

SYNAPSE: Learning Preferential Concepts from Visual Demonstrations

Sadanand Modak¹, Noah Patton¹, Isil Dillig¹, and Joydeep Biswas¹

The University of Texas at Austin, TX USA
 {smodak11, npatt, isil, joydeepb}@cs.utexas.edu

Abstract. This paper addresses the problem of *preference learning*, which aims to learn user-specific preferences (e.g., “good parking spot”, “convenient drop-off location”) from visual input. Despite its similarity to learning *factual* concepts (e.g., “red cube”), preference learning is a fundamentally harder problem due to its subjective nature and the paucity of person-specific training data. We address this problem using a new framework called SYNAPSE, which is a neuro-symbolic approach designed to efficiently learn preferential concepts from limited demonstrations. SYNAPSE represents preferences as neuro-symbolic programs in a domain-specific language (DSL) that operates over images, and leverages a novel combination of visual parsing, large language models, and program synthesis to learn programs representing individual preferences. We evaluate SYNAPSE through extensive experimentation including a user case study focusing on mobility-related concepts in mobile robotics and autonomous driving. Our evaluation demonstrates that SYNAPSE significantly outperforms existing baselines as well as its own ablations. The code and other details can be found on the project website <https://amrl.cs.utexas.edu/synapse>.

Keywords: Concept learning · Neuro-symbolic programming · Program Synthesis · Visual Reasoning

1 Introduction

Imagine trying to come up with a definition of “a good taxi drop-off location”. One person may consider a spot to be a good drop-off location depending on whether it is close to the door of a building, while someone else might want it in the shade. Such concepts vary from person to person and inherently depend on their preferences. We call them *preferential concepts*, and we are interested in the problem of *preference learning* from visual input. Learning preferences is important because we want systems that are customizable and that can adapt to end-users. This problem is quite related to the task of visual concept learning, wherein much of the work focuses on learning concepts such as *having the color red* or *being to the left of another object* [5, 13, 23, 25, 26, 31, 33, 35, 40, 41, 43, 44, 49, 52, 55, 61]. All such prior work assumes there is a *ground-truth* for the concept, *i.e.*, the definition of the concept does not differ among people, and as

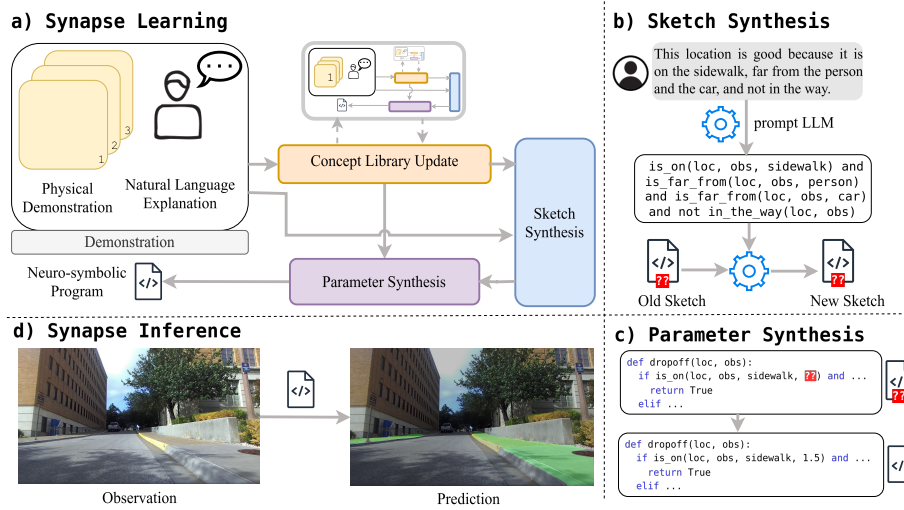


Fig. 1: Illustration of SYNAPSE for identifying good contingency locations. (a) **SYNAPSE** learns a neuro-symbolic program that represents the preferential concept based on user demonstrations, which include both a NL explanation and a physical demonstration. The learning algorithm consists of three steps, namely updating the concept library, synthesizing a program sketch, and performing parameter synthesis; (b) The new sketch gets synthesized based on the previous sketch and the current natural language description; (c) Parameter synthesis determines suitable values for numeric parameters in the program sketch; (d) **SYNAPSE** evaluates the program on a new query image to return a preference (in this case, boolean) mask over the input image.

a consequence, sufficiently many examples are available, and can be objectively evaluated. We refer to such concepts as *factual concepts*. While most prior work that learns visual concepts exploits the availability of large datasets such as CLEVR [29], those methods cannot be applied to preference learning because it is a data-impovertised setting by its very nature: a single individual can put up with providing only so much data. This limitation is also present in most of the preference learning work in the reinforcement learning literature as well, where human preferences are represented as neural networks [10, 59] or latent reward models [2, 7, 14, 15, 48, 60]. Furthermore, because preferences are inherently individual, they can depend on entirely different concepts, such as in the drop-off location example above (*i.e.*, based on *proximity to door* as opposed to *being in the shade*). This requires learning novel visual concepts in a *hierarchical* manner. Lastly, coming up with a complete definition of a preferential concept at once is itself a hard problem: it is much easier for someone to *show* examples that satisfy their intuition as humans tend to *build* their notion of a preferential concept over time. Thus, preference learning calls for an approach that can handle *incremental learning* from visual demonstrations.

To address these challenges, we present SYNAPSE, a novel framework that learns human preferences in a data-efficient manner. In contrast to prior preference learning approaches [2, 7, 10, 14, 15, 19, 48, 59, 60, 64] which take in weak reward signals to learn preferences, we use a more direct form of a preference signal, which consists of a physical demonstration including visual data, and a natural language (NL) explanation for the preference. We use NL input from the human to identify new concepts to be learned as well as how to compose them. However, in addition to learning new concepts or composing existing ones, preferences also have a quantitative aspect. For example, to be a good drop-off spot, you should be close to a door, but exactly how close is a personal preference. This is where the demonstrations come into play and allow us to infer quantitative aspects of the preference that are hard to capture via natural language alone. Finally, to allow *incremental, data-efficient* learning, SYNAPSE expresses preferential concepts as programs in *neuro-symbolic* domain specific language (DSL) operating over images, and learns these programs based on demonstrations. Such a programmatic representation also facilitates *life-long learning*, allowing incremental changes to the learnt program as new demonstrations arrive.

Fig. 1 shows a schematic of our proposed SYNAPSE framework. Given a user demonstration (*i.e.*, the physical demonstration and NL input), the general workflow of SYNAPSE has three main components: First, SYNAPSE leverages the user’s NL explanation, along with SYNAPSE’s existing *concept library*, to ground the visual *concepts* needed to represent the user’s preference. If the NL explanation contains auxiliary concepts that are not part of SYNAPSE’s existing concept library, SYNAPSE may query the user for additional demonstrations of the auxiliary concept, which are then used to update SYNAPSE’s concept library. Once the library contains all required concepts, SYNAPSE uses the NL explanation to generate a so-called *sketch* which is a program in our DSL with missing values for numeric parameters. Finally, SYNAPSE uses constrained optimization techniques (based on *maximum satisfiability* [24]) to find values of the numeric parameters that are maximally consistent with the user’s physical demonstrations.

To demonstrate the effectiveness of our framework, we evaluate it on three mobility-related visual preferential concepts which find their use in mobile robotics and autonomous driving: a) **CONTINGENCY**: What is a good spot for a robot to pull over to in case of contingency?, b) **DROPOFF**: What is a good location for an autonomous car to stop and drop-off the customer?, and c) **PARKING**: Where can the autonomous car be parked?. We carry out extensive experiments, including a user study, spanning multiple baselines. Empirical results show that SYNAPSE outperforms the baselines by a significant margin, especially when evaluated on out-of-distribution data — even when SYNAPSE is trained on an order of magnitude fewer examples than baseline Neural-Network (NN) approaches. A case study with multi-user preferences demonstrates that our method can learn personalized preferences unique to each user with just a few demonstrations.

