\bar{M} = M \vee R$.

Masked Reconstruction. Once unreliable voxels are identified through the error candidate mask, the next step is to correct their predictions. The core idea is to mask out these uncertain predictions and reconstruct them through a masked reconstruction process [8]. Specifically, we form masked predictions $\bar{Q}$ by substituting the unreliable predictions in Q with a learnable mask token $T \in \mathbb{R}^K$, which serves as a placeholder indicating that the original prediction at its position has been masked and needs to be reconstructed by the refiner: $\bar{Q} = (\mathbf{1} - \bar{M}) \odot Q + \bar{M} \odot T$, where $\odot$ denotes element-wise multiplication. Analogous to the mask token in masked autoencoders [8], T is jointly optimized during training, allowing the network to learn an effective representation for masked regions. The refiner then takes channel-wise concatenation of X and $\bar{Q}$ as input and outputs refined predictions: $\hat{Q} = g(X, \bar{Q})$.

Training Refiner on Labeled Data. On labeled data, the refiner is trained to reconstruct ground-truth labels of the masked predictions. The loss for training the refiner on labeled data, $\mathcal{L}_{\text{rsup}}$, is the same as the supervised learning objective $\mathcal{L}_{\text{ssup}}$ in Eq. (4), except that it is applied to only the masked regions (*i.e.*, $\forall \omega \in \Omega$ s.t. $[\bar{M}_i]_{\omega} = 1$) of the refined predictions $\hat{Q}_i$.

Training Refiner on Unlabeled Data. On unlabeled data, the refiner is further trained with a negative learning signal [12]. Rather than enforcing hard pseudo-labels, we suppress predictions on implausible classes by taking the teacher’s top- k predictions as plausible candidates and encouraging the refiner to assign low probability to the remaining classes. This negative learning strategy offers reliable supervisory signals even when pseudo-supervision may not be sufficiently accurate on unlabeled data. Formally, the negative learning loss is defined as the average cross-entropy penalty over unlabeled scenes in D_U :

$$\mathcal{L}_{\text{runl}} = \frac{1}{N_u} \sum_{j=1}^{N_u} \frac{1}{|\Omega|} \sum_{\omega \in \Omega} \frac{1}{|\mathcal{N}_j(\omega)|} \sum_{k \in \mathcal{N}_j(\omega)} \{ -\log(1 - [\hat{Q}_j]_{k,\omega}) \}, \quad (5)$$

where $\mathcal{N}_j(\omega)$ denotes the set of implausible classes for voxel ω .

Training Refiner on Mix of Labeled and Unlabeled Data. Training the refiner benefits from diverse prediction errors of the teacher, but the limited labeled data restrict such diversity. To strengthen the supervision, we mix labeled and unlabeled scenes so that the refiner reconstructs labels under richer variability in distance, geometry, and density. This produces augmented predictions with higher error rates, and allows the refiner to partially experience errors on unlabeled data. We adopt LaserMix [13] with a single inclination plane to fuse labeled and unlabeled scans at ratio $r \in (0, 1)$. Let a mix of labeled and unlabeled scenes be indexed by $m \in \{1, \dots, N_m\}$. Given a labeled scene (X_i, Y_i) and an unlabeled one X_j for the m -th mix, LaserMix generates a selector mask $S_m \in \{0, 1\}^{H \times W \times L}$ and produces the mixed input and output $(X_m, Y_m) = (S_m \odot X_i + (\mathbf{1} - S_m) \odot X_j, S_m \odot Y_i)$. We then compute student and teacher predictions on X_m , $P_m = f(X_m)$ and $Q_m = f^\tau(X_m)$, respectively. Following the unreliable prediction identification procedure in Sec. 3.3, we construct the error mask $\tilde{M}_m$ restricted to the labeled prediction. The supervised learning losses are then applied to voxels of the labeled scene marked as unreliable, *i.e.*, $\forall \omega \in \Omega$ s.t. $[S_m]_\omega = [\tilde{M}_m]_\omega = 1$:

$$\mathcal{L}_{\text{rmix}} = \frac{1}{N_m} \sum_{m=1}^{N_m} \left\{ \mathcal{L}_{\text{ce}}(\hat{Q}_m, Y_m) + \lambda_{\text{ls}} \mathcal{L}_{\text{ls}}(\hat{Q}_m, Y_m) \right\}. \quad (6)$$

Total Training Objective for Refiner. Each iteration optimizes the pseudo-label refiner with the summation of the three losses, $\mathcal{L}_{\text{rsup}} + \mathcal{L}_{\text{runl}} + \mathcal{L}_{\text{rmix}}$, without balancing hyper-parameters.

3.4 Semi-supervised Learning with Pseudo-label Refiner

The student network leverages the refined pseudo-labels from the refiner to improve the quality of supervision when learning with unlabeled data.

Pseudo-label Refinement. For an unlabeled input $X_j \in D_U$, we obtain predictions from both student and teacher (P_j, Q_j) and construct the error mask M_j as before but without random masking. The refined pseudo-label $\tilde{Y}_j$ is then generated voxel-wise: reliable voxels follow the teacher’s prediction, while unreliable ones are replaced with the refiner’s output. Specifically, we form masked predictions $\bar{Q}_j$ by substituting the unreliable predictions in Q_j with a learnable mask token T : $\bar{Q}_j = (\mathbf{1} - M_j) \odot Q_j + M_j \odot T$ and obtain refined predictions $\hat{Q}_j = g(X_j, \bar{Q}_j)$, which are combined with teacher predictions on reliable voxels. The refined pseudo-label $\tilde{Y}_j$ is then obtained by choosing the best class per voxel.

