



D. Calanzone, DISI, University of Trento
F. Merlo, CIMEC, University of Trento



- VLMs present **sparse concepts and weak compositional relations** [5, 6].
- **Comparison** allows humans to attend to relevant features in inputs, to highlight differences and thus to learn abstractions.
- **RQ1:** *can we teach logical concepts through comparison?*
- **RQ2:** *can we find an efficient multi-task architecture for this task?*

- We base on **Simulated Objects for Language Aquisition (SOLA)** [1].
- As in [1], we teach concepts by grouping objects that express them (**similarity batch**) and objects that do not (**dissimilarity batch**).

RED

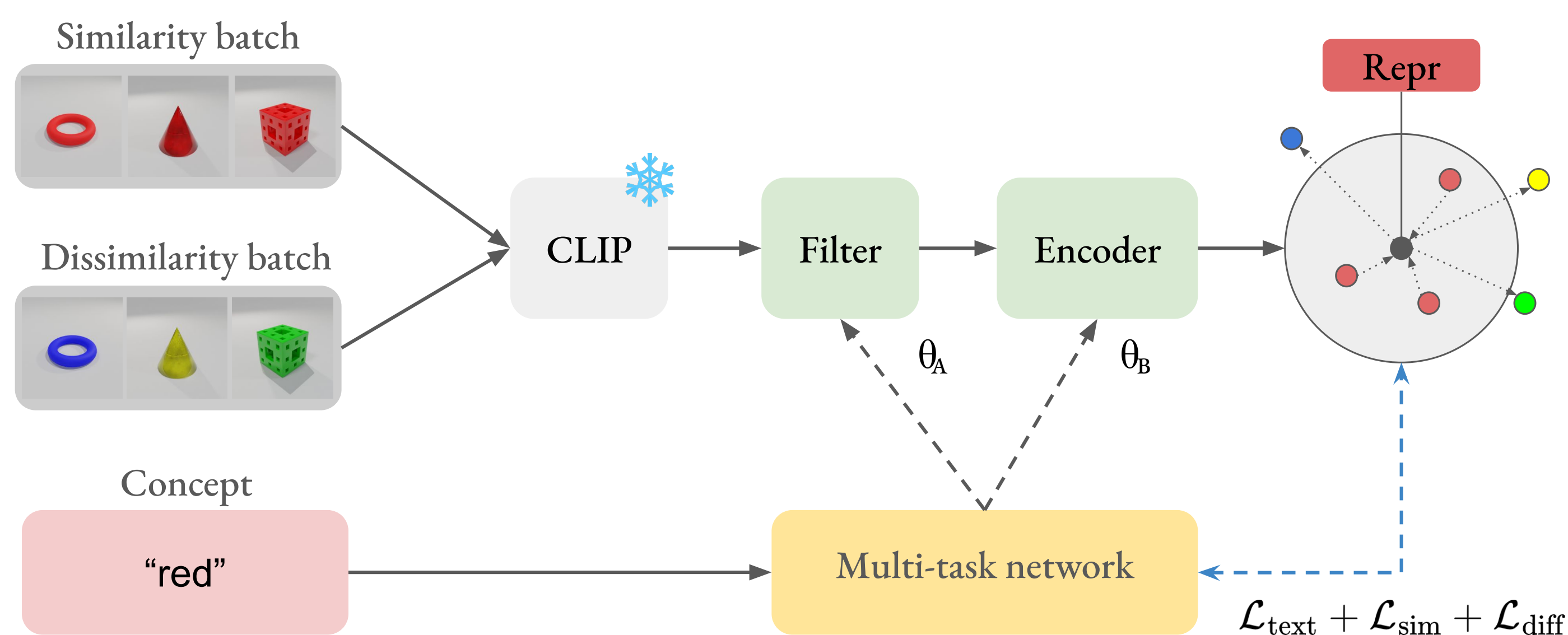


The image shows two groups of 3D objects. The left group, labeled 'RED', contains a red ring, a red cone, and a red cube. The right group contains a blue ring, a yellow cone, and a green cube.