In summary, this paper contributes:

1. SYNAPSE, a neuro-symbolic framework to learn and evaluate preferences

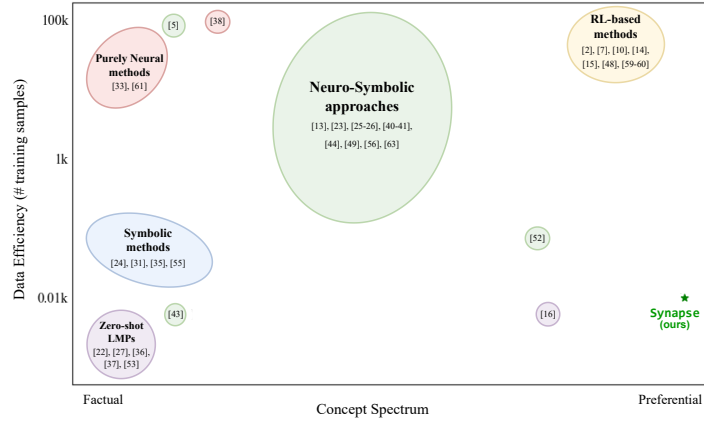


Fig. 2: A taxonomy of (visual) concept learning and VQA

2. A new method for hierarchical lifelong learning from visual demonstrations and natural language
3. A comprehensive experimental evaluation of the proposed approach, including a multi-user case study and comparison against baselines and ablations

2 Related work

Our framework positions itself in the larger field of concept learning and visual question answering (VQA), a broad taxonomy depicted in Fig. 2. While there exist Reinforcement-Learning based methods for preference learning, most of them fall in the imitation learning setting, where human preferences are represented via neural policies [10, 59] or latent reward models [2, 7, 14, 15, 48, 60]. Further, they do not deal with natural language, but rather take some form of weak preference signal as input. In the following discussion, we focus on work that is most closely-related to SYNAPSE.

Language Model Programs (LMPs). Generating executable programs from natural language is not a new idea. Many earlier works [5, 13, 23, 35, 40, 41, 44, 52, 63] use custom semantic parsers to perform specific tasks. However, with the advent of Large Language Models (LLMs) such as GPT-4 [1], LMPs have gained significant attention [16, 22, 25–27, 37, 50, 53, 54, 56] due to the extensive knowledge that these foundation models possess. Code-as-Policies [36] pioneered the effort in this direction and demonstrated that LLMs can generate simple Python programs for tasks ranging from drawing shapes to tabletop manipulation. This approach uses recursive prompting as a strategy to get rid of invalid function calls. Voyager [55] builds an LLM-powered embodied agent that can learn a diverse set of skills in a lifelong manner in the game of Minecraft.

Neuro-symbolic visual learning and reasoning. Neuro-symbolic approaches couple the interpretability of rule-based symbolic AI with the flexibility of neural networks. This idea of composing neural networks into classical symbolic programs to learn a particular concept for VQA dates back to Neural Module Networks (NMN) [5]. NMNs compose neural modules into an end-to-end differentiable network, which are then jointly trained. While methods like LLaVa [38] and GPT-4V aim to tackle visual-language problems in a completely neural way, works such as VisProg [22] and ViperGPT [53] build on top of these and take the zero-shot route by composing the available pretrained models. One work [52] uses a trained semantic parser to first extract useful feature definitions from a few statements describing the concept, which are then evaluated for each datapoint to build feature vectors on which standard classification can be done. NS-VQA [63] is another approach that uses a separately trained visual parser to generate a structured representation of objects in the image and a semantic parser trained to parse the question into a predefined DSL format, which is followed by Pythonic program execution to generate the answer. NS-CL [40] uses the same framework as NS-VQA [63], but instead of answering questions given the trained modules, it represents concepts as neural operators and tries to learn them given the question-answer pairs. DCL [13] goes one step further and extracts keyframes from video and tracks object trajectories to build latent feature vectors that describe dynamic concepts such as collision. Similarly, a host of other methods have been developed [23, 25, 26, 31, 33, 35, 41, 43, 44, 49, 55, 61], of which a few [31, 35] also use learning from demonstrations such as GUI feedback in a mobile app to guide interactive learning.

Program synthesis. There is a rich literature on synthesizing programs from user-provided specifications in the programming languages community [3, 4, 9, 11, 12, 17, 18, 20, 21, 28, 30, 45, 51, 57, 58]. Closer to our application domain, LDIPS [24] is a recent approach that makes use of the dimensionality of variables to synthesize Action Selection Policies (ASPs) given user demonstrations. From the viewpoint of visual reasoning and concept learning, [43] tries to solve Visual Discrimination Puzzles (VDP) by first constructing scene graphs for all images in the puzzle, and then synthesizing a discriminator expressed in first-order logic by performing a full-blown discrete search. However, this can quickly become inefficient as problem size scales. To tackle this, our method uses natural language informed sketch generation and performs synthesis over the space of parameters.

3 Method

We define a *preference task* $\mathcal{T} := \langle O, Q, P \rangle$ as a tuple consisting of an observation space O , a query space Q , and a preference space P . A preference evaluator π accepts an observation and a query, and returns a preference value: $\pi : O \times Q \rightarrow P$. The goal of *preference learning* is to synthesize a suitable evaluator π that accurately predicts a person’s visual preferences given an image. As stated earlier, a distinguishing feature of preference learning is that it must be performed using

Inputs	Constants	Terms
$q \in Q$ $o \in O$	$v \in \{\text{Int, Real}, \dots\}$ $p \in P$	$t := q \mid o \mid v$ $\mid f(t_1, \dots, t_n) \text{ where } f \in \mathcal{C}_f$
Conditions	Programs	
$\phi := p_c(t_1, \dots, t_n) \text{ where } p_c \in \mathcal{C}_b$ $\mid \neg\phi \mid \phi \wedge \phi \mid \phi \vee \phi$	$\pi := p \mid \text{if } (\phi) \text{ then } \pi \text{ else } \pi'$	

Fig. 3: The SYNAPSE neuro-symbolic DSL parametrized over concept library $\mathcal{C}$, which consists of a set of predicates $\mathcal{C}_b$ and functions $\mathcal{C}_f$.

small amounts of training data due to the subjective nature of preferences. To enable data-efficient learning, our proposed SYNAPSE approach represents the preference evaluator π as a neuro-symbolic program in a DSL and synthesizes π from a small number of user demonstrations, where each demonstration includes a sequence of images along with a natural language explanation for the user’s preference. In the following discussion, we first present our DSL for representing preferences and then describe our learning algorithm.

3.1 Representing Preferences

We represent preference evaluation functions π (or preferences for short) in the neuro-symbolic DSL shown in Fig. 3. This DSL is parameterized over a so-called *concept library* $\mathcal{C}$, which includes both predicates $\mathcal{C}_b$ as well as non-boolean functions $\mathcal{C}_f$. The concept library includes built-in operators and predicates (e.g., $+$, $\times$, $\leq$, $\dots$), pre-trained neural networks (e.g., for object classification and terrain detection), zero-shot visual language models (VLMs), as well as previously learned concepts and functions (expressed in the same DSL). We sometimes use the notation $\pi_{\mathcal{C}}$ to denote programs using concept library $\mathcal{C}$ and omit the subscript $\mathcal{C}$ where it is clear from context.

At a high level, a program π consists of (nested) if-then-else statements and is therefore conceptually similar to a decision tree. Each leaf of this decision tree is a preference (e.g., *good*, *bad*, *neutral*) drawn from the preference space P , which is assumed to be a finite set. Internal nodes of the decision tree are neuro-symbolic conditions ϕ , which include boolean combinations of predicates of the form $p(t_1, \dots, t_n)$ where each t_i is a neuro-symbolic term and p is a predicate drawn from $\mathcal{C}_b$, which could be a built-in relation (e.g., $\leq$), result of a neural classifier, or a previously-learned concept (e.g., *close-to*).

