

Comparing Visual and Omniscient Models of Collective Crowd Motion

James Falandays; William Warren

+ Author Affiliations & Notes

Journal of Vision August 2023, Vol.23, 5124. doi:<https://doi.org/10.1167/jov.23.9.5124>

Abstract

Collective motion can emerge in human crowds when pedestrians match the heading and speed of nearby neighbors. Existing computational models of this phenomenon are “omniscient,” making the assumption that pedestrians have direct knowledge of the 3D positions and velocities of their neighbors, and are thus not cognitively plausible. Here we compare two models of collective motion in human crowds, based on either (1) distal physical or (2) proximal optical variables. In our original omniscient model (Rio, Dachner & Warren, PRSB 2