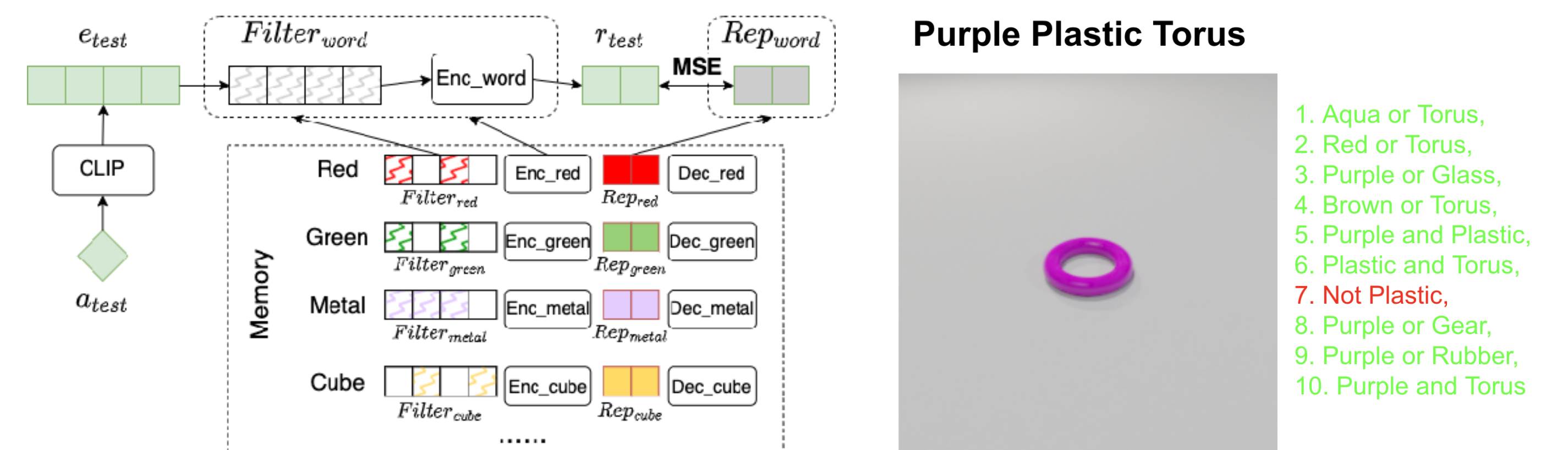
-

Figure 1 shows a 2x3 grid of images. The top row contains three images of a blue sphere, a green cube, and a red cube, respectively. The bottom row contains three images of a blue sphere, a green cone, and a red ring, respectively. Each image shows the object from a different perspective, illustrating the concept of 'different views'.

- We adopt for **RQ1** concept-specific networks as in [1].
- Fine-grained Contrastive Learning with additional attention maps.
- For **RQ2**, we introduce multi-task adapters for CLIP modulated by the input text: a hyper-network [4] and modular skill sharing [3], in comparison.



- For **primitive concepts**, we test our networks on Multi-Attribute Recognition (MAR) [1].
- For **complex logical expressions**, we modify MAR and thus define Logical Pattern Recognition (LPR).



-
- The graph displays the performance of a logical pattern recognition system across eight different conditions. The x-axis represents the number of logical relations retrieved (0 to 65), and the y-axis represents the accuracy score (0.3 to 1.0). The conditions are categorized into four groups: 'Tot' (Total), 'Not' (Not), 'And' (And), and 'Or' (Or), each with 'New' and 'Variation' sub-conditions.
- | Number of Logical Relations Retrieved | Tot New Objects | Not New Objects | And New Objects | Or New Objects | Tot Variation | Not Variation | And Variation | Or Variation |
|---------------------------------------|-----------------|-----------------|-----------------|----------------|---------------|---------------|---------------|--------------|
| 0 | 0.95 | 0.98 | 0.75 | 0.95 | 0.95 | 0.98 | 0.75 | 0.95 |
| 10 | 0.88 | 0.94 | 0.58 | 0.90 | 0.88 | 0.94 | 0.58 | 0.90 |
| 20 | 0.85 | 0.94 | 0.53 | 0.89 | 0.85 | 0.93 | 0.55 | 0.88 |
| 30 | 0.82 | 0.93 | 0.49 | 0.87 | 0.82 | 0.92 | 0.50 | 0.86 |
| 40 | 0.79 | 0.92 | 0.44 | 0.85 | 0.79 | 0.91 | 0.42 | 0.83 |
| 50 | 0.76 | 0.91 | 0.38 | 0.82 | 0.76 | 0.90 | 0.39 | 0.80 |
| 60 | 0.72 | 0.90 | 0.32 | 0.78 | 0.72 | 0.89 | 0.33 | 0.75 |
| 65 | 0.68 | 0.89 | 0.28 | 0.75 | 0.68 | 0.88 | 0.30 | 0.72 |

- **RQ2:** Modular skill sharing (**MSS**) [3] is comparable in performance to **hyper-networks**, but much more parameter efficient ($\sim 3x$).
- Learning in sequence requires experience replay (DER++ [2]).
- Materials have subtle visual features and thus are harder to learn.

Split	Model	Color	Material	Shape
D_{test_nc}	Concept-specific	0.96	0.48	0.98
	HyperNet	0.37	0.25	0.73
	HyperNet (DER++)	0.71	0.28	0.89
	Modular Skills Sharing	0.72	0.21	0.67

- **RQ1:** concept-specific networks can learn worlds satisfying a logical rules with CL. *Next: testing noisy environments.*
- **RQ2:** multi-task adapters well adapt CLIP for the task [4]. Reducing the skills could nudge learning more generic patterns. *Next: extending to logical pattern recognition.*

- [1] Human Inspired Progressive Alignment and Comparative Learning for Grounded Word Acquisition. Bao et al. 2023.
- [2] Dark Experience for General Continual Learning: a Strong, Simple Baseline. Buzzega et al. 2020.
- [3] Combining Modular Skills in Multitask Learning. Ponti et al. 2022.
- [4] Parameter-efficient Multi-task Fine-tuning for Transformers via Shared Hypernetworks. Mahabadi et al. 2021.
- [5] Do Vision-Language Pretrained Models Learn Composable Primitive Concepts? Yun et al. 2022.
- [6] When and why vision-language models behave like bags-of-words, and what to do about it? Yuksekgonul et al. 2022.

Diego Calanzone
University of Trento
DISI
diego.calanzone@studenti.unitn.it
<https://halixness.github.io>