Semi-supervised Learning of Student. The student network is trained with three objectives, the supervised learning objective applied to labeled data (*i.e.*, $\mathcal{L}_{\text{ssup}}$ in Sec. 3.2), and two additional objectives for learning with unlabeled data. For the segmentation losses on unlabeled data, as the refined pseudo-labels are used directly without additional filtering, we adopt the symmetric cross-entropy [38]

instead of the standard cross-entropy to ensure more stable training. The symmetric formulation is more robust to potential noise in pseudo-labels as it penalizes over-confident predictions and provides regularization through bidirectional loss computation. Together with the Lovász-Softmax loss, this objective is applied to all voxels using the refined pseudo-labels:

$$\mathcal{L}_{\text{sunl}} = \frac{1}{N_u} \sum_{j=1}^{N_u} \left\{ \frac{1}{2} \left(\mathcal{L}_{\text{ce}}(P_j, \tilde{Y}_j) + \mathcal{L}_{\text{ce}}(\tilde{Y}_j, P_j) \right) + \lambda_{\text{ls}} \mathcal{L}_{\text{ls}}(P_j, \tilde{Y}_j) \right\}. \quad (7)$$

In addition to training with pseudo-labels, we augment training by mixing a labeled scene (X_i, Y_i) with an unlabeled scene X_j paired with pseudo-labels $\tilde{Y}_j$. We employ the LaserMix operator with a single inclination plane. This strategy combines clean supervision from labeled data with broader coverage from pseudo-labels, exposing the student to reliable signals and diverse structures. Given the selector mask S_m by LaserMix, the mixed input and target are defined by combining ground-truth labels in the labeled region and refined pseudo-labels in the unlabeled region:

$$(X_m, \tilde{Y}_m) = (S_m \odot X_i + (\mathbf{1} - S_m) \odot X_j, S_m \odot Y_i + (\mathbf{1} - S_m) \odot \tilde{Y}_j). \quad (8)$$

For the mixed input X_m , the student network produces predictions $P_m = f(X_m)$. The loss for the mixed sample is then computed by

$$\mathcal{L}_{\text{smix}} = \frac{1}{N_m} \sum_{m=1}^{N_m} \left\{ \frac{1}{2} \left(\mathcal{L}_{\text{ce}}(P_m, \tilde{Y}_m) + \mathcal{L}_{\text{ce}}(\tilde{Y}_m, P_m) \right) + \lambda_{\text{ls}} \mathcal{L}_{\text{ls}}(P_m, \tilde{Y}_m) \right\}. \quad (9)$$

Total Training Objective for Student. In each iteration, the student network is optimized with the summation of the three losses, $\mathcal{L}_{\text{ssup}} + \mathcal{L}_{\text{sunl}} + \mathcal{L}_{\text{smix}}$, with no balancing hyper-parameter. The student network is optimized jointly with the pseudo-label refiner. We stop gradients between their optimization paths to prevent interference.

3.5 Theoretical Analysis

This section rigorously analyzes if the pseudo-label refinement is truly helpful. We first show whether the pseudo-label refinement is easier than generating high-quality pseudo-labels from scratch in Proposition 1.

Proposition 1 (Refinement Difficulty). *Consider two segmentation tasks, the original task $Z : X \rightarrow Y$ and the refinement task $Z' : (X, T) \rightarrow Y$, where X denotes input 3D LiDAR point data, Y denotes segmentation labels, and T represents additional information such as teacher predictions $f(X)$. The difficulties of the two tasks $D(Z)$ and $D(Z')$ hold the following inequality:*

$$D(Z') = H(Y \mid X, T) \leq H(Y \mid X) = D(Z). \quad (10)$$

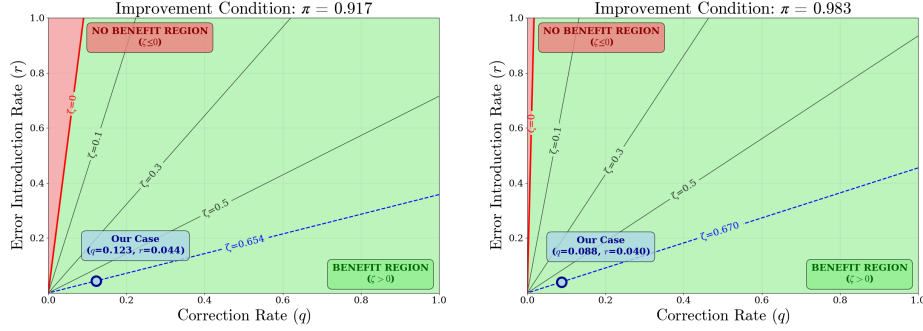


Fig. 2: Visualization of the improvement condition from Eq. (11) on the SemanticKITTI dataset given π . The two values of π were derived during the actual experiments on the dataset: 1% labeled data ($\pi = 0.917$) and 50% labeled data ($\pi = 0.983$). Green areas indicate combinations of q and r where the refinement yields a net benefit ($\zeta > 0$), while red areas show such combinations leading to detriment ($\zeta \leq 0$). The results suggest that the refinement remains beneficial across a broad range of q and r , and that REPL clearly helps improve the quality of pseudo-labels.

