

PLAD: Learning to Infer Shape Programs with Pseudo-Labels and Approximate Distributions

R. Kenny Jones
Brown University

Homer Walke
UC Berkeley

Daniel Ritchie
Brown University

Abstract

Inferring programs which generate 2D and 3D shapes is important for reverse engineering, editing, and more. Training models to perform this task is complicated because paired (shape, program) data is not readily available for many domains, making exact supervised learning infeasible. However, it is possible to get paired data by compromising the accuracy of either the assigned program labels or the shape distribution. Wake-sleep methods use samples from a generative model of shape programs to approximate the distribution of real shapes. In self-training, shapes are passed through a recognition model, which predicts programs that are treated as ‘pseudo-labels’ for those shapes. Related to these approaches, we introduce a novel self-training variant unique to program inference, where program pseudo-labels are paired with their executed output shapes, avoiding label mismatch at the cost of an approximate shape distribution. We propose to group these regimes under a single conceptual framework, where training is performed with maximum likelihood updates sourced from either Pseudo-Labels or an Approximate Distribution (PLAD). We evaluate these techniques on multiple 2D and 3D shape program inference domains. Compared with policy gradient reinforcement learning, we show that PLAD techniques infer more accurate shape programs and converge significantly faster. Finally, we propose to combine updates from different PLAD methods within the training of a single model, and find that this approach outperforms any individual technique.

1. Introduction

Having access to a procedure which generates a visual datum reveals its underlying structure, facilitating high-level manipulation and editing by a person or autonomous agent. Thus, inferring such programs from visual data is an important problem. In $\mathbb{R}^2$, inferring shape programs has applications in the design of diagrams, icons, and other 2D graphics. In $\mathbb{R}^3$, it has applications in reverse engineering of

CAD models, procedural modeling for 3D games, and 3D structure understanding for autonomous agents.

We formally define shape program inference as obtaining a latent program $\mathbf{z}$ which generates a given observed shape $\mathbf{x}$. We model $p(\mathbf{z}|\mathbf{x})$ with deep neural networks that train over a distribution of real shapes in order to amortize the cost of shape program inference on unseen shapes (e.g. a test set). This is a challenging problem: it is a structured prediction problem whose output is high-dimensional and can feature both discrete and continuous components (i.e. program control flow vs. program parameters). Nevertheless, learning $p(\mathbf{z}|\mathbf{x})$ becomes tractable provided that one has access to paired $(\mathbf{X}, \mathbf{Z})$ data (i.e. a dataset of shapes and the programs which generate them) [31]. Unfortunately, while shape data is increasingly available in large quantities [2], these shapes do not typically come with their generating program.

To circumvent this data problem, researchers have typically *synthesized* paired data by generating synthetic programs and pairing them with the shapes they output [24, 28]. However, as there is typically significant distributional mismatch between these synthetic shapes and “real” shapes from the distribution of interest, S^* , various techniques must be employed to fine-tune $p(\mathbf{z}|\mathbf{x})$ models towards S^* .

A number of these fine-tuning strategies attempt to directly propagate gradients from geometric similarity measures back to $p(\mathbf{z}|\mathbf{x})$. When the program executor is not a black-box, it may be possible to do this by implementing a differentiable relaxation of its behavior [16], but this requires knowledge of its functional form. One can also try learning a differentiable proxy of the executor’s behavior [28], but this approximation introduces errors. Moreover, as shape programs typically involve many discrete structural decisions, training such a model end-to-end is usually infeasible in many domains. Thus, many prior works often resort to general-purpose policy gradient reinforcement learning [8, 24], which treats the program executor as a (non-differentiable) black-box. The downside of this strategy is that RL is notoriously unstable and slow to converge.

In this paper, we study a collection of methods that create

(shape, program) data pairs used to train $p(\mathbf{z}|\mathbf{x})$ models with maximum likelihood estimation (MLE) updates while treating the program executor as a black-box. As discussed, ground-truth (shape, program) pairs are often unavailable, so these techniques must make compromises in how they formulate paired data. In wake-sleep, a generative model $p(\mathbf{z})$ is trained to convergence on alternating cycles with respect to $p(\mathbf{z}|\mathbf{x})$. When training $p(\mathbf{z}|\mathbf{x})$, paired data can be created by sampling from $p(\mathbf{z})$. Each program label $\mathbf{z}$ is valid with respect to its associated $\mathbf{x}$ shape, but there is often a distributional mismatch between the generated set of shapes, $\mathbf{X}$, and shapes from the target distribution, S^* . In self-training, one uses $p(\mathbf{z}|\mathbf{x})$ to infer latent $\mathbf{z}$'s for unlabeled input $\mathbf{x}$'s; these $\mathbf{z}$'s then become "pseudo-labels" which are treated as ground truth for another round of supervised training. In this paradigm, there is no distributional shift between $\mathbf{X}$ and S^* , but each $\mathbf{z}$ is only an approximately correct label with respect to its paired $\mathbf{x}$.

We observe that shape program inference has a unique property that makes it especially well-suited for self-training: the distribution $p(\mathbf{x}|\mathbf{z})$ is known *a priori*—this is a delta distribution defined by the program executor. When using a model $p(\mathbf{z}|\mathbf{x})$ to infer a program $\mathbf{z}$ from some shape $\mathbf{x}^*$ of interest, one can use this executor to produce a shape $\mathbf{x}$ that is consistent with the program $\mathbf{z}$: in the terminology of self-training, $\mathbf{z}$ is guaranteed to be the "correct label" for $\mathbf{x}$. However, similar to wake-sleep, formulating $\mathbf{X}$ as shape executions produced by model inferred programs can cause a distributional shift between $\mathbf{X}$ and S^* . Since this variant of self-training involves executing the inferred latent program $\mathbf{z}$, we call this procedure latent execution self-training (LEST).

As all of the aforementioned fine-tuning regimes use either Pseudo-Labels or Approximate Distributions to formulate (shape, program) pairs, we group them under a single conceptual framework: PLAD. We evaluate PLAD methods experimentally, using them to fine-tune shape program inference models in multiple shape domains: 2D and 3D constructive solid geometry (CSG), and assembly-based modeling with ShapeAssembly, a domain-specific language for structures of manufactured 3D objects [14]. We find that PLAD training regimes offer substantial advantages over the de-facto approach of policy gradient reinforcement learning, achieving better shape reconstruction performance while requiring significantly less computation time. Further, we explore combining training updates from a mixture of PLAD methods, and find that this approach leads to better performance compared with any individual method. Code for our method and experiments can be found at <https://github.com/rkjones4/PLAD>.