3.2 Learning Preferences

We now discuss our learning algorithm, **SYNAPSE-Learn**, for synthesizing preference evaluation functions from a set of demonstrations $\mathcal{D}$. As **SYNAPSE-Learn** is meant to be used in a life-long-learning setting, we present it as an incremental algorithm that takes one additional demonstration in each invocation and returns

Algorithm 1: SYNAPSE-Learn

Input: *a set of previously seen demonstrations $\mathcal{D}$, the new demonstration d_n , a potentially empty previous sketch $\hat{\pi}_o$, a previous concept library $\mathcal{C}$*
Output: *The new demonstrations set $\mathcal{D}'$, a neuro-symbolic preference evaluator π parameterized by the new concept library $\mathcal{C}'$, the sketch $\hat{\pi}$ used to generate π*

```

1:  $\text{Learn}(\mathcal{D}, d_n, \hat{\pi}_o, \mathcal{C})$ 
2:   # Update concept library with new natural language utterance
3:    $\mathcal{C}' \leftarrow \text{UpdateConceptLibrary}(d_n.e, \mathcal{C})$ 
4:   # Get the updated sketch from the new demonstration and previous sketch
5:    $\hat{\pi} \leftarrow \text{SketchSynth}(d_n, \hat{\pi}_o, \mathcal{C}')$ 
6:   # Fill the holes in  $\hat{\pi}$  based off of the demonstrations  $\mathcal{D}'$ 
7:    $\pi \leftarrow \text{ParamSynth}(\hat{\pi}, \mathcal{D} \cup \{d_n\})$ 
8:   return  $\mathcal{D} \cup \{d_n\}, \pi, \mathcal{C}', \hat{\pi}$ 

```

Algorithm 2: Update Concept Library

Input: *A new natural language explanation e and the previous concept library $\mathcal{C}$*
Output: *A new concept library $\mathcal{C}'$*

```

1:  $\text{UpdateConceptLibrary}(e, \mathcal{C})$ 
2:    $\mathcal{C}' \leftarrow \mathcal{C}$  # Initialize new concept library with old concept library
3:   # Extract new visual groundings from natural language and
4:   # add them to concept library
5:    $g \leftarrow \text{ExtractEntities}(e, \mathcal{C})$ 
6:    $\mathcal{C}' \leftarrow \mathcal{C}' \cup g$ 
7:   # Extract new predicates from  $e$ 
8:    $\text{Preds} \leftarrow \text{ExtractPredicates}(e, \mathcal{C})$ 
9:   # Recursively update concepts with user feedback
10:  for  $\text{pred} \in \text{Preds}$  where  $\text{pred} \notin \mathcal{C}$ 
11:     $\mathcal{D} \leftarrow \emptyset$  #Empty initial demonstration set
12:     $\mathcal{C}'' \leftarrow \mathcal{C}'$ 
13:     $\hat{\pi} \leftarrow \text{None}$ 
14:    do
15:       $d \leftarrow \text{QueryUserForDemonstration}(\text{pred})$ 
16:       $\mathcal{D}, \pi, \mathcal{C}'', \hat{\pi} \leftarrow \text{Learn}(\mathcal{D}, d, \hat{\pi}, \mathcal{C}'')$ 
17:    while  $d$ 
18:       $\mathcal{C}' \leftarrow \mathcal{C}''$ 
19:  return  $\mathcal{C}'$ 

```

an updated preference evaluation function. As mentioned earlier, we represent each demonstration d as a pair (t, e) where t is a physical demonstration consisting of a sequence of images and LiDAR point clouds and e is a natural language explanation for the preference. Given a demonstration d , we write $d.t$ and $d.e$ to denote its physical demonstration and explanation component respectively.

In addition to the new demonstration d_n , **SYNAPSE-Learn** takes three additional arguments, namely the previous set of demonstrations $\mathcal{D}$, the previously learnt preference evaluation function π_o (**None** for the first invocation), and the current concept library $\mathcal{C}$, which is initialized to contain only a set of built-in concepts. **SYNAPSE-Learn** uses the old program π_o to bootstrap the learning process, and the previous demonstrations are required to ensure that the updated program is consistent with *all* demonstrations provided thus far.

At a high level, the learning procedure consists of three steps, which are explained in more detail in the remainder of this section:

1. **Updating the concept library:** SYNAPSE first checks whether the existing concept library $\mathcal{C}$ is sufficient for successfully learning the desired preference

Algorithm 3: Parameter Synthesis

Input: A program sketch $\hat{\pi}$, a set of demonstrations $\mathcal{D}$
Output: A complete program π

```

1: ParamSynth( $\hat{\pi}, \mathcal{D}$ )
2:    $\varphi \leftarrow \text{true}$  # Initialize condition
3:   for  $d \in \mathcal{D}$ 
4:     # Perform partial evaluation on the sketch and demonstration
5:     # to get a simplified sketch  $\hat{\pi}'$  and expected result  $r$ 
6:      $(\hat{\pi}', r) \leftarrow \text{PartialEval}(\hat{\pi}, d)$ 
7:     # Merge with condition
8:      $\varphi \leftarrow \varphi \wedge \llbracket \hat{\pi}' \rrbracket^r$ 
9:     # Include negation for each parameter  $\in P$ 
10:    for  $i \in P$ 
11:      if  $(i \neq r)$   $\varphi \leftarrow \varphi \wedge \neg \llbracket \hat{\pi}' \rrbracket^i$ 
12:    # Use solver to fill holes over symbolic features
13:     $\pi \leftarrow \text{Model}(\varphi)$ 
14:  return  $\pi$ 

```

evaluation function. For example, if the natural language explanation uses the term “far away” but the concept library does not contain a suitable definition, **SYNAPSE-Learn** interactively queries the user for clarification and updates its concept library as needed.

2. **Synthesizing a program sketch:** If the concept library is sufficient for representing the preference, **SYNAPSE-Learn** proceeds to synthesize a so-called *program sketch*, which is a program with missing constants to be synthesized. We differentiate between program sketches and complete programs because the user’s natural language explanation is often sufficient to understand the general structure of the preference evaluation function but not its numeric parameters, which can only be accurately learned from the physical demonstrations. Thus, **SYNAPSE-Learn** generates the program sketch based only on the natural language explanation.
3. **Parameter synthesis:** The final phase of the learning algorithm utilizes all physical demonstrations provided thus far to synthesize the unknown numeric parameters of the sketch using a constraint-solving approach. For example, if the user’s NL explanation mentions “not too close to the sidewalk”, the physical demonstrations are needed to understand what the user considers “too close”. For this reason, **SYNAPSE-Learn** utilizes a separate parameter synthesis procedure to determine suitable numeric parameters from the physical demonstrations.

Concept Library Update. SYNAPSE analyzes the user’s natural language explanation e to extract concepts of interest. We differentiate between two types of concepts: (1) entities (e.g., car, door, sidewalk) and (2) predicates (e.g., far, near). Because SYNAPSE uses an open-vocabulary visual language model to find entities of interest in the current observation, new entity concepts do not require interacting with the user. On the other hand, if the natural language explanation contains new predicates that are not part of the existing concept library, SYNAPSE needs to query the user to provide suitable demonstrations.

Algorithm 2 summarizes this discussion as an algorithm. As shown in lines 5-6, the **ExtractEntities** procedure grounds the entities used in the NL description and cross-references them against existing entities in the concept library. Any new entities are added to the concept library without requiring user interaction, as we assume that any entity can be extracted from the observation using a modern VLM. Lines 8-17, on the other hand, extract new *predicates* from the natural language description. Since the semantics of these predicates are not known a priori (unless they are already in the concept library), we must query the user to learn their semantics. Thus, the **QueryUserForDemonstration** procedure obtains new demonstrations, which are then used to synthesize the implementation of the new predicate through recursive invocation of **SYNAPSE-Learn** at line 15. Thus, when the **UpdateConceptLibrary** procedure terminates, the new concept library $\mathcal{C}'$ contains all entities and predicates of interest.

Program Sketch Synthesis. Once SYNAPSE has all the required concepts as part of its library, it uses a large language model to synthesize a suitable program sketch based on the natural language description and concept library. In particular, it first prompts the LLM to translate the NL explanation e to a pair (Φ, p) where Φ is a formula (in conjunctive normal form) over the predicates in the concept library and p is the user’s preference. Then, in a second step, SYNAPSE prompts the LLM to update the previous sketch $\hat{\pi}$ to a new one $\hat{\pi}'$ such that $\hat{\pi}'$ returns p when Φ evaluates to **True**. We found this two-stage process of first converting the NL explanation to a CNF formula and then prompting the LLM to repair the old sketch to work better in practice compared to prompting the LLM directly with all inputs (see Section 4).