This result implies that the refinement may have potential for improving pseudo-label quality. Next, we derive a practical condition required for net performance gains in Proposition 2.

Proposition 2 (Improvement Condition). *For the j -th scene, an error-candidate mask M_j divides the voxel grid into an unreliable region E_j and a reliable region C_j . Let π_j denote the precision of this mask, which is the fraction of voxels misclassified by the teacher in E_j . The refiner operates on E_j , where q_j and r_j denote the rates at which it corrects misclassifications and incorrectly changes correct predictions, respectively. Then, if $q_j + r_j > 0$, the refinement improves the accuracy for scene j if and only if*

$$\zeta_j := \pi_j - \frac{r_j}{q_j + r_j} > 0. \quad (11)$$

This proposition characterizes the trade-off between error correction and error introduction, and its conclusion is a condition that is mild and easily satisfied by REPL even with the simple error estimation method in Sec. 3.4, as will be demonstrated below. The proofs for the two propositions are presented in the supplementary material.

Empirical Analysis on the Improvement Condition. To examine the practical implication of the condition in Eq. (11), we analyze the relationship between the averaged correction rate q and the averaged error introduction rate r using experimental results on the SemanticKITTI dataset. We consider two scenarios: labeling 1% of training data yielding $\pi = 0.917$, and labeling 50% achieving $\pi = 0.983$. Fig. 2 visualizes combinations of q and r where the refinement yields a net benefit ($\zeta > 0$) versus those where it does not ($\zeta \leq 0$). The results suggest that the refinement remains beneficial across a broad range of q and r . For

instance, in the case of $\pi = 0.917$, the refinement remains beneficial as long as the error introduction rate stays below approximately *eleven times* the correction rate ($r < 11.05 \cdot q$), allowing the refinement to be effective across a broad range of performance levels. In both cases, REPL falls within the benefit regions ($\zeta > 0$), showing that it satisfies the theoretical condition for the improvement of pseudo-label quality even with the simple error estimation strategy based on teacher-student agreement.

4 Experiments

4.1 Experimental Setup

Datasets. We trained and evaluated our model on two outdoor LiDAR semantic segmentation benchmarks: nuScenes-lidarseg [6] and SemanticKITTI [1]. nuScenes-lidarseg is a large-scale outdoor LiDAR dataset with point-wise annotations for 16 classes. We adopted the official split of 700 sequences for training and 150 for validation, which resulted in 28,130 training and 6,019 validation point cloud scenes. SemanticKITTI is a LiDAR segmentation benchmark with 19 classes. It consists of 22 sequences, of which 10 were used for training, and 1 for validation, yielding 19,130 training and 4,071 validation scenes.

Network Architecture. Following previous work [13, 21, 24], we used Cylinder3D [44] for both the segmentation models and pseudo-label refiner, fixing the intermediate layer at 16 dimensions as specified by [13] and [24].

Implementation Details. Our method was implemented in PyTorch [28], and trained on 8 NVIDIA RTX 6000 Ada GPUs with AdamW [25] and weight decay of $1e-3$. The batch size was 8 on nuScenes-lidarseg and 4 on SemanticKITTI. The learning rate was $5e-3$ with cosine annealing for both the segmentation network and pseudo-label refiner. The weight update ratio α was 0.994. Following [24], the loss coefficient λ_{ls} was set to 3.0. We set the confidence percentile κ to 40%, k of top- k classes for negative learning to 3, random masking probability σ to 0.15, and mixed data at a mixing ratio r of 0.7 in the selector mask S . Note that all hyper-parameters except the batch size were set identically across the two benchmarks.

4.2 Comparison with State of the Art

REPL was compared with latest methods using Cylinder3D as the backbone, namely AIScene [21], IT2 [24], FrustrumMix [39], and conventional methods employing semi-supervision [3, 13, 14, 34, 45]. For a more comprehensive comparison, we further evaluated against Seal [23], SuperFlow [40], and SLidR [32], which leverage external sources or additional representation learning, and Lim3D [20], which uses a distinct backbone based on Cylinder3D. Table 1 summarizes the results on the validation sets of nuScenes-lidarseg [6] and SemanticKITTI [1] with various ratio of labeled data: 1%, 10%, 20%, and 50%. *Sup-only* denotes the baseline performance of training only with labeled data. On nuScenes-lidarseg, REPL outperformed all competing methods. Compared to IT2, the second-best,

Table 1: Comparison of different semi-supervised learning methods on nuScenes-lidarseg and SemanticKITTI while varying the ratio of labeled data. The best results in each column are shown in bold, and the second best are underlined. Backbones marked with an asterisk (*) indicate that additional representation learning or knowledge distillation from external sources has been applied.