In summary, our contributions are:

1. Proposing the PLAD conceptual framework to group a family of related self-supervised learning techniques for shape program inference.

2. Introducing latent execution self-training, a PLAD method, to take advantage of the unique properties of the shape program inference problem.
3. Experiments across multiple 2D and 3D shape program inference domains, demonstrating that (i) fine-tuning under PLAD regimes outperforms policy gradient reinforcement learning and (ii) combining PLAD methods is better than any individual technique.

2. Related Work

Program synthesis is a broad field that has employed many techniques throughout its history. A program synthesizer takes as input a domain-specific language (DSL) and a specification; it outputs a program in the DSL that meets the specification. Machine learning has been used to improve performance on program synthesis tasks by e.g. performing a neurally-guided search over all possible programs or letting a recurrent network predict program text directly [1, 5, 7, 22, 27, 32].

In this work, we are interested in the sub-problem of *shape program inference*, which is a type of *visual program induction* problem [4]. In our case, the specification is a visual representation of an object which the inferred program's output must geometrically match. Here, we discuss prior work that has attacked this problem, organized by methodology used to learn $p(\mathbf{z}|\mathbf{x})$; see Table 1 for an overview. As discussed, a common practice in this prior work is to start with a model that has been pretrained on synthetically generated (shape, program) pairs with supervised learning, and then perform fine-tuning towards a distribution of interest [8, 9, 19, 24, 28].

Policy Gradient Reinforcement Learning The most general method for fine-tuning a pretrained $p(\mathbf{z}|\mathbf{x})$ is reinforcement learning: treating $p(\mathbf{z}|\mathbf{x})$ as a policy network and using policy gradient methods [30]. The geometric similarity of the inferred program's output to its input is the reward function; the program executor $p(\mathbf{x}|\mathbf{z})$ can be treated as a (non-differentiable) black-box. CSG-Net uses RL for fine-tuning [24, 25], as does other recent work on inferring CSG programs from input geometry [8]. While CSG-Net has been improved to allow it to converge without supervised pretraining [34], not starting from the supervised model results in worse performance. The main problem with policy gradient RL is its instability due to high variance gradients, leading to slow convergence. Like RL, PLAD methods treat the program executor as a black-box, but as we show experimentally, they converge faster and achieve better reconstruction performance.

Differentiable Executor If the functional form of the program executor $p(\mathbf{x}|\mathbf{z})$ is known and differentiable, then the gradient of the reward with respect to the parameters of $p(\mathbf{z}|\mathbf{x})$ can be computed, making policy gradient unnecessary. Shape programs are typically not fully differentiable, as

Method	Models	Black-Box $p(\mathbf{x} \mathbf{z})$?	$X = S^*$	Low variance, unbiased gradients
Policy gradient RL	$p(\mathbf{z} \mathbf{x})$	✓	✓	X
Differentiable executor	$p(\mathbf{z} \mathbf{x})$	X	✓	✓
Variational Bayes	$p(\mathbf{z} \mathbf{x}), p(\mathbf{z})$	X	✓	✓
Wake-sleep, EM	$p(\mathbf{z} \mathbf{x}), p(\mathbf{z})$	✓	X	✓
Self-training	$p(\mathbf{z} \mathbf{x})$	✓	✓	X
LEST	$p(\mathbf{z} \mathbf{x})$	✓	X	✓

Table 1. Comparison of different methods for fine-tuning $p(\mathbf{z}|\mathbf{x})$, in terms of the models that must be trained, if they treat the program executor as black-box, if their distribution of training shapes matches the distribution real shapes ($X = S^*$), and if their loss gradients are unbiased with low-variance. The last three rows describe methods that fall under the PLAD framework.

they often involve discrete choices (e.g. which type of primitives to create). UCSGNet uses a differentiable relaxation to circumvent this issue [16]. Other work trains a differentiable network to approximate the behavior of the program executor [28], which introduces errors. PLAD regimes do not require the program executor to be differentiable, yet they perform better than other approaches (e.g. policy gradient RL) that share this desirable property.

Generative Model Learning Shape program inference has also been explored in the context of learning a generative model $p(\mathbf{x}, \mathbf{z})$ of programs and the shapes they produce. The most popular approach for training such models is variational Bayes, in particular the variational autoencoder [17]. This method simultaneously trains a generative model $p(\mathbf{x}, \mathbf{z})$ and a recognition model $p(\mathbf{z}|\mathbf{x})$ by optimizing a lower bound on the marginal likelihood $p(\mathbf{x})$. When the $\mathbf{z}$ ’s are shape programs, the program executor is $p(\mathbf{x}|\mathbf{z})$, so learning the generative model reduces to learning a prior over programs $p(\mathbf{z})$. Training such models with gradient descent requires that the executor $p(\mathbf{x}|\mathbf{z})$ be differentiable. When this is not possible, the wake-sleep algorithm is a viable alternative [13]. This approach alternates training steps of the generative and recognition models, training one on samples produced by the other. Recent work has used wake-sleep for visual program induction [10, 12]. If one trains the generative model and the inference model to convergence before switching to training the other, this is equivalent to expectation maximization (viewed as alternating maximization [21]).

Self-Training Traditionally, self-training has been employed in weakly-supervised learning paradigms to increase the predictive accuracy of simple classification models [20, 23, 33]. Recently, renewed interest in self-training-inspired data augmentation approaches have demonstrated empirical performance improvements for neural models in domains such as large-scale image classification, machine translation, and speech recognition [11, 15, 35]. But while self-training has been shown to yield practical gains for some domains, for others it can actually lead to worse performance, as training on too many incorrect pseudo-labels can

cause learning to degrade [3, 26]. For self-training within the PLAD framework, the assigned pseudo-label for each example changes during fine-tuning whenever the inference model discovers a program that better explains the input shape; similar techniques have been proposed for learning programs that perform semantic parsing under the view of iterative maximum likelihood [18]. To our knowledge, self-training has not been applied for fine-tuning visual program inference models, likely because it is somewhat unintuitive to view a program as a “label” for a visual datum.