Parameter Synthesis. As mentioned earlier, a program sketch contain unknown numeric parameters that arise from the ambiguity of natural language (e.g., what does “close” mean in terms of distances between objects?) Thus, the last step of the SYNAPSE pipeline utilizes the user’s physical demonstrations to synthesize numeric parameters in the sketch. Our parameter synthesis algorithm is summarized in Algorithm 3 and constructs a logical formula φ whose models are guaranteed to be consistent with all demonstrations. The algorithm constructs this formula φ , which is initialized to true at line 2, as follows: First, for each physical demonstration d , it partially evaluates $\hat{\pi}$ by concretely evaluating all expressions without unknowns. For example, if the sketch contains the predicate **distanceTo(car)**, we can use the observation from d to compute the actual distance between the subject and the car. This partial evaluation (line 6) yields a much simpler sketch containing only unknowns to be synthesized but no other variables. Now, given a sketch $\hat{\pi}$, let $\llbracket \hat{\pi} \rrbracket^i$ denote the condition under which $\hat{\pi}$ returns preference $p_i \in P$, and suppose that the current demonstration d illustrates preference class p_r . Since we would like the synthesized program to return p_r for demonstration d , $\llbracket \hat{\pi} \rrbracket^r$ should evaluate to true, and for all other preference classes p_i where $i \neq r$, $\llbracket \hat{\pi} \rrbracket^i$ should evaluate to false. Thus, the loop in lines 10-12 iteratively strengthens formula φ by conjoining it with $\llbracket \hat{\pi} \rrbracket^r$ and

the negation of $\llbracket \hat{\pi} \rrbracket^i$ for any i distinct from r . Finally, we use an off-the-shelf constraint solver to obtain a model of the resulting formula.¹ By substituting the unknowns in the sketch with the model returned by the solver, SYNAPSE obtains a program that is maximally consistent with the user’s demonstrations.

4 Evaluation

To test the effectiveness of SYNAPSE, we evaluate it on three mobility-related preferential concepts relevant to mobile robotics and autonomous driving domains: a) **CONTINGENCY**: *What is a good spot for a robot to pull over to in case of an emergency?*, b) **DROPOFF**: *What is a good location for an autonomous taxi to stop and drop-off a customer?*, and c) **PARKING**: *What is a good location for parking an autonomous car?*. The following subsection gives some details on how we ground SYNAPSE to learn such mobility-related concepts. Further information is provided in supplementary material.

4.1 Implementation Details

Inputs. Human demonstrations include the robot trajectories of the user driving the robot to the desired location using a joystick based on the user’s preference, and a natural language description to explain the rationale for choosing that location. We use RGB-camera images as well as the LIDAR sensor data for the entire trajectory to construct a symbolic representation of the scene.

Models. We use Grounded-SAM [32, 39, 46] zero-shot VLM for object detection and a custom terrain model with SegFormer architecture [62] trained to segment terrains from sensor data. We use both of these models to get segmentation masks of the neural concepts (*i.e.*, objects and terrains) currently in the concept library. We use GPT-4 [1] as the language model for sketch synthesis in SYNAPSE.

Concept Library Initialization. We impart some basic capabilities to the learning framework by seeding the concept library with a total of six functions based on neural models (*i.e.*, terrain model and VLM) as well as some domain-specific predicates. For example, `terrain_at` returns the terrain type at a specific pixel location on a provided image, while `project_map_to_pixel` transforms real-world 3D coordinates to pixel coordinates.

Querying the user. As mentioned earlier in Sec. 3.2, SYNAPSE can interactively query the user to clarify new concepts that are present in the user’s NL explanation but not in the current concept library. In principle, SYNAPSE can

¹ In general, the demonstration may be noisy, meaning that φ could be unsatisfiable. Since this is quite often the case, our implementation uses a MaxSMT solver [8] to find a solution that maximizes the number of satisfied clauses in the formula.

query the user for both physical demonstrations and NL explanations. To reduce the burden on the user, SYNAPSE, by default, only queries the user for NL explanations of auxiliary concepts and performs synthesis of auxiliary concepts using NL explanations alone.

4.2 Results

To evaluate SYNAPSE, we create a dataset of 815 labeled images taken from the UT Austin campus area, where the labels mark the locations on the images that are consistent with the intended user preference for each of the three tasks. We split the dataset into three sets: train, in-distribution test, and out-of-distribution test sets. The train and in-distribution sets belong to the same part of the UT Austin campus region, while the out-of-distribution set belongs to a different part of the region. Table 1 shows the comparison against various baselines. We use mean Intersection-Over-Union (mIOU) as the metric and evaluate the following baselines: (1) pure neural models based on SegFormer [62] architecture with pretrained weights, with and without depth input, fine-tuned on our custom dataset; (2) GPT4 [1] with vision capabilities, and (3) VisProg [22]. The ‘+’ for the latter two indicate additional prompting that provides additional information about the underlying reasoning behind the preferential concept. More details about the dataset and how different baselines are trained and/or evaluated can be found in the supplementary material.

We find that SYNAPSE outperforms all baselines, and improves on the closest baseline by a significant margin on out-of-distribution test data — 74.07 *vs.* 57.42 for CONTINGENCY, 80.72 *vs.* 55.04 for DROPOFF, and 62.75 *vs.* 52.90 for PARKING. Further, even though SYNAPSE is trained on an order of magnitude fewer samples (for instance, 29 demonstration for CONTINGENCY) than neural baselines (for instance, 224 images for CONTINGENCY), it performs at-par, if not better, on the train dataset.

4.3 Ablations

We investigate three types of ablations: (1) NN-ablations, in which we compare the performance of neural baselines against SYNAPSE when trained on the *same* number of samples (*i.e.*, 29), (2) LLM-based ablations, where we see how using different LLMs affect the performance of SYNAPSE, and (3) framework ablations where we test different design choices. We investigate ablations based on these framework features: (1) *feat1*: whether it queries the user for auxiliary concepts, (2) *feat2*: whether it performs lifelong learning by building on its concept library, and (3) *feat3*: whether it has got *direct access* to the full concept library learned by SYNAPSE. More details are provided in the supplementary material.

Table 2 show how the ablation methods differ based on these features as well their performance when evaluated on the CONTINGENCY dataset. It can be seen that the neural ablations perform poorly since they are only exposed to so few training samples that they aren’t able to generalize well to the full dataset. Changing the program synthesis process of SYNAPSE (*i.e.*, doing direct synthesis

Table 1: Mean IOU (%) results for the three concepts. The train set represents the training set for neural baselines – SYNAPSE only needs 29 demonstrations (from the train area), and the visual-language baselines have not been fine-tuned.

	CONTINGENCY			DROPOFF			PARKING	
	train	in-test	out-test	train	in-test	out-test	train	out-test
Synapse	77.64	76.29	74.07	79.32	80.18	80.72	68.60	62.76
SF-RGB-b0	70.49	62.75	57.42	72.99	68.13	52.21	57.68	49.66
SF-RGB-b5	74.59	70.48	56.00	77.26	72.83	55.04	68.63	52.91
SF-RGBD-b0	72.17	67.23	54.75	74.33	69.80	54.90	60.90	50.69
SF-RGBD-b5	76.48	67.81	56.11	77.69	70.70	52.39	71.06	49.99
GPT4V	26.56	29.71	30.45	30.54	32.97	37.65	38.41	40.18
GPT4V+	28.73	28.96	33.92	39.38	38.34	39.14	41.38	39.77
VisProg	45.62	45.63	44.98	46.29	47.49	48.07	39.43	40.99
VisProg+	38.94	39.21	41.83	39.17	39.44	43.14	38.88	38.99

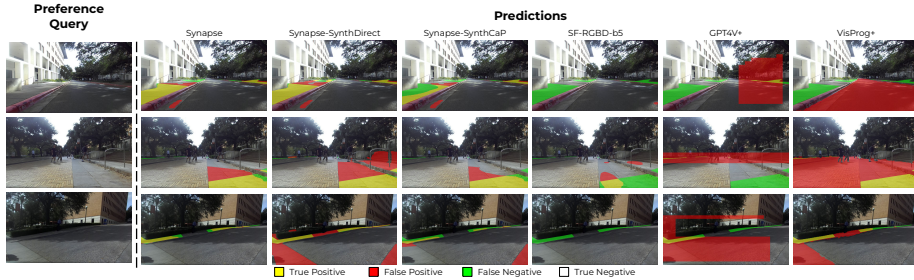


Fig. 4: An illustrative comparison between SYNAPSE, baselines, and ablations for CONTINGENCY. Color coding shows the overlap of the predictions with the *ground-truth*.

instead of two-step) or the LLM which in turn affects the accuracy of the program sketch being synthesized, also has a significant impact on the performance. Examples showing the difference in inference outputs for a few of the baselines and the ablations is shown in Fig. 4.