Method	Backbone	nuScenes-lidarseg [6]					SemanticKITTI [1]				
		1%	10%	20%	50%	Avg.	1%	10%	20%	50%	Avg.
<i>Sup-only</i>	Cylinder3D	50.9	65.9	66.6	71.2	63.7	45.4	56.1	57.8	58.7	54.5
Seal [23]	MinkUNet*	45.8	63.0	-	-	-	46.6	-	-	-	-
SuperFlow [40]	MinkUNet*	48.1	64.5	-	-	-	48.4	-	-	-	-
SLiDR [32]	Cylinder3D*	39.0	58.8	-	-	-	44.6	-	-	-	-
Lim3D [20]	LiM3D	-	-	-	-	-	58.4	59.5	63.1	63.6	61.2
MT [34]	Cylinder3D	51.6	66.0	67.1	71.7	64.1	45.4	57.1	59.2	60.0	55.4
CBST [45]	Cylinder3D	53.0	66.5	69.6	71.6	65.2	48.8	58.3	59.4	59.7	56.6
CPS [3]	Cylinder3D	52.9	66.3	70.0	72.5	65.4	46.7	58.7	59.6	60.5	56.4
LaserMix [13]	Cylinder3D	55.3	69.9	71.8	73.2	67.6	50.6	60.0	61.9	62.3	58.7
IT2 [24]	Cylinder3D	57.5	72.1	73.5	74.1	69.3	52.0	61.4	62.1	62.5	59.5
AIScene [21]	Cylinder3D	56.6	70.2	72.8	73.9	68.4	54.5	63.3	63.7	64.3	61.5
FrustrumMix [39]	Cylinder3D	60.0	70.0	72.6	<u>74.1</u>	69.2	<u>55.7</u>	<u>62.5</u>	63.0	64.9	<u>61.5</u>
LaserMix++ [14]	Cylinder3D	58.5	71.1	72.8	74.0	69.1	56.2	62.3	62.9	63.4	61.2
REPL (Ours)	Cylinder3D	60.0	74.4	75.0	75.8	71.3	54.7	<u>62.5</u>	<u>63.2</u>	65.9	61.6

it achieved an average gain of +2.0 in mIoU. REPL also showed strong results on SemanticKITTI, achieving the best in average mIoU. These results are particularly impressive regarding that REPL fixes hyperparameters across both benchmarks, unlike the strong competitors like AIScene [21] and FrustrumMix [39] that tune key hyperparameters per benchmark. Table 2 shows the results on the ScribbleKITTI [36] benchmark, which uses scribble-based weak annotations as labeled supervision. REPL achieved the highest average mIoU and outperformed all competing methods at 1%, 10%, and 50% labeled data ratios. Figure 3 qualitatively compares pseudo-labels before and after the refinement by REPL at the end of training on the unlabeled data of nuScenes-lidarseg.

4.3 In-depth Analysis

This section studies the contribution of each component of REPL, and investigates the aspects of the pseudo-label refinement in details. All experiments were conducted on the validation set, except for the pseudo-label refinement analysis, which used the unlabeled training data.

Ablation Study on Loss Components. We assessed the contribution of individual loss terms by incrementally adding them for the refiner and segmentation network. For the refiner (Table 3), the supervised-only baseline yielded 50.9 mIoU. Adding $\mathcal{L}_{rsup}$ improved performance to 57.2 mIoU, and including $\mathcal{L}_{runl}$ further raised it to 58.7 mIoU. With all three objectives, including $\mathcal{L}_{rmix}$, accuracy reached 60.0 mIoU, confirming their complementary effect. The averaged improvement condition ζ also consistently increased as each loss component is added. For the segmentation network (Table 4), the supervised-only baseline also scored 50.9 mIoU. Adding the semi-supervised loss, $\mathcal{L}_{sunl}$, improved performance to 58.1 mIoU, while its variant without symmetric cross-entropy ($\blacktriangle$) with $\mathcal{L}_{smix}$ gave 58.0 mIoU. Using all three training objectives yielded the best result

Table 2: Comparison of different semi-supervised learning methods on ScribbleKITTI while varying the ratio of labeled data. The best results are in bold, and the second best underlined.

Method	1%	10%	20%	50%	Avg.
<i>Sup-only</i>	39.2	48.0	52.1	53.8	48.3
MT [34]	41.0	50.1	52.8	53.9	49.5
CBST [45]	41.5	50.6	53.3	54.5	50.0
CPS [3]	41.4	51.8	53.9	54.8	50.5
LaserMix [13]	44.2	53.7	55.1	56.8	52.5
IT2 [24]	47.9	56.7	57.5	58.3	55.1
FrustrumMix [39]	45.6	55.7	58.2	60.8	55.1
LaserMix++ [14]	47.3	56.7	57.6	59.8	55.4
RePL (Ours)	48.7	57.3	57.6	61.0	56.1

Table 4: Impact of the losses for learning the LiDAR semantic segmentation network. In $\mathcal{L}_{\text{sunl}}$, $\blacktriangle$ denotes the omission of the symmetric cross-entropy.

$\mathcal{L}_{\text{ssup}}$	$\mathcal{L}_{\text{sunl}}$	$\mathcal{L}_{\text{smix}}$	mIoU
✓			50.9
✓	✓		58.1
✓	$\blacktriangle$	✓	58.0
✓	✓	✓	60.0

of 60.0 mIoU, highlighting the benefit of jointly optimizing supervised learning on labeled data, semi-supervised learning with refined pseudo-labels, and mixed scene training.

Sensitivity to the Quality of Error Candidate Mask. We analyzed how the quality of the error candidate mask influences inference performance by replacing our heuristic error mask with different alternatives on the validation set. As shown in Table 5, random masks yielded modest improvements over the baseline (no refinement) of the teacher. Our heuristic error mask provided a clear gain, while an oracle error mask derived from ground-truth labels further improved performance to 67.3 mIoU. These results indicate that even a simple heuristic achieves competitive improvements, with more accurate error mask offering substantial room for further gains.