3. Method

In this section, we describe the PLAD framework: a conceptual grouping of fine-tuning methods for shape program inference models. Our formulation assumes three inputs: a training dataset of shapes from the distribution of interest, S^* , a program inference model, $p(\mathbf{z}|\mathbf{x})$, and a program executor that converts programs into shapes, E . Throughout this paper, we assume that the $p(\mathbf{z}|\mathbf{x})$ passed as input has undergone supervised pretraining on a distribution of synthetically generated shapes. Methods within the PLAD framework return a fine-tuned $p(\mathbf{z}|\mathbf{x})$ specialized to the distribution of interest from which S^* was sampled.

We depict the PLAD procedure both algorithmically and pictorially in Figure 1. To fine-tune $p(\mathbf{z}|\mathbf{x})$ towards S^* , PLAD methods iterate through the following steps: (1) use $p(\mathbf{z}|\mathbf{x})$ to find visually similar programs to S^* , (2) construct a dataset of shape and program pairs $(\mathbf{X}, \mathbf{Z})$ using the inferred programs, and (3) fine-tune $p(\mathbf{z}|\mathbf{x})$ with maximum likelihood estimation updates on batches from $(\mathbf{X}, \mathbf{Z})$. Through successive iterations, these steps bootstrap one another, forming a virtuous cycle: improvements to $p(\mathbf{z}|\mathbf{x})$ create $(\mathbf{X}, \mathbf{Z})$ pairs that more closely match the statistics of S^* , and training on better $(\mathbf{X}, \mathbf{Z})$ pairs specializes $p(\mathbf{z}|\mathbf{x})$ to S^* .

Methods that fall within the PLAD framework differ in how the paired $(\mathbf{X}, \mathbf{Z})$ data is created within each round. We detail this process for wake-sleep (Section 3.1), self-training (Section 3.2), and latent execution self-training (Section 3.3). In Section 3.4 we explain our program inference procedure

Input: $(S^*, p(z|x), E)$
Output: $p(z|x)$ fine-tuned on S^*

```

 $P^{BEST} \leftarrow \{\}$ 
for Number of Rounds do
  // Update Best Programs
   $P^{BEST} \leftarrow \text{inferProgs}(p(z|x), S^*, P^{BEST})$ 
  // Create Training Data
  if Self-Train then
     $Z \leftarrow P^{BEST}$ 
     $X \leftarrow S^*$ 
  else if LEST then
     $Z \leftarrow P^{BEST}$ 
     $X \leftarrow \{E(z) \mid z \in Z\}$ 
  else if Wake-Sleep then
     $p(z) \leftarrow \text{trainGenerative}(P^{BEST})$ 
     $Z \leftarrow \text{sample}(p(z), |S^*|)$ 
     $X \leftarrow \{E(z) \mid z \in Z\}$ 
  end if
  // Train inference model
   $p(z|x) \leftarrow \text{trainMLE}(X, Z)$ 
end for

```

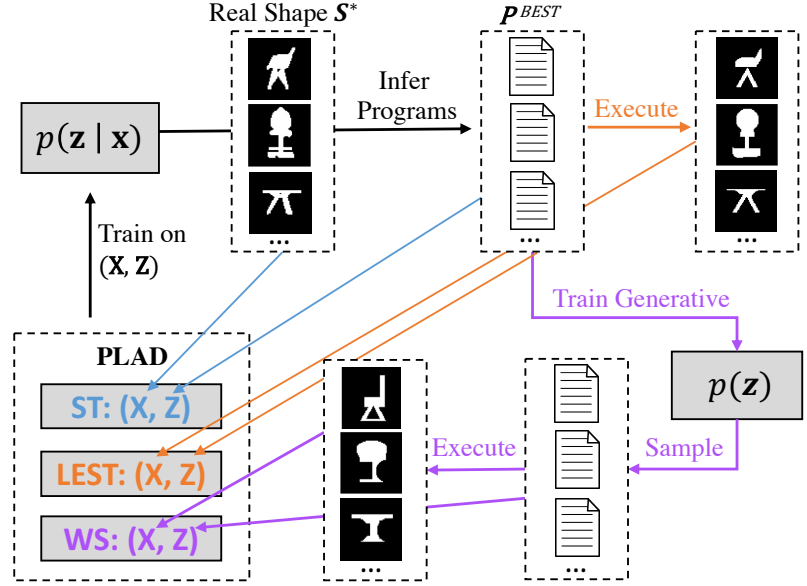


Figure 1. (Left) Pseudocode for fine-tuning shape program inference models, $p(z|x)$, towards a shape distribution of interest, S^* , with Pseudo-Labels and Approximate Distributions (PLAD). PLAD methods iterate through three steps: infer programs for S^* with $p(z|x)$, create a dataset of (X, Z) shape-program pairs, and train $p(z|x)$ on batches from (X, Z) . Self-training, latent execution self-training, and wake-sleep differ in how (X, Z) is constructed. (Right) A visual illustration of the algorithm’s dataflow.

(inferProgs, Fig 1). Finally, in Section 3.5 we discuss how a single $p(z|x)$ can be fine-tuned by multiple PLAD methods.

3.1. Wake-Sleep (X, Z) Construction

Wake-sleep uses a generative model, $p(z)$ to construct (X, Z) . In our implementation, we choose to model $p(z)$ as a variational auto-encoder (VAE) [17], where the encoder consumes visual data and the decoder outputs a program. To create data for $p(z)$, we take the current best programs discovered for S^* , P^{BEST} , and execute each program to form a set of shapes X^G . $p(z)$ is then trained on pairs from (X^G, P^{BEST}) in the typical VAE framework. Note that the design space for $p(z)$ is quite flexible, for instance, $p(z)$ can be trained without access to X^G if implemented with a program encoder.

Once $p(z)$ has converged, we use it to sample $|S^*|$ programs by decoding normally distributed random vectors. This set of programs becomes Z , and X is formed by executing each program in Z . In this set of (X, Z) programs, Z is always the correct label for X , so the gradient estimates during $p(z|x)$ training will be low-variance and unbiased. However, X is not guaranteed to be close to S^* , it is only an approximate distribution. Note though, that as P^{BEST} better approximates S^* , the distributional mismatch should

become smaller, as long as the generative model has enough capacity to properly model $p(z)$.