4.4 Case study with multi-user preferences

We also perform a case study² to evaluate the alignment and customizability characteristics of SYNAPSE with inputs from different users. We ask three participants to provide demonstrations for their preferences in the CONTINGENCY task. We also each participant to label a set of 15 images with their *ground-truth* preference to be used for evaluation, which are not seen by the algorithm during training. The results are summarized in Tab. 3. For each user, the highest performance is attained by the program that learned from the same user’s demonstrations, which indicates good alignment. The results also indicate that

² Reviewed by IRB and the determination was ‘Not Human Research’

Table 2: The results for the ablation studies. It highlights the importance of different elements of SYNAPSE. All methods trained on 29 images (for neural) or 29 demonstrations (for SYNAPSE-based). Evaluation is on the full CONTINGENCY dataset.

	feat1	feat2	feat3	LLM	mIOU (%)
Synapse	✓	✓	✗	GPT-4	76.11
Synapse-SynthDirect	✗	✗	✗	GPT-4	60.74
Synapse-SynthDirect+	✗	✗	✓	GPT-4	68.86
Synapse-SynthCaP	✗	✓	✗	GPT-4	64.11
Synapse-CodeLLama	✓	✓	✗	CodeLLama [47]	69.88
Synapse-StarCoder	✓	✓	✗	StarCoder [34]	63.62
Synapse-PaLM2	✓	✓	✗	PaLM2 [6]	71.62
SF-RGB-b0	-	-	-	-	54.58
SF-RGB-b5	-	-	-	-	56.30
SF-RGBD-b0	-	-	-	-	55.84
SF-RGBD-b5	-	-	-	-	63.81

Table 3: The results (mean IOU (%)) for case study. Rows represent the learned programs for the particular user, and columns represent the *ground-truth* for that user.

	user-1	user-2	user-3
user-1	76.85	66.29	67.74
user-2	65.66	73.75	69.28
user-3	65.46	68.97	73.82

the preference programs trained on one user’s input performs poorly on a different user’s evaluation set, indicating that there are indeed user-specific preferences captured by SYNAPSE.

4.5 Reordering of demonstrations

We investigate if SYNAPSE is susceptible to performance degradation if the order of demonstrations is altered. For this evaluation, we provide the demonstrations in a randomized order for the CONTINGENCY concept and then evaluate the learned program. The variations in the mean IOU are shown in Fig. 5. It is clearly visible that after about 10-15 demonstrations, the effect of re-ordering diminishes which shows the robustness of SYNAPSE.

5 Discussion of Limitations

Our experimental results demonstrate that SYNAPSE significantly outperforms baselines when it comes to generalizing to new data, which is crucial for achieving transferability of learned models. However, SYNAPSE has three potential limitations that we discuss below.

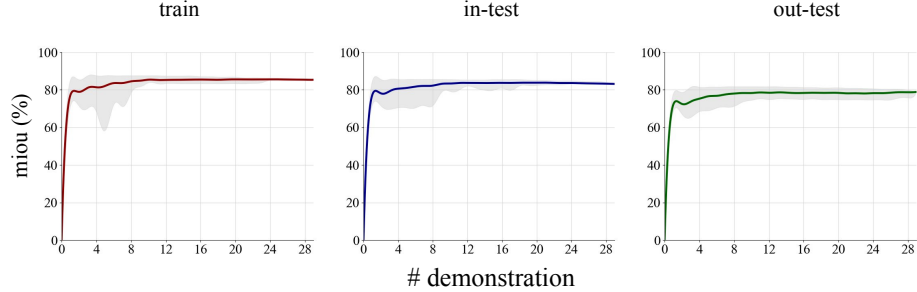


Fig. 5: Plot showing susceptibility of SYNAPSE to reordering of demonstrations. *Gray* area represents the mean IOU (%) variation as SYNAPSE sees more demonstrations.

First, SYNAPSE relies substantially on how good the underlying neural modules are and their capabilities, specifically, the zero-shot LLMs and VLMs. In our experiments, we observe that a careful selection of parameters and verbose prompting is needed to achieve best performance. Additionally, due to the poor performance of existing VLMs for terrain detection, our implementation of the SYNAPSE framework uses a custom terrain model.

Second, SYNAPSE relies on the quality of the user’s physical demonstrations for accurate parameter synthesis. In practice, the demonstrations may be noisy and imperfect, and SYNAPSE tries to compensate for slight inconsistencies in the user demonstration by using a constrained optimization approach based on MaxSMT. While poor-quality demonstrations can still negatively impact overall accuracy, our experiments show that SYNAPSE is effective for demonstrations provided by participants in our user study and that it is not particularly sensitive to the order of demonstrations. Third, another limitation of operating on real-world data with SYNAPSE is the incomplete depth information – SYNAPSE approximates depth at all locations by interpolation, however, that introduces noise into the quantitative evaluation. Better approaches for scene completion [42] would reduce such noise.

6 Conclusion

We presented SYNAPSE, a data-efficient, neuro-symbolic framework for learning preferential concepts from a small number of human demonstrations. The framework utilises a novel combination of visual parsing, large language models, and program synthesis to represent preferences as interpretable programs that can be synthesized from demonstrations. We experimentally showed that SYNAPSE achieves strong generalization on new data and that it outperforms the baselines by a large margin ($\approx 15\%$ mIOU). Further, we demonstrated that SYNAPSE is able to align well with user preferences and that it is robust to the order in which demonstrations are given. Finally, we demonstrated the importance of key design choices underlying SYNAPSE through ablation studies.

Acknowledgement

This work is supported in part by NSF awards OIA-2219236, CCF-1762299, CCF-1918889, CNS-1908304, CCF-1901376, CNS-2120696, CCF- 2210831, and CCF-2319471, as well as Amazon and JP Morgan. We thank the members of Autonomous Mobile Robotics Laboratory (AMRL) at UT Austin for participating in the user case study. We are also grateful to the amazing Hugging Face ecosystem for simplifying the use of state-of-the-art neural models.