Impact of the Random Masking Strategy. We investigated whether training with random masking improves performance on the validation set. As shown in Table 6, incorporating random masking yielded higher performance (60.0

Table 3: Impact of the losses for the pseudo-label refiner. Each row shows the average improvement condition ζ and mean IoU when training with different subsets of the losses.

$\mathcal{L}_{\text{rsup}}$	$\mathcal{L}_{\text{runl}}$	$\mathcal{L}_{\text{rmix}}$	ζ	mIoU
			-	50.9
✓			0.327	57.2
✓	✓		0.353	58.7
✓	✓	✓	0.430	60.0

Table 5: Sensitivity to the quality of the error candidate mask in LiDAR semantic segmentation accuracy. Different error mask generation strategies were compared at inference time.

Setting	Baseline	Random				Oracle	Ours
		25%	50%	75%			
mIoU	57.0	57.6	58.2	58.7	67.3	60.0	

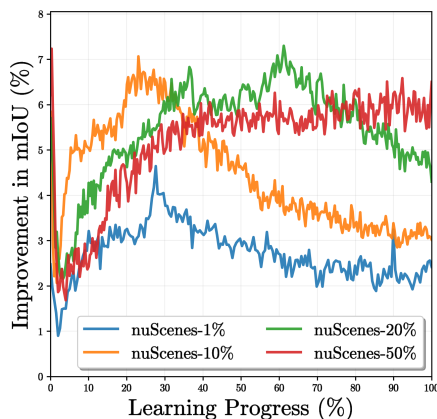


Fig. 5: Pseudo-label quality improvement by the refiner during training on nuScenes-lidarseg.

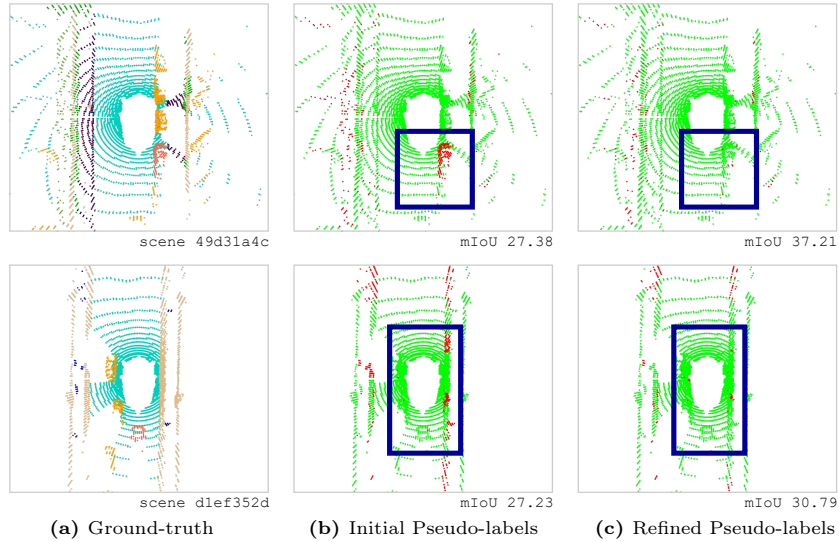


Fig. 3: Qualitative results of refined pseudo-labels and their initial predictions on the unlabeled set of nuScenes-lidarseg, at the end of training. Correct and incorrect predictions are shown in green and red, respectively. The model was trained with a 1% label ratio.

Table 6: Impact of the random masking strategy for training the pseudo-label refiner in the segmentation quality of the final model.

Setting	mIoU
w/o Random Masking	57.7
w/ Random Masking	60.0

Table 7: Computational cost analysis on nuScenes-lidarseg.

Method	Latency (s)	Memory (MB)	mIoU
Baseline	0.43	1231	50.9
Baseline + Refiner	0.68	1627	60.0
Δ	+0.25	+396	+9.1

mIoU) compared to training without it (57.7 mIoU). This indicates that random masking serves as a regularizer, helping the network handle erroneous predictions more effectively and improving performance during inference.

Analysis on Pseudo-label Quality Improvement throughout Training.

We report the trend of pseudo-label quality improvement throughout training for different labeled data ratios (1%, 10%, 20%, 50%) on the unlabeled data of nuScenes-lidarseg in Fig. 5. During early stage of training, the improvement was relatively low across all ratios as the refiner learns error correction from scratch. As training progresses, the improvement increased as the refiner learns to correct errors more effectively. However, the improvement gradually declined in later stages as the segmentation network itself becomes accurate, leaving less room for the refiner to provide meaningful corrections to the already high-quality predictions. REPL showed effectiveness across all labeled data ratios, with better performance and scalability at higher ratios.