3.2. Self-Training (X, Z) Construction

Self-training constructs (X, Z) by assigning labels from the current best program set, P^{BEST} , to shape instances from S^* . Formally, $X \leftarrow S^*$ and $Z \leftarrow P^{BEST}$. This framing maintains the nice property that $X = S^*$, so there will never be distributional mismatch between these two sets. The downside is that unless programs from P^{BEST} exactly recreate their paired shapes from S^* when executed, we know that the pseudo-labels from P^{BEST} are ‘incorrect’. From this perspective, we can consider gradient estimates that come from such (X, Z) pairs to be biased. However, as we will show experimentally, when X forms a good approximation to S^* , sourcing gradient estimates from these pseudo-labels leads to strong reconstruction performance.

3.3. LEST (X, Z) Construction

A unique property of shape program inference is that the distribution $p(x|z)$ is readily available in the form of the program executor E . We leverage this property to propose LEST, a variant of self-training that does not create mismatch between pseudo-labels and their associated visual data. Sim-

ilar to the self-training paradigm, LEST first constructs $\mathbf{Z}$ as the current best program set, P^{BEST} . Then, differing from self-training, LEST constructs $\mathbf{X}$ as the executed version of each program in $\mathbf{Z}$. By construction, the labels in $\mathbf{Z}$ will now be correct for their paired shapes in $\mathbf{X}$. The downside is that, like wake-sleep, LEST may introduce a distributional mismatch between $\mathbf{X}$ and S^* . But once again, as P^{BEST} better approximates S^* , the mismatch between the two distributions will decrease.

3.4. Inferring Programs with $p(z|x)$

During each round of fine-tuning, PLAD methods rely on $p(z|x)$ to infer programs that approximate S^* . We propose to train PLAD methods such that the best matching inferred programs for S^* are maintained across rounds. Specifically, we construct a data structure P^{BEST} that maintains a program for each training shape in S^* . In this way, as more iterations are run, P^{BEST} always forms a closer approximation to S^* . There is an alternative framing, where P^{BEST} is reset each epoch, but we show experimental results in the supplemental material that this can lead to worse generalization.

To update P^{BEST} each round, we employ an inner-loop search procedure. For each shape in S^* , $p(z|x)$ suggests high-likelihood programs, and the P^{BEST} entry is updated to keep the program whose execution obtains the highest similarity to the input shape; the specific similarity metric varies by domain. While there are many ways to structure this inner-loop search, we choose beam-search, as we find it offers a good trade-off between speed and performance. Experimentally, we demonstrate that PLAD methods are capable of performing well even as the time spent on inner-loop search is varied (Section 4.4).

3.5. Training $p(z|x)$ with multiple PLAD methods

As detailed in the preceding sections, the main difference between PLAD approaches is in how they construct the $(\mathbf{X}, \mathbf{Z})$ dataset used for fine-tuning $p(z|x)$. However, there is no strict requirement that these different $(\mathbf{X}, \mathbf{Z})$ distributions be kept separate. We explore how $p(z|x)$ behaves under fine-tuning from multiple PLAD methods, such as combining LEST and self-training. We implement these mixtures on a per-batch basis. Before $p(z|x)$ training, we construct distinct $(\mathbf{X}, \mathbf{Z})$ distributions for each method in the combination. Then, during training, each batch is randomly sampled from one of the $(\mathbf{X}, \mathbf{Z})$ distributions. We experimentally validate the effectiveness of this approach in the next section.

4. Results

We evaluate a series of methods on their ability to fine-tune shape program inference models across multiple domains. We describe the different domains in Section 4.1 and details of our experimental design in Section 4.2. In Section 4.3, we compare the reconstruction accuracy of each method,

and study how they are affected by varying the time spent on inner-loop search (Section 4.4) and the size of the training set (Section 4.5). Finally, we explore the convergence speed of each method in Section 4.6.

4.1. Shape Program Domains

We run experiments across three shape program domains: 2D Constructive Solid Geometry (CSG), 3D CSG, and ShapeAssembly. Details can be found in the supplemental.

In CSG, shapes are created by declaring parametric primitives (e.g. circles, boxes) and combining them with Boolean operations (union, intersection, difference). CSG inference is non-trivial: as CSG uses non-additive operations (intersection, difference), inferring a CSG program does not simply reduce to primitive detection. For 2D CSG, we follow the grammar defined by CSGNet [24], using 400 shape tokens that correspond to randomly placed circles, triangles and rectangles on a 64 x 64 grid. For 3D CSG, we use a grammar that has individual tokens for defining primitives (ellipsoids and cuboids), setting primitive attributes (position and scales), and the three Boolean operators. Attributes are discretized into 32 bins.

ShapeAssembly is designed for specifying the part structure of manufactured 3D objects. It creates objects by declaring cuboid part geometries and assembling those parts together via attachment and symmetry operators. Our grammar contains tokens for each command type and parameter value; to handle continuous values, we discretize them into 32 bins.

4.2. Experimental Design

Fine-Tuning Methods We compare the ability of the following training schemes to fine-tune a model on a specific domain of interest:

- **SP:** $p(z|x)$ with supervised pretraining.
- **RL:** Fine-tuning with REINFORCE.
- **WS:** Fine-tuning with wake-sleep.
- **ST:** Fine-tuning with self-training.
- **LEST:** Fine-tuning with latent execution self-training.
- **LEST+ST:** combining LEST and ST.
- **LEST+ST+WS:** combining LEST, ST and WS.