References

1. Achiam, J., Adler, S., Agarwal, S., Ahmad, L., Akkaya, I., Aleman, F.L., Almeida, D., Altenschmidt, J., Altman, S., Anadkat, S., et al.: Gpt-4 technical report. arXiv preprint arXiv:2303.08774 (2023)
2. Akrou, R., Schoenauer, M., Sebag, M., Souplet, J.C.: Programming by feedback. In: International Conference on Machine Learning. pp. 1503–1511. No. 32, JMLR. org (2014)
3. Alur, R., Bodík, R., Dallal, E., Fisman, D., Garg, P., Juniwal, G., Kress-Gazit, H., Madhusudan, P., Martin, M.M.K., Raghothaman, M., Saha, S., Seshia, S.A., Singh, R., Solar-Lezama, A., Torlak, E., Udupa, A.: Syntax-guided synthesis. In: Dependable Software Systems Engineering, pp. 1–25 (2015)
4. Alur, R., Černý, P., Radhakrishna, A.: Synthesis through unification. In: Kroening, D., Păsăreanu, C.S. (eds.) Computer Aided Verification. pp. 163–179. Springer International Publishing, Cham (2015)
5. Andreas, J., Rohrbach, M., Darrell, T., Klein, D.: Neural module networks. In: Proceedings of the IEEE conference on computer vision and pattern recognition. pp. 39–48 (2016)
6. Anil, R., Dai, A.M., Firat, O., et al.: PaLM 2 Technical Report (2023)
7. Bestick, A., Pandya, R., Bajcsy, R., Dragan, A.D.: Learning human ergonomic preferences for handovers. In: 2018 IEEE international conference on robotics and automation (ICRA). pp. 3257–3264. IEEE (2018)
8. Bjørner, N., Phan, A.D., Fleckenstein, L.: ν z-an optimizing smt solver. In: Tools and Algorithms for the Construction and Analysis of Systems: 21st International Conference, TACAS 2015, Held as Part of the European Joint Conferences on Theory and Practice of Software, ETAPS 2015, London, UK, April 11–18, 2015, Proceedings 21. pp. 194–199. Springer (2015)
9. Bornholt, J., Torlak, E., Grossman, D., Ceze, L.: Optimizing synthesis with metas-ketches. In: ACM SIGPLAN Notices. vol. 51, pp. 775–788. ACM (2016)
10. Busa-Fekete, R., Szörényi, B., Weng, P., Cheng, W., Hüllermeier, E.: Preference-based evolutionary direct policy search. In: ICRA Workshop on autonomous learning. vol. 2 (2013)
11. Chasins, S.E., Mueller, M., Bodik, R.: Rousillon: Scraping distributed hierarchical web data. In: Proceedings of the 31st Annual ACM Symposium on User Interface Software and Technology. p. 963–975. UIST ’18, Association for Computing Machinery, New York, NY, USA (2018). <https://doi.org/10.1145/3242587.3242661>, <https://doi.org/10.1145/3242587.3242661>
12. Chen, Q., Lamoreaux, A., Wang, X., Durrett, G., Bastani, O., Dillig, I.: Web question answering with neurosymbolic program synthesis. In: Proceedings of the

- 42nd ACM SIGPLAN International Conference on Programming Language Design and Implementation. p. 328–343. PLDI 2021, Association for Computing Machinery, New York, NY, USA (2021). <https://doi.org/10.1145/3453483.3454047>, <https://doi.org/10.1145/3453483.3454047>
13. Chen, Z., Mao, J., Wu, J., Wong, K.Y.K., Tenenbaum, J.B., Gan, C.: Grounding physical concepts of objects and events through dynamic visual reasoning. arXiv preprint arXiv:2103.16564 (2021)
 14. Christiano, P.F., Leike, J., Brown, T., Martic, M., Legg, S., Amodei, D.: Deep reinforcement learning from human preferences. *Advances in neural information processing systems* **30** (2017)
 15. Cui, Y., Niekum, S.: Active reward learning from critiques. In: 2018 IEEE international conference on robotics and automation (ICRA). pp. 6907–6914. IEEE (2018)
 16. Ding, Y., Zhang, X., Paxton, C., Zhang, S.: Task and motion planning with large language models for object rearrangement. arXiv preprint arXiv:2303.06247 (2023)
 17. Feng, Y., Martins, R., Bastani, O., Dillig, I.: Program Synthesis Using Conflict-Driven Learning. In: Proceedings of the 39th ACM SIGPLAN Conference on Programming Language Design and Implementation. p. 420–435. PLDI 2018, Association for Computing Machinery, New York, NY, USA (2018). <https://doi.org/10.1145/3192366.3192382>, <https://doi.org/10.1145/3192366.3192382>
 18. Feser, J.K., Chaudhuri, S., Dillig, I.: Synthesizing data structure transformations from input-output examples. *ACM SIGPLAN Notices* **50**(6), 229–239 (2015)
 19. Fürnkranz, J., Hüllermeier, E., Cheng, W., Park, S.H.: Preference-based reinforcement learning: a formal framework and a policy iteration algorithm. *Machine learning* **89**, 123–156 (2012)
 20. Gulwani, S.: Automating string processing in spreadsheets using input-output examples. In: Proc. of POPL. pp. 317–330 (2011)
 21. Gulwani, S., Polozov, O., Singh, R.: Program synthesis. vol. 4, pp. 1–119 (2017)
 22. Gupta, T., Kembhavi, A.: Visual programming: Compositional visual reasoning without training. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 14953–14962 (2023)
 23. Han, C., Mao, J., Gan, C., Tenenbaum, J., Wu, J.: Visual concept-metaconcept learning. *Advances in Neural Information Processing Systems* **32** (2019)
 24. Holtz, J., Guha, A., Biswas, J.: Robot action selection learning via layered dimension informed program synthesis. In: Conference on Robot Learning. pp. 1471–1480. PMLR (2021)
 25. Hsu, J., Mao, J., Tenenbaum, J.B., Wu, J.: What’s left? concept grounding with logic-enhanced foundation models. arXiv preprint arXiv:2310.16035 (2023)
 26. Hsu, J., Mao, J., Wu, J.: Ns3d: Neuro-symbolic grounding of 3d objects and relations. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 2614–2623 (2023)
 27. Hu, Z., Lucchetti, F., Schlesinger, C., Saxena, Y., Freeman, A., Modak, S., Guha, A., Biswas, J.: Deploying and evaluating llms to program service mobile robots. arXiv preprint arXiv:2311.11183 (2023)
 28. Jha, S., Gulwani, S., Seshia, S.A., Tiwari, A.: Oracle-guided component-based program synthesis. In: 2010 ACM/IEEE 32nd International Conference on Software Engineering. vol. 1, pp. 215–224. IEEE (2010)
 29. Johnson, J., Hariharan, B., Van Der Maaten, L., Fei-Fei, L., Lawrence Zitnick, C., Girshick, R.: Clevr: A diagnostic dataset for compositional language and elementary visual reasoning. In: Proceedings of the IEEE conference on computer vision and pattern recognition. pp. 2901–2910 (2017)

30. Kalyan, A., Mohta, A., Polozov, O., Batra, D., Jain, P., Gulwani, S.: Neural-guided deductive search for real-time program synthesis from examples (2018). <https://doi.org/10.48550/ARXIV.1804.01186>, <https://arxiv.org/abs/1804.01186>
31. Kane, B., Gervits, F., Scheutz, M., Marge, M.: A system for robot concept learning through situated dialogue. In: Proceedings of the 23rd Annual Meeting of the Special Interest Group on Discourse and Dialogue. pp. 659–662 (2022)
32. Kirillov, A., Mintun, E., Ravi, N., Mao, H., Rolland, C., Gustafson, L., Xiao, T., Whitehead, S., Berg, A.C., Lo, W.Y., Dollár, P., Girshick, R.: Segment anything. arXiv:2304.02643 (2023)
33. Koh, P.W., Nguyen, T., Tang, Y.S., Mussmann, S., Pierson, E., Kim, B., Liang, P.: Concept bottleneck models. In: International conference on machine learning. pp. 5338–5348. PMLR (2020)
34. Li, R., Allal, L.B., Zi, Y., et al.: StarCoder: may the source be with you! (2023)
35. Li, T.J.J., Radensky, M., Jia, J., Singarajah, K., Mitchell, T.M., Myers, B.A.: Interactive task and concept learning from natural language instructions and gui demonstrations. arXiv preprint arXiv:1909.00031 (2019)
36. Liang, J., Huang, W., Xia, F., Xu, P., Hausman, K., Ichter, B., Florence, P., Zeng, A.: Code as policies: Language model programs for embodied control. In: 2023 IEEE International Conference on Robotics and Automation (ICRA). pp. 9493–9500. IEEE (2023)
37. Liu, B., Jiang, Y., Zhang, X., Liu, Q., Zhang, S., Biswas, J., Stone, P.: Llm+ p: Empowering large language models with optimal planning proficiency. arXiv preprint arXiv:2304.11477 (2023)
38. Liu, H., Li, C., Wu, Q., Lee, Y.J.: Visual instruction tuning. arXiv preprint arXiv:2304.08485 (2023)
39. Liu, S., Zeng, Z., Ren, T., Li, F., Zhang, H., Yang, J., Li, C., Yang, J., Su, H., Zhu, J., et al.: Grounding dino: Marrying dino with grounded pre-training for open-set object detection. arXiv preprint arXiv:2303.05499 (2023)
40. Mao, J., Gan, C., Kohli, P., Tenenbaum, J.B., Wu, J.: The neuro-symbolic concept learner: Interpreting scenes, words, and sentences from natural supervision. arXiv preprint arXiv:1904.12584 (2019)
41. Mei, L., Mao, J., Wang, Z., Gan, C., Tenenbaum, J.B.: Falcon: fast visual concept learning by integrating images, linguistic descriptions, and conceptual relations. arXiv preprint arXiv:2203.16639 (2022)
42. Meng, X., Hatch, N., Lambert, A., Li, A., Wagener, N., Schmitt, M., Lee, J., Yuan, W., Chen, Z., Deng, S., et al.: Terrainnet: Visual modeling of complex terrain for high-speed, off-road navigation. arXiv preprint arXiv:2303.15771 (2023)
43. Murali, A., Sehgal, A., Krogmeier, P., Madhusudan, P.: Composing neural learning and symbolic reasoning with an application to visual discrimination. arXiv preprint arXiv:1907.05878 (2019)
44. Namasivayam, K., Singh, H., Bindal, V., Tuli, A., Agrawal, V., Jain, R., Singla, P., Paul, R.: Learning neuro-symbolic programs for language guided robot manipulation. In: 2023 IEEE International Conference on Robotics and Automation (ICRA). pp. 7973–7980. IEEE (2023)
45. Patton, N., Rahmani, K., Missula, M., Biswas, J., Dillig, I.: Programming-by-demonstration for long-horizon robot tasks. Proc. ACM Program. Lang. **8**(POPL) (jan 2024). <https://doi.org/10.1145/3632860>, <https://doi.org/10.1145/3632860>
46. Ren, T., Liu, S., Zeng, A., Lin, J., Li, K., Cao, H., Chen, J., Huang, X., Chen, Y., Yan, F., Zeng, Z., Zhang, H., Li, F., Yang, J., Li, H., Jiang, Q., Zhang, L.: Grounded sam: Assembling open-world models for diverse visual tasks (2024)