Analysis on Computational Cost. To quantify the additional overhead by the refiner, we measured the latency and memory usage during inference on

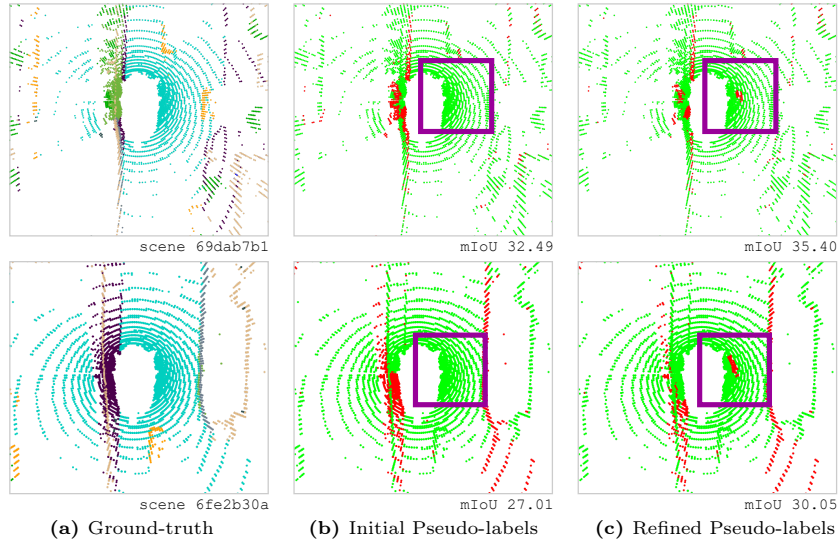


Fig. 4: Failure cases on the unlabeled set of nuScenes-lidarseg at the end of training with a 1% label ratio. Correct and incorrect predictions are shown in green and red, respectively.

the validation set using a single batch. As shown in Table 7, the refiner adds approximately 0.25 seconds of latency and 396 MB of memory, while providing a substantial improvement of +9.1 mIoU from the supervised-only baseline. These results demonstrate that the added computational cost is moderate relative to the significant accuracy gains.

Ablation Study on Unreliable Voxel Identification. We analyzed the sensitivity of an unreliable voxel identification in REPL on the validation set. As shown in Figure 6, the confidence percentile $\kappa = 0.4$ yields the best performance at 60.0 mIoU, while $\kappa = 0.2, 0.3, 0.5$, and 0.6 result in sub-optimal performance at 55.1, 56.8, 55.7, and 58.4 mIoU, respectively.

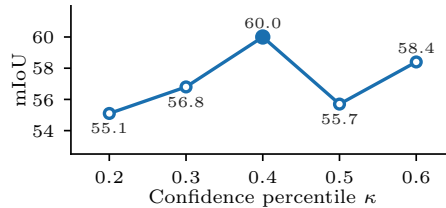


Fig. 6: Sensitivity analysis of a hyper-parameter κ .

Analysis on Failure Cases. Despite overall improvements, REPL occasionally introduces errors by over-correcting initially accurate predictions. Figure 4 shows representative failure cases (purple boxes). Nevertheless, the mIoU gain indicates that successful corrections outweigh these localized failures, leading to overall enhancement of the pseudo-labels.

Discard Baseline. To verify that performance gains are derived from the refinement module and not from discarding uncertain voxels, we compared against a discard-only baseline that uses the same error masks as REPL but removes detected unreliable voxels from the pseudo-label loss instead of reconstructing them. On nuScenes-lidarseg 1%, this baseline achieves 53.5 mIoU, below self-

Table 8: Per-class pseudo-label quality improvement on the unlabeled set of nuScenes-lidarseg with 1% labeled ratio. The column marked as % indicates class frequency.

Class	%	Δ IoU	Class	%	Δ IoU	Class	%	Δ IoU	Class	%	Δ IoU
barrier	1.10	+3.45	bus	0.55	+1.04	const. veh.	0.18	+2.54	motorcycle	0.05	+2.55
bicycle	0.02	+2.68	car	4.42	+1.60	pedestrian	0.27	+6.47	traffic cone	0.09	+7.86
trailer	0.58	+0.60	truck	1.83	+1.02	driv. surf.	37.6	-1.47	flat other	1.02	+2.60
sidewalk	8.31	-2.40	terrain	8.34	+4.14	manmade	21.1	+4.41	vegetation	14.5	+5.35

training with original pseudo-labels (54.3 mIoU) and far below REPL (60.0 mIoU). This indicates that the improvements are driven by the proposed refinement module.

Per-class Analysis. As shown in Table 8, REPL improved pseudo-label quality for 14 of 16 classes on the unlabeled set of nuScenes-lidarseg with 1% labeled ratio. Rare and small classes such as traffic cone (0.09% freq., +7.86 mIoU), pedestrian (0.27%, +6.47), vegetation (14.5%, +5.35), and manmade (21.1%, +4.41) benefited the most. Overall, the mIoU improved from 59.42 to 62.07.

Range-based Analysis. To show that pseudo-label refinement is not limited to near-range regions, we measured pseudo-label quality gains across distance bins on nuScenes-lidarseg 1%. REPL improved pseudo-label quality consistently across all evaluated distance bins (0–40 m), with gains of +1.95 (0–10 m), +1.30 (10–20 m), +1.34 (20–30 m), and +1.26 (30–40 m). The results confirm that the refinement also benefits far-range regions.

Generalization to Other Backbones. To assess whether REPL is tied to Cylinder3D, we replaced it with SphereFormer [17] without architecture-aware tuning. On nuScenes-lidarseg 1%, REPL with SphereFormer improved from the supervised-only baseline of 50.7 to 56.5 mIoU. This demonstrates that the proposed framework is not architecture-specific.

