

# UC Merced

## Proceedings of the Annual Meeting of the Cognitive Science Society

### Title

Learning Abstract Categories through Linguistic Feature Explanations

### Permalink

<https://escholarship.org/uc/item/7v28m9df>

### Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

### Authors

Xie, Haodong

Wang, Xuena

Maharjan, Rahul Singh

et al.

### Publication Date

2026

### Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

# Learning Abstract Categories through Linguistic Feature Explanations

Haodong Xie<sup>1</sup>, Xuena Wang<sup>2</sup>, Rahul Singh Maharjan<sup>1</sup>, Federico Tavella<sup>1</sup> & Angelo Cangelosi<sup>1</sup>

<sup>1</sup>Centre for Robotics and AI, The University of Manchester

<sup>2</sup>The School of Psychology and Cognitive Science, East China Normal University

## Abstract

Abstract categories often lack a single, bounded perceptual referent, making them harder to learn than highly concrete categories. Explanations that identify properties shared across category members may support category learning by highlighting features that generalise beyond individual examples. This study investigates the role of explanatory concepts in category learning using a label-free Concept Bottleneck Model framework. We generate human-interpretable feature concepts for both basic-level and super-ordinate categories, then learn a projection from visual representations into this concept space. A classifier then predicts category labels from these concept activation to evaluate whether explanatory features support category learning. We evaluate the approach on CIFAR-100, CUB-200, and ImageNet, organising labels into basic-level and super-ordinate categories. Across datasets, concept-based models achieve accuracy comparable to standard classifiers while providing interpretable concepts. Visualisations reveal distinct category clusters in concept space, and evaluations on held-out subordinate classes suggest that explanatory concepts transfer effectively to previously unseen subclasses.

**Keywords:** Concept Learning; Abstract Concept; Concept Bottleneck Model