47. Rozière, B., Gehring, J., Gloeckle, F., et al.: Code Llama: Open Foundation Models for Code (2023)
48. Sadigh, D., Dragan, A.D., Sastry, S., Seshia, S.A.: Active preference-based learning of reward functions (2017)
49. Silver, T., Athalye, A., Tenenbaum, J.B., Lozano-Perez, T., Kaelbling, L.P.: Learning neuro-symbolic skills for bilevel planning. *arXiv preprint arXiv:2206.10680* (2022)
50. Singh, I., Blukis, V., Mousavian, A., Goyal, A., Xu, D., Tremblay, J., Fox, D., Thomason, J., Garg, A.: Progprompt: Generating situated robot task plans using large language models. In: 2023 IEEE International Conference on Robotics and Automation (ICRA). pp. 11523–11530. IEEE (2023)
51. Solar-Lezama, A.: Program synthesis by sketching. Citeseer (2008)
52. Srivastava, S., Labutov, I., Mitchell, T.: Joint concept learning and semantic parsing from natural language explanations. In: Proceedings of the 2017 conference on empirical methods in natural language processing. pp. 1527–1536 (2017)
53. Suris, D., Menon, S., Vondrick, C.: Vipergpt: Visual inference via python execution for reasoning. *arXiv preprint arXiv:2303.08128* (2023)
54. Wake, N., Kanehira, A., Sasabuchi, K., Takamatsu, J., Ikeuchi, K.: Gpt-4v (ision) for robotics: Multimodal task planning from human demonstration. *arXiv preprint arXiv:2311.12015* (2023)
55. Wang, G., Xie, Y., Jiang, Y., Mandlekar, A., Xiao, C., Zhu, Y., Fan, L., Anandkumar, A.: Voyager: An open-ended embodied agent with large language models. *arXiv preprint arXiv:2305.16291* (2023)
56. Wang, R., Mao, J., Hsu, J., Zhao, H., Wu, J., Gao, Y.: Programmatically grounded, compositionally generalizable robotic manipulation. *arXiv preprint arXiv:2304.13826* (2023)
57. Wang, Y., Shah, R., Criswell, A., Pan, R., Dillig, I.: Data migration using datalog program synthesis. *VLDB* (2020)
58. Wang, Y., Wang, X., Dillig, I.: Relational program synthesis. *PACMPL 2(OOPSLA)*, 155:1–155:27 (2018)
59. Wilson, A., Fern, A., Tadepalli, P.: A bayesian approach for policy learning from trajectory preference queries. *Advances in neural information processing systems* **25** (2012)
60. Wirth, C., Fürnkranz, J., Neumann, G.: Model-free preference-based reinforcement learning. In: Proceedings of the AAAI Conference on Artificial Intelligence. vol. 30 (2016)
61. Xie, A., Finn, C.: Lifelong robotic reinforcement learning by retaining experiences. In: Conference on Lifelong Learning Agents. pp. 838–855. PMLR (2022)
62. Xie, E., Wang, W., Yu, Z., Anandkumar, A., Alvarez, J.M., Luo, P.: Segformer: Simple and efficient design for semantic segmentation with transformers. *Advances in Neural Information Processing Systems* **34**, 12077–12090 (2021)
63. Yi, K., Wu, J., Gan, C., Torralba, A., Kohli, P., Tenenbaum, J.: Neural-symbolic vqa: Disentangling reasoning from vision and language understanding. *Advances in neural information processing systems* **31** (2018)
64. Zhang, R., Bansal, D., Hao, Y., Hiranaka, A., Gao, J., Wang, C., Martín-Martín, R., Fei-Fei, L., Wu, J.: A dual representation framework for robot learning with human guidance. In: Conference on Robot Learning. pp. 738–750. PMLR (2023)

Appendix

A Implementation Details

Inputs. For learning the program, we collected 29 demonstrations where each demonstration had a trajectory of robot poses, the associated RGB camera image, and pointcloud data from Ouster-128 LiDAR sensor. A high-resolution LiDAR was needed to achieve as accurate a mapping as possible from 2D to 3D. A few of the natural language prompts from the user are listed below:

1. It is safe since it is on a sidewalk, and is far from any person and the pole, and it is not on grass.
2. This looks good since it is on a sidewalk, and is far from the approaching person and the pole.
3. This is not safe since it is not far from bushes, even though its on a sidewalk.

Models. The details are as follows:

- a. Grounded-SAM: Using a set of hyperparameters for all objects does not always yield good results due to the limitations of current VLMs, *i.e.*, their accuracy and sensitivity differs a lot between different classes of objects. We use Grounded-SAM [46] with the object-specific hyperparameters as shown in Table 4 to perform zero-shot object detection and segmentation on images.
- b. Terrain segmentation: Our experiments showed that present VLMs do not do so well on terrain segmentation, which was a domain-specific essential capability to be able to represent the preferential concept well enough. Thus, we finetuned the SegFormer-b5 [62] model, with pretrained weights from the HuggingFace transformers library, on our custom UT campus dataset. The details about the models can be found on the project website.
- c. We use GPT-4 [1] as our language model for the sketch synthesis module. We set a temperature of 0.0 and seed it for enhancing reproducibility. Prompts for doing different tasks in our framework (*i.e.*, grounding, synthesis *etc.*) can be found in the codebase. Note that they have certain placeholders like `<!...!>` and `<dyn!...!dyn>` which refer to other prompt instances or are replaced dynamically in the code based on generated outputs.

Concept Library Initialization. We impart some basic neural and domain-specific capabilities to the learning framework by seeding it with the functions shown in Fig. 6.