5 Conclusion

We presented REPL, a semi-supervised learning framework for LiDAR semantic segmentation that refines pseudo-labels through a two-stage mechanism of error estimation and masked reconstruction. The framework integrates a teacher-student segmentation network with a pseudo-label refiner to identify unreliable predictions and reconstruct them into cleaner supervision signals. We also provided theoretical analysis establishing the mathematical conditions under which refinement improves pseudo-label quality. With this design, our method achieved state-of-the-art results on nuScenes-lidarseg and SemanticKITTI across various label ratios. A current limitation is that the consensus-based error estimation may miss errors on which the teacher and student are jointly overconfident. Replacing the heuristic error mask with a learned error-localization approach such as ELN [16] is a promising direction for further closing the gap to the oracle.

Acknowledgments. This work was supported by the Scaleup TIPS grant (RS-2023-00321784: Development of Novel Generative AI Technology to Generate Domain-Specific Synthetic Data) and the IITP grants (RS-2022-II220290, RS-2022-II220926, RS-2024-00509258, RS-2024-00469482, RS-2024-00457882, RS-2019-II191906) funded by Ministry of Science and ICT, Korea.

References

1. Behley, J., Garbade, M., Milioto, A., Quenzel, J., Behnke, S., Stachniss, C., Gall, J.: SemanticKITTI: A Dataset for Semantic Scene Understanding of LiDAR Sequences. In: Proc. IEEE International Conference on Computer Vision (ICCV) (2019)
2. Berman, M., Rannen Triki, A., Blaschko, M.B.: The lovász-softmax loss: A tractable surrogate for the optimization of the intersection-over-union measure in neural networks. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2018)
3. Chen, X., Yuan, Y., Zeng, G., Wang, J.: Semi-supervised semantic segmentation with cross pseudo supervision. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (June 2021)
4. Choy, C., Gwak, J., Savarese, S.: 4d spatio-temporal convnets: Minkowski convolutional neural networks. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2019)
5. Cortinhal, T., Tzelepis, G., Erdal Aksoy, E.: Salsanext: Fast, uncertainty-aware semantic segmentation of lidar point clouds. In: Bebis, G., Yin, Z., Kim, E., Bender, J., Subr, K., Kwon, B.C., Zhao, J., Kalkofen, D., Baci, G. (eds.) *Advances in Visual Computing* (2020)
6. Fong, W.K., Mohan, R., Hurtado, J.V., Zhou, L., Caesar, H., Beijbom, O., Valada, A.: Panoptic nusenes: A large-scale benchmark for lidar panoptic segmentation and tracking. arXiv preprint arXiv:2109.03805 (2021)
7. Geiger, A., Lenz, P., Urtasun, R.: Are we ready for autonomous driving? the kitti vision benchmark suite. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2012)
8. He, K., Chen, X., Xie, S., Li, Y., Dollár, P., Girshick, R.: Masked autoencoders are scalable vision learners. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2022)
9. Hu, Q., Yang, B., Xie, L., Rosa, S., Guo, Y., Wang, Z., Trigoni, N., Markham, A.: Randla-net: Efficient semantic segmentation of large-scale point clouds. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2020)
10. Jiang, L., Shi, S., Tian, Z., Lai, X., Liu, S., Fu, C.W., Jia, J.: Guided point contrastive learning for semi-supervised point cloud semantic segmentation. In: Proc. IEEE International Conference on Computer Vision (ICCV) (2021)
11. Jiang, L., Zhao, H., Liu, S., Shen, X., Fu, C.W., Jia, J.: Hierarchical point-edge interaction network for point cloud semantic segmentation. In: Proc. IEEE International Conference on Computer Vision (ICCV) (October 2019)
12. Kim, Y., Yim, J., Yun, J., Kim, J.: Nlnl: Negative learning for noisy labels. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (October 2019)
13. Kong, L., Ren, J., Pan, L., Liu, Z.: Lasermix for semi-supervised lidar semantic segmentation. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2023)
14. Kong, L., Xu, X., Ren, J., Zhang, W., Pan, L., Chen, K., Ooi, W.T., Liu, Z.: Multi-modal data-efficient 3d scene understanding for autonomous driving. *IEEE Transactions on Pattern Analysis and Machine Intelligence (TPAMI)* (2025)
15. Kwon, D., Kwak, S.: Semi-supervised semantic segmentation with error localization network. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2022)