Shape Datasets Fine-tuning methods learn to specialize $p(z|x)$ against a distribution of real shapes S^* . For each domain, we construct a dataset of shapes S^* , and split it into train, validation, and test sets. We perform early-stopping with respect to the validation set. For 2DCSG, we use the CAD dataset from CSGNet [24], which consists of front and side views of chairs, desks, and lamps from the Trimble 3D warehouse. We split the dataset into 10K shapes for training, 3K shapes for validation, and 3K shapes for testing. For 3D CSG and ShapeAssembly, we use CAD shapes from the chair, table, couches, and benches categories of ShapeNet; voxelizations are provided by [6]. We split the dataset into

Method	2D CSG CD ↓	3D CSG IoU ↑	ShapeAssembly IoU ↑
Supervised Pretraining (SP)	1.580	41.0	37.6
REINFORCE (RL)	1.097	53.4	50.8
Wake-Sleep (WS)	1.118	67.4	57.2
Self-Training (ST)	0.841	67.3	61.3
LEST	0.976	69.8	56.5
LEST+ST	0.829	70.8	66.0
LEST+ST+WS	0.811	74.3	66.4

Table 2. Test-set reconstruction performance across multiple shape program inference domains. The top row contains results for the pretrained $p(\mathbf{z}|\mathbf{x})$ model fine-tuned by the other methods. For 2D CSG the metric is Chamfer distance (CD, lower is better). For 3D CSG and ShapeAssembly the metric is intersection over union (IoU, higher is better). Individual PLAD methods outperform RL, and combining PLAD methods achieves the best performance across all domains (LEST+ST+WS).

10K shapes for training, 1K shapes for validation, and 1K shapes for testing.

Model Architectures For all experiments, we model $p(\mathbf{z}|\mathbf{x})$ in an encoder-decoder framework, although the particular architectures vary by domain. In all cases, the encoder is a CNN that converts visual data into a latent variable, and the decoder is an auto-regressive model that decodes the latent variables into a sequence of tokens. For 2D CSG, we use the same $p(\mathbf{z}|\mathbf{x})$ architecture as CSGNet. A CNN consumes 64×64 binary mask shape images to produce a latent code that initializes a GRU-based recurrent decoder. For 3D CSG and ShapeAssembly, we use a 3D CNN that consumes $32 \times 32 \times 32$ occupancy voxels. This CNN outputs a latent code that is attended to by a Transformer decoder network [29]; this network also attends over token sequences in a typical auto-regressive fashion.

Supervised Pretraining Before fine-tuning, $p(\mathbf{z}|\mathbf{x})$ undergoes supervised pretraining on synthetically generated programs until it has converged on that set. For 2D CSG, we follow CSGNet’s approach. For 3D CSG, we construct valid programs by (i) sampling a set of primitives within the allotted grid (ii) identifying potential overlaps (iii) constructing a binary tree of boolean operations using these overlaps. For ShapeAssembly, we propose programs by sampling random grammar expansions according to the language’s typing system. We then employ a validation step where a program is rejected if any of its part are not the sole occupying part of at least 8 voxels (to discourage excessive part overlaps). For 3D CSG and ShapeAssembly we sample 2 million synthetic programs and train until convergence on a validation set of 1000 programs. Full details provided in the supplemental.

4.3. Reconstruction Accuracy

We evaluate the performance of each fine-tuning method according to reconstruction accuracy: how closely the output of a shape program matches the input shape from which it was inferred, on a held out set of test shapes. The specific metric varies by domain. For 2D CSG, we follow CSGNet

and use Chamfer Distance (CD), where lower distances indicate more similar shapes. For 3D CSG and ShapeAssembly, we use volumetric intersection over union (IoU).

For each domain, we run each fine-tuning method to convergence, starting with the same $p(\mathbf{z}|\mathbf{x})$ model that has undergone supervised pretraining. The reward for RL models follows the similarity metric in each domain: CD for 2D CSG; IoU for 3D CSG and ShapeAssembly. For PLAD fine-tuning methods, the similarity metric in each domain determines which program is kept during updates to P^{BEST} . At inference time, when evaluating the reconstruction performance of each $p(\mathbf{z}|\mathbf{x})$, we employ a beam search procedure, decoding multiple programs in parallel, and choosing the program that achieves the highest similarity to the target shape. We use a beam size of 10, unless otherwise stated.

We present quantitative results in Table 2. Looking at the middle four rows, the two self-training variants (ST and LEST) outperform RL as a fine-tuning method in all the domains we studied. The wake-sleep variant (WS) also outperforms RL for both 3D CSG and ShapeAssembly. These are more challenging domains with larger token spaces, posing difficulties for policy gradient fine-tuning. As demonstrated by the last two rows, further improvement can be had by combining multiple methods: for each domain, the best performance is achieved by LEST+ST+WS. In fact, for 2DCSG, the test-set reconstruction accuracy achieved by LEST+ST+WS (0.811) outperforms previous state-of-the-art results, CSGNet (1.14) [24] and CSGNetStack (1.02) [25], for paradigms where the executor is treated as a black-box.

Mixing updates from multiple PLAD methods is beneficial because, in this joint paradigm, each method can cover the other’s weaknesses. For instance, employing ST ensures that some samples of $\mathbf{X}$ are sourced from S^* , and employing LEST ensures that some samples of $\mathbf{X}$ have paired $\mathbf{Z}$ programs which are exact labels. We present qualitative results in Figure 3. The reconstructions from the PLAD combination methods better reflect the input shapes, reinforcing the quantitative trends.

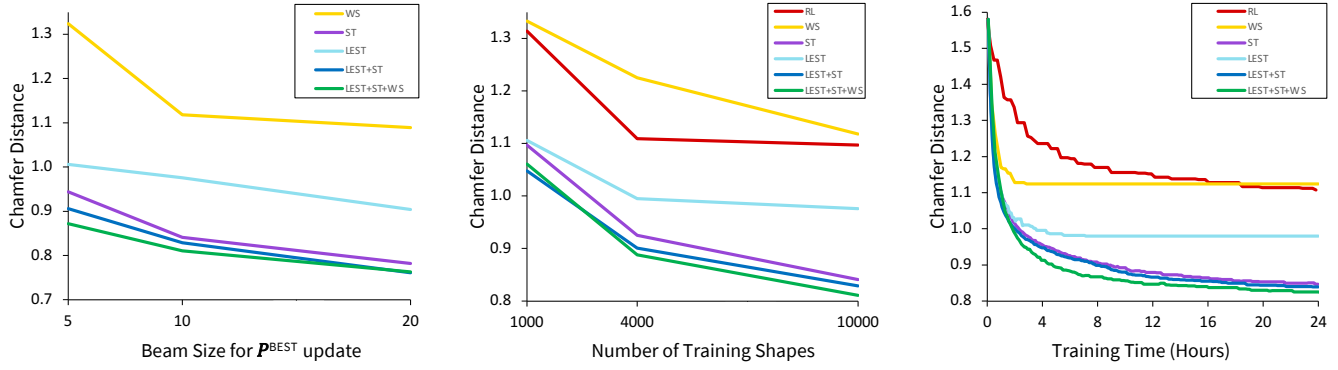


Figure 2. Experiments exploring properties of PLAD methods on 2D CSG. On the X-axis we plot the beam size used during the P^{BEST} update (*Left*), the number of training shapes (*Middle*), and the training time (*Right*). The Y-axis of each plot measures reconstruction accuracy on test-set shapes.