Table 4: Object-specific hyperparameters for VLM

	box-threshold	text-threshold	nms-threshold
barricade	0.5	0.5	0.3
board	0.3	0.3	0.5
bush	0.4	0.4	0.4
car	0.3	0.3	0.3
entrance	0.3	0.3	0.2
person	0.25	0.25	0.6
pole	0.4	0.4	0.5
staircase	0.25	0.25	0.4
tree	0.4	0.4	0.45
wall	0.5	0.5	0.4

```

def project_pixel_to_map(pixel_loc: Tuple[int, int]) -> Tuple[float, float]:
    """Given a pixel location tuple, returns the corresponding location in the global map frame."""

def distance_to_nearest_object(loc: Tuple[float, float], object_class: str) -> float:
    """Given a location tuple and an object class, returns the real-world distance of the loc to the nearest
    object of that class."""

def extend(flat_object: str, alpha: float) -> None:
    """Given a flat object, extends all instances of the object in the normal direction in global map frame
    by alpha amount."""

def terrain_at(pixel_loc: Tuple[int, int]) -> str:
    """Given a pixel location tuple, returns the terrain class at that location."""

def is_traversable(terrain: str) -> bool:
    """Given a terrain class, returns whether the terrain is traversable."""

def slope_at(loc: Tuple[float, float]) -> float:
    """Given a location tuple, returns the slope at that location."""

```

Fig. 6: Concept library seeded with some basic predicates

B Detailed Results & Explanations

B.1 Baselines

We curate a custom dataset on UT campus region, consisting of 815 pixel-level labeled images (375 each for CONTINGENCY and DROPOFF, and 65 for PARKING). We evaluate the following baselines:

- * SegFormer b0 and b5 models, both with or without depth input. For taking in depth, we only modify the input layer, and still retain all other pretrained weights. We take measures such as early stopping to prevent overfitting. The hyperparameters used to finetune each of these models can be found on their model cards in HuggingFace (links to which can be found on the project website).
- * GPT4-vision: Due to token limitations, we query GPT4-vision to output a 20×20 output class array given the image. For reporting the IOU for GPT4-vision, we downsample the ground-truth to the same size for the comparison

Table 5: Results for CONTINGENCY concept

	iou_pos (%)			iou_neg (%)			miou (%)		
	train	in-test	out-test	train	in-test	out-test	train	in-test	out-test
Synapse	57.22	54.46	50.18	98.06	98.11	97.96	77.64	76.29	74.07
SF-RGB-b0	43.54	28.46	17.77	97.44	97.03	97.07	70.49	62.75	57.42
SF-RGB-b5	51.63	43.87	19.04	97.54	97.08	92.96	74.59	70.48	56.00
SF-RGBD-b0	46.71	37.07	13.49	97.62	97.39	96.00	72.17	67.23	54.75
SF-RGBD-b5	55.10	38.48	15.71	97.85	97.13	96.50	76.48	67.81	56.11
GPT4V	01.73	01.91	02.74	51.38	57.51	58.15	26.56	29.71	30.45
GPT4V+	02.29	02.18	02.59	55.16	55.74	65.24	28.73	28.96	33.92
VisProg	04.99	01.25	05.04	86.24	90.01	84.92	45.62	45.63	44.98
VisProg+	08.13	06.36	07.15	69.75	72.06	76.51	38.94	39.21	41.83

Table 6: Results for DROPOFF concept

	iou_pos (%)			iou_neg (%)			miou (%)		
	train	in-test	out-test	train	in-test	out-test	train	in-test	out-test
Synapse	60.64	62.05	63.13	97.99	98.31	98.31	79.32	80.18	80.72
SF-RGB-b0	48.59	38.93	08.12	97.39	97.32	96.30	72.99	68.13	52.21
SF-RGB-b5	56.73	48.00	15.94	97.78	97.66	94.13	77.26	72.83	55.04
SF-RGBD-b0	51.09	42.17	13.89	97.57	97.43	95.90	74.33	69.80	54.90
SF-RGBD-b5	57.43	43.86	08.93	97.95	97.54	95.84	77.69	70.70	52.39
GPT4V	02.40	02.25	03.09	58.67	63.69	72.20	30.54	32.97	37.65
GPT4V+	04.43	01.90	03.02	74.33	74.77	75.25	39.38	38.34	39.14
VisProg	01.48	01.94	06.58	91.10	93.03	89.56	46.29	47.49	48.07
VisProg+	08.62	06.86	08.70	69.71	72.01	77.57	39.17	39.44	43.14

to be fair. The ‘+’ variant essentially means that we prompt it with additional information about the user’s ground-truth, *i.e.*, we provide it the program that SYNAPSE has learned, in natural language. All prompts can be found in the codebase.

- * VisProg and VisProg+ come from a related approach to VQA [22]. We again follow the same methodology of prompting as for GPT4-vision. Here too, no finetuning is done and the evaluation is zero-shot. All prompts can be found in the codebase.

Tables 5 to 7 show the full evaluation for the three concepts, with the IOU for each class. Due to unavoidable data imbalance on the pixel-level (*i.e.*, $\approx 90\%$ negative pixels against 10% positive class pixels), it is much easier to learn the negative class.

B.2 Ablations

We test the following ablations for our framework:

Table 7: Results for PARKING concept

	iou_pos (%)		iou_neg (%)		miou (%)	
	train	out-test	train	out-test	train	out-test
Synapse	39.14	27.24	98.06	98.27	68.60	62.76
SF-RGB-b0	21.28	04.91	94.07	94.40	57.68	49.66
SF-RGB-b5	38.99	08.13	98.27	97.68	68.63	52.91
SF-RGBD-b0	24.80	04.67	97.00	96.70	60.90	50.69
SF-RGBD-b5	43.97	03.50	98.14	96.48	71.06	49.99
GPT4V	01.10	01.11	75.72	79.25	38.41	40.18
GPT4V+	01.56	00.64	81.19	78.90	41.38	39.77
VisProg	04.72	01.69	74.13	80.29	39.43	40.99
VisProg+	03.31	03.08	74.44	74.89	38.88	38.99

* Framework ablations: We test three alternative ways to generate the program sketch from given natural language input from the user:

- a. **Synapse-SynthDirect**: given only the basic predicates, we adopt a one-step approach to synthesis, where we ask LLM to update the program sketch based on the new NL input and previous program sketch, while utilising only the basic predicates, *i.e.*, it has no way to hierarchically build and retain higher-level predicates, as well as it does not do any CNF extraction.
- b. **Synapse-SynthDirect+**: we provide the higher-level concepts learned by our main framework and then given these already learned higher-level predicates, we again adopt a one-step approach to synthesis, where we ask LLM to update the program sketch based on the new NL input and previous program sketch, *i.e.*, it still does not do any CNF extraction. Note, however, this is *only* possible for a post-learning ablation study, as in an actual learning framework, we do not have apriori access to the higher-level learned predicates.
- c. **Synapse-SynthCaP**: we disallow the framework from querying the user for auxiliary demonstrations to learn auxiliary concepts, *i.e.*, the framework is forced to generate code (using the concept library) for the auxiliary concepts solely based on the information available, which essentially is the *name* of that particular concept. This is similar to the recursion performed in Code-as-Policies [36].

* NN-ablations: We finetune the four SegFormer models on the same number of samples as SYNAPSE, which is 29 for the CONTINGENCY concept.

* LLM-ablations: We test the performance of our overall framework using other language models available: (1) CodeLLama [47], (2) StarCoder [34], and (3) PaLM2 [6]. We see that it is essential for the LLM to be strong enough for the performance of SYNAPSE to be good.

The results reaffirm the design choices made in SYNAPSE.



Fig. 7: Illustration of reasoning about failures using the learned program

B.3 User-study

The user case study involved each of the three participants using SYNAPSE to have it learn their notion about the preferential concept under consideration. We chose CONTINGENCY for the study. Each of the three participants gave 10-15 demonstrations after which each of them labeled 15 images with their *ground-truth*.

It was observed that the learned programs differed on two levels: (1) symbolic concepts, and (2) quantitative parameters. While many of the symbolic concepts were similar across users, there were a few such as `is_too_far` or `is_next_to` that differed. We then evaluated each learned program on the ground-truth labels of all the users, which showed good one-one correspondence between learned programs and ground-truth of users, indicating good alignment.

C Error Analysis using SYNAPSE

An added benefit of the neuro-symbolic representation of the preferential concept with SYNAPSE is that it facilitates reasoning about failures. Fig. 7 shows one such example where we can easily reason about the false positives and false negatives comparing the prediction with the individual component outputs of the program.