16. Kwon, D., Kwak, S.: Semi-supervised semantic segmentation with error localization network. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 9957–9967 (June 2022)
17. Lai, X., Chen, Y., Lu, F., Liu, J., Jia, J.: Spherical transformer for lidar-based 3d recognition. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 17545–17555 (2023). <https://doi.org/10.1109/CVPR52729.2023.01683>
18. Lai, X., Liu, J., Jiang, L., Wang, L., Zhao, H., Liu, S., Qi, X., Jia, J.: Stratified transformer for 3d point cloud segmentation. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2022)
19. Li, J., Dong, Q.: Density-guided semi-supervised 3d semantic segmentation with dual-space hardness sampling. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2024)
20. Li, L., Shum, H.P., Breckon, T.P.: Less is more: Reducing task and model complexity for 3d point cloud semantic segmentation. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2023)
21. Liu, C., Weng, X., Jiang, S., Li, P., Yu, L., Xia, G.S.: Exploring scene affinity for semi-supervised lidar semantic segmentation. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2025)
22. Liu, M., Zhou, Y., Qi, C.R., Gong, B., Su, H., Anguelov, D.: Less: Label-efficient semantic segmentation for lidar point clouds. arXiv preprint arXiv:2210.08064 (2022)
23. Liu, Y., Kong, L., Cen, J., Chen, R., Zhang, W., Pan, L., Chen, K., Liu, Z.: Segment any point cloud sequences by distilling vision foundation models. Proc. Neural Information Processing Systems (NeurIPS) (2023)
24. Liu, Y., Chen, Y., Wang, H., Belagiannis, V., Reid, I., Carneiro, G.: Ittakestwo: Leveraging peer representations for semi-supervised lidar semantic segmentation. In: Proc. European Conference on Computer Vision (ECCV) (2024)
25. Loshchilov, I., Hutter, F.: Decoupled weight decay regularization. In: Proc. International Conference on Learning Representations (ICLR) (2019)
26. Milioto, A., Vizzo, I., Behley, J., Stachniss, C.: RangeNet++: Fast and Accurate LiDAR Semantic Segmentation. In: IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS) (2019)
27. Nunes, L., Wiesmann, L., Marcuzzi, R., Chen, X., Behley, J., Stachniss, C.: Temporal consistent 3d lidar representation learning for semantic perception in autonomous driving. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2023)
28. Paszke, A., Gross, S., Chintala, S., Chanan, G., Yang, E., DeVito, Z., Lin, Z., Desmaison, A., Antiga, L., Lerer, A.: Automatic differentiation in pytorch. In: AutoDiff, NIPS Workshop (2017)
29. Pendleton, S.D., Andersen, H., Du, X., Shen, X., Meghjani, M., Eng, Y.H., Rus, D., Ang, M.H.: Perception, planning, control, and coordination for autonomous vehicles. Machines (2017)
30. Qi, C., Su, H., Mo, K., Guibas, L.J.: Pointnet: Deep learning on point sets for 3d classification and segmentation. Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2016)
31. Roriz, R., Cabral, J., Gomes, T.: Automotive lidar technology: A survey. IEEE Transactions on Intelligent Transportation Systems (2022)
32. Sautier, C., Puy, G., Gidaris, S., Boulch, A., Bursuc, A., Marlet, R.: Image-to-lidar self-supervised distillation for autonomous driving data. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2022)

33. Serfling, B., Reichert, H., Bayerlein, L., Doll, K., Radkhah-Lens, K.: Lidar based semantic perception for forklifts in outdoor environments (2025)
34. Tarvainen, A., Valpola, H.: Mean teachers are better role models: Weight-averaged consistency targets improve semi-supervised deep learning results. In: Proc. Neural Information Processing Systems (NeurIPS) (2017)
35. Unal, O., Dai, D., Hoyer, L., Can, Y.B., Van Gool, L.: 2d feature distillation for weakly- and semi-supervised 3d semantic segmentation. In: Proc. IEEE Winter Conf. on Applications of Computer Vision (WACV) (January 2024)
36. Unal, O., Dai, D., Van Gool, L.: Scribble-supervised lidar semantic segmentation. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (June 2022)
37. Wang, N., Guo, R., Shi, Ziyue Wang, C., Zhang, H., Lu, H., Zheng, Z., Chen, X.: SegNet4D: Efficient Instance-Aware 4D Semantic Segmentation for LiDAR Point Cloud. arXiv preprint (2024)
38. Wang, Y., Ma, X., Chen, Z., Luo, Y., Yi, J., Bailey, J.: Symmetric cross entropy for robust learning with noisy labels. In: Proc. IEEE International Conference on Computer Vision (ICCV). pp. 322–330 (2019). <https://doi.org/10.1109/ICCV.2019.00041>
39. Xu, X., Kong, L., Shuai, H., Liu, Q.: Frnet: Frustum-range networks for scalable lidar segmentation. IEEE Transactions on Image Processing (TIP) (2025)
40. Xu, X., Kong, L., Shuai, H., Zhang, W., Pan, L., Chen, K., Liu, Z., Liu, Q.: 4d contrastive superflows are dense 3d representation learners. In: Proc. European Conference on Computer Vision (ECCV) (2024)
41. Yang, L., Zhuo, W., Qi, L., Shi, Y., Gao, Y.: ST++: Make self-training work better for semi-supervised semantic segmentation. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2022)
42. Zhang, Y., Zhou, Z., David, P., Yue, X., Xi, Z., Gong, B., Foroosh, H.: Polarnet: An improved grid representation for online lidar point clouds semantic segmentation. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (June 2020)
43. Zhang, Z., Bai, M., Li, E.: Implicit surface contrastive clustering for lidar point clouds. In: Proc. IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2023)
44. Zhu, X., Zhou, H., Wang, T., Hong, F., Li, W., Ma, Y., Li, H., Yang, R., Lin, D.: Cylindrical and asymmetrical 3d convolution networks for lidar-based perception. IEEE Transactions on Pattern Analysis and Machine Intelligence (TPAMI) (2022)
45. Zou, Y., Yu, Z., Kumar, B.V., Wang, J.: Unsupervised domain adaptation for semantic segmentation via class-balanced self-training. In: Proc. European Conference on Computer Vision (ECCV) (September 2018)