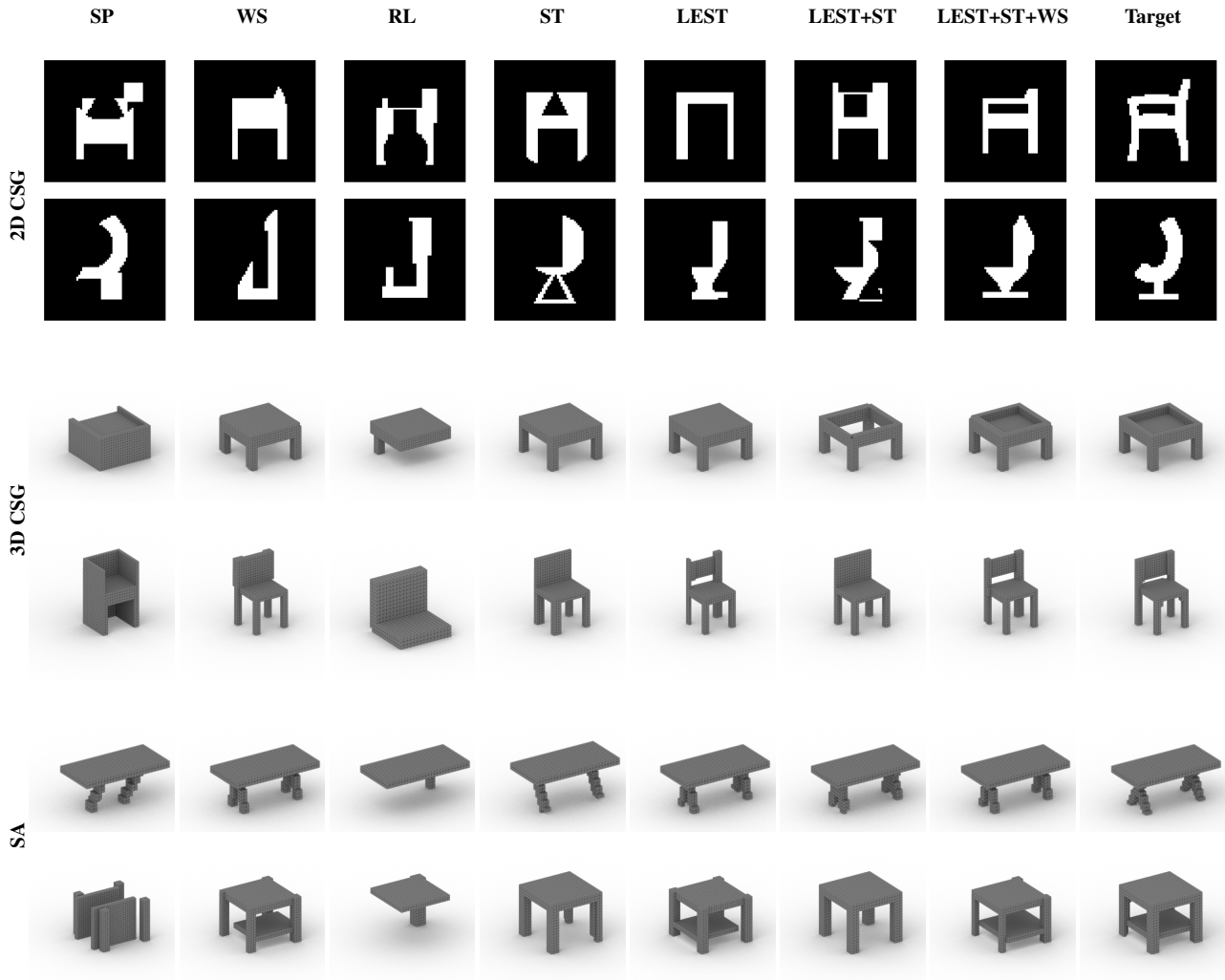


Figure 3. Qualitative comparisons of shape programs inferred for test-set shapes made by different fine-tuning methods for 2D CSG (*Top*), 3D CSG (*Middle*), and ShapeAssembly (*Bottom*). We provide additional qualitative results in the supplemental.

4.4. Inner-loop Search Time

PLAD methods make use of P^{BEST} to generate $(\mathbf{X}, \mathbf{Z})$ datasets that train $p(\mathbf{z}|\mathbf{x})$. To study how time spent on inner-loop search affects each technique, we ran an experiment using different beam sizes to update P^{BEST} . We present results in Figure 2, left. On the X-axis we plot beam size; on the Y-axis we plot test-set reconstruction Chamfer distance. Unsurprisingly, spending more time on the inner-loop search leads to better performance; finding better programs for training shapes improves test time generalization. That said, across all beam sizes, we find that it is always best to train under a combination of PLAD methods; the LEST+ST and the LEST+ST+WS variants always outperform any individual fine-tuning scheme. Note, RL is not included in this experiment because REINFORCE, as defined, has no inner-loop search mechanism. In this way, PLAD provides an additional control lever, where time spent on inner-loop search modulates a trade-off between convergence speed and test-set reconstruction performance.

4.5. Number of Training Shapes from S^*

All the fine-tuning methods make use of a training distribution of shapes that are sampled from S^* . For some domains, the size of available samples from S^* may be limited. We run an experiment on 2D CSG to see how different fine-tuning methods are affected by training data size. We present the results of this experiment in Figure 2, middle. We plot the number of training shapes on the X-axis and test set reconstruction accuracy on the Y-axis. All fine-tuning methods improve as the training size of S^* increases, but once again, combining multiple PLAD methods leads to the best performance in all regimes. This study also demonstrates the sample efficiency of PLAD combinations: LEST+ST and LEST+ST+WS trained on 1,000 shapes achieve better test set generalization than RL trained on 10,000 shapes.

