

Learning Generalizations and Exceptions: the Good, the Bad and the Unpredictable

Karina Tachihara

Princeton University, Princeton, New Jersey, United States

Adele Goldberg

Princeton University, Princeton, New Jersey, United States

Kenneth Norman

Princeton University, Princeton, New Jersey, United States

Abstract

How are exceptions to a generalization learned? 20 participants were exposed to a mini-artificial language in which each word (prefix + wordstem) was associated with a unique image. One of two prefixes generalized probabilistically: it appeared with 40 stems associated with faces and 8 exceptions, which were associated with scenes. The other prefix occurred with 8 faces and 8 scenes. The prefixes and the image categories (faces vs. scenes) were counterbalanced across participants. Participants performed a 2 alternative-forced choice task on all items, with feedback, over 6 repeating blocks. Results show that image-word pairs that included the generalizable prefix were learned better than those which appeared with the other prefix, despite having 48 items in the first class and 16 in the other ($d = 1.28$, $d = 0.80$, $p = 0.023$). We investigate the neural representation of these words and how they change over the course of learning.