4.6. Convergence Speed

Beyond reconstruction accuracy, we are also interested in the convergence properties of a fine-tuning method. Policy gradient RL is notoriously unstable and slow to converge, which is undesirable. For 2D CSG, we record the convergence speed of each method and present these results in Figure 2, right. We plot reconstruction accuracy (Y-axis) as a function of training wall-clock time (X-axis); all timing information was collected on a machine with a GeForce RTX 2080 Ti GPU and an Intel i9-9900K CPU. All PLAD techniques converge faster than policy gradient RL. For instance, RL took 36 hours to reach its converged test-set CD of 1.097, while LEST matched this performance at 1.1 hours (32x faster) and LEST+ST matched this performance at 0.85 hours (42x faster).

5. Conclusion

We presented the PLAD framework to group a family of techniques for fine-tuning shape program inference models with Pseudo-Labels and Approximate Distributions. Within this framework, we proposed LEST: a self-training variant that creates a shape distribution $\mathbf{X}$ approximating the real distribution S^* by executing inferred latent programs. Experiments on 2D CSG, 3D CSG, and ShapeAssembly demonstrate that PLAD methods achieve better reconstruction accuracy and converge faster than policy gradient RL, the current standard approach for black-box fine-tuning. Finally, we found that combining updates from multiple PLAD methods outperforms any individual technique.

While fine-tuning $p(\mathbf{z}|\mathbf{x})$, PLAD methods construct $(\mathbf{X}, \mathbf{Z})$ sets approximating the statistics of S^* , specializing $p(\mathbf{z}|\mathbf{x})$ towards S^* . As a consequence, $p(\mathbf{z}|\mathbf{x})$ may not generalize as well to shapes outside of S^* ; we explore this phenomenon in the supplemental material. Training a general-purpose inference model for all shapes expressible under the grammar is an interesting line of future work.

While our work focuses on reconstruction quality, producing programs with ‘good’ structure matters just as much, if the program is to be used for editing tasks. Currently, the synthetic pretraining data is the only place where knowledge about what constitutes “good program structure” can be injected. Such knowledge must be expressed in procedural form, which may be harder to elicit from domain experts than declarative knowledge (i.e. “a good program has these properties” vs. “this is how you write a good program”). Finding efficient ways to elicit and inject such knowledge is an important future direction.

Finally, we believe that ideas from the PLAD framework are applicable to a broader class of program inference problems than those we evaluated in this paper. In principle, these approaches can be used to train an approximate inference model $p(\mathbf{z}|\mathbf{x})$ for any domain in which (1) an executor $p(\mathbf{x}|\mathbf{z})$ is available and (2) the executed output $\mathbf{x}$ takes the form of some concrete artifact which can be encoded via neural network (image, audio, text, etc.). For example, one could imagine using PLAD techniques to infer graphics shader programs which produce certain textures or audio synthesizers which sound like certain real-world instruments.

Acknowledgments

We would like to thank the anonymous reviewers for their helpful suggestions. This work was funded in parts by NSF award #1941808 and a Brown University Presidential Fellowship. Daniel Ritchie is an advisor to Geopipe and owns equity in the company. Geopipe is a start-up that is developing 3D technology to build immersive virtual copies of the real world with applications in various fields, including games and architecture

References

- [1] Rudy Bunel, Matthew Hausknecht, Jacob Devlin, Rishabh Singh, and Pushmeet Kohli. Leveraging grammar and reinforcement learning for neural program synthesis. *ICLR*, 2018. [2](#)
- [2] Angel X. Chang, Thomas Funkhouser, Leonidas Guibas, Pat Hanrahan, Qixing Huang, Zimo Li, Silvio Savarese, Manolis Savva, Shuran Song, Hao Su, Jianxiong Xiao, Li Yi, and Fisher Yu. ShapeNet: An Information-Rich 3D Model Repository. *arXiv:1512.03012*, 2015. [1](#)
- [3] Eugene Charniak. Statistical parsing with a context-free grammar and word statistics. In *Proceedings of the Fourteenth National Conference on Artificial Intelligence and Ninth Conference on Innovative Applications of Artificial Intelligence*, AAAI’97/IAAI’97, page 598–603. AAAI Press, 1997. [3](#)
- [4] Siddhartha Chaudhuri, Daniel Ritchie, Jiajun Wu, Kai Xu, and Hao Zhang. Learning Generative Models of 3D Structures. *Computer Graphics Forum*, 2020. [2](#)
- [5] Xinyun Chen, Chang Liu, and Dawn Song. Execution-guided neural program synthesis. In *International Conference on Learning Representations*, 2019. [2](#)
- [6] Zhiqin Chen, Kangxue Yin, Matthew Fisher, Siddhartha Chaudhuri, and Hao Zhang. Bae-net: Branched autoencoder for shape co-segmentation. *Proceedings of International Conference on Computer Vision (ICCV)*, 2019. [5](#)
- [7] Jacob Devlin, Jonathan Uesato, Surya Bhupatiraju, Rishabh Singh, Abdel-rahman Mohamed, and Pushmeet Kohli. Robustfill: Neural program learning under noisy i/o. In *Proceedings of the 34th International Conference on Machine Learning - Volume 70*, ICML’17, page 990–998. JMLR.org, 2017. [2](#)
- [8] Kevin Ellis, Maxwell Nye, Yewen Pu, Felix Sosa, Josh Tenenbaum, and Armando Solar-Lezama. Write, execute, assess: Program synthesis with a repl. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2019. [1](#), [2](#)
- [9] Kevin Ellis, Daniel Ritchie, Armando Solar-Lezama, and Josh Tenenbaum. Learning to Infer Graphics Programs from Hand-Drawn Images. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2018. [2](#)
- [10] Kevin Ellis, Catherine Wong, Maxwell Nye, Mathias Sable-Meyer, Luc Cary, Lucas Morales, Luke Hewitt, Armando Solar-Lezama, and Joshua B. Tenenbaum. Dreamcoder: Growing generalizable, interpretable knowledge with wake-sleep bayesian program learning, 2020. [3](#)
- [11] Junxian He, Jiatao Gu, Jiajun Shen, and Marc’Aurelio Ran-zato. Revisiting self-training for neural sequence generation. In *International Conference on Learning Representations*, 2020. [3](#)
- [12] Luke B. Hewitt, Tuan Anh Le, and Joshua B. Tenenbaum. Learning to learn generative programs with memoised wake-sleep, 2020. [3](#)
- [13] Geoffrey E Hinton, Peter Dayan, Brendan J Frey, and Radford M Neal. The “wake-sleep” algorithm for unsupervised neural networks. *Science*, 268(5214):1158–1161, 1995. [3](#)
- [14] R. Kenny Jones, Theresa Barton, Xianghao Xu, Kai Wang, Ellen Jiang, Paul Guerrero, Niloy J. Mitra, and Daniel Ritchie. Shapeassembly: Learning to generate programs for 3d shape structure synthesis. *ACM Transactions on Graphics (TOG), Siggraph Asia 2020*, 39(6):Article 234, 2020. [2](#)
- [15] Jacob Kahn, Ann Lee, and Awni Hannun. Self-training for end-to-end speech recognition. In *ICASSP 2020-2020 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 7084–7088. IEEE, 2020. [3](#)
- [16] Kacper Kania, Maciej Zieba, and Tomasz Kajdanowicz. Ucs-gnet – unsupervised discovering of constructive solid geometry tree. In *arXiv*, 2020. [1](#), [3](#)
- [17] Diederik P. Kingma and Max Welling. Auto-Encoding Variational Bayes. In *International Conference on Learning Representations (ICLR)*, 2014. [3](#), [4](#)
- [18] Chen Liang, Jonathan Berant, Quoc Le, Kenneth D Forbus, and Ni Lao. Neural symbolic machines: Learning semantic parsers on freebase with weak supervision. In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, volume 1, pages 23–33, 2017. [3](#)
- [19] Yunchao Liu, Zheng Wu, Daniel Ritchie, William T. Freeman, Joshua B. Tenenbaum, and Jiajun Wu. Learning to Describe Scenes with Programs. In *International Conference on Learning Representations (ICLR)*, 2019. [2](#)
- [20] David McClosky, Eugene Charniak, and Mark Johnson. Effective self-training for parsing. In *Proceedings of the Human Language Technology Conference of the NAACL, Main Conference*, pages 152–159, New York City, USA, June 2006. Association for Computational Linguistics. [3](#)
- [21] Radford M. Neal and Geoffrey E. Hinton. A new view of the em algorithm that justifies incremental and other variants. In *Learning in Graphical Models*, pages 355–368. Kluwer Academic Publishers, 1993. [3](#)
- [22] Emilio Parisotto, Abdel-rahman Mohamed, Rishabh Singh, Lihong Li, Dengyong Zhou, and Pushmeet Kohli. Neuro-symbolic program synthesis. *ICLR*, 2017. [2](#)
- [23] H. Scudder. Probability of error of some adaptive pattern-recognition machines. *IEEE Transactions on Information Theory*, 11(3):363–371, 1965. [3](#)
- [24] Gopal Sharma, Rishabh Goyal, Difan Liu, Evangelos Kalogerakis, and Subhransu Maji. CSGNet: Neural Shape Parser for Constructive Solid Geometry. In *IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2018. [1](#), [2](#), [5](#), [6](#)
- [25] Gopal Sharma, Rishabh Goyal, Difan Liu, Evangelos Kalogerakis, and Subhransu Maji. Neural shape parsers for constructive solid geometry. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2020. [2](#), [6](#)
- [26] Mark Steedman, Miles Osborne, Anoop Sarkar, Stephen Clark, Rebecca Hwa, Julia Hockenmaier, Paul Ruhlen, Steven Baker, and Jeremiah Crim. Bootstrapping statistical parsers from small datasets. In *10th Conference of the European Chapter of the Association for Computational Linguistics*, Budapest, Hungary, Apr. 2003. Association for Computational Linguistics. [3](#)
- [27] Shao-Hua Sun, Hyeonwoo Noh, Sriram Somasundaram, and Joseph Lim. Neural program synthesis from diverse demonstration videos. In *Proceedings of the 35th International Conference on Machine Learning*, 2018. [2](#)

- [28] Yonglong Tian, Andrew Luo, Xingyuan Sun, Kevin Ellis, William T. Freeman, Joshua B. Tenenbaum, and Jiajun Wu. Learning to Infer and Execute 3D Shape Programs. In *International Conference on Learning Representations (ICLR)*, 2019. [1](#), [2](#), [3](#)
- [29] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *Proceedings of the 31st International Conference on Neural Information Processing Systems, NIPS’17*, page 6000–6010, Red Hook, NY, USA, 2017. Curran Associates Inc. [6](#)
- [30] Ronald J. Williams. Simple statistical gradient-following algorithms for connectionist reinforcement learning. *Machine Learning*, 8, 1992. [2](#)
- [31] Karl D. D. Willis, Yewen Pu, Jieliang Luo, Hang Chu, Tao Du, Joseph G. Lambourne, Armando Solar-Lezama, and Wojciech Matusik. Fusion 360 gallery: A dataset and environment for programmatic cad construction from human design sequences. *ACM Transactions on Graphics (TOG)*, 40(4), 2021. [1](#)
- [32] Jiajun Wu, Joshua B Tenenbaum, and Pushmeet Kohli. Neural scene de-rendering. In *IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2017. [2](#)
- [33] David Yarowsky. Unsupervised word sense disambiguation rivaling supervised methods. In *33rd Annual Meeting of the Association for Computational Linguistics*, pages 189–196, Cambridge, Massachusetts, USA, June 1995. Association for Computational Linguistics. [3](#)
- [34] Chenghui Zhou, Chun-Liang Li, and Barnabas Poczos. Unsupervised program synthesis for images using tree-structured lstm, 2020. [2](#)
- [35] Barret Zoph, Golnaz Ghiasi, Tsung-Yi Lin, Yin Cui, Hanxiao Liu, Ekin D. Cubuk, and Quoc V. Le. Rethinking pre-training and self-training, 2020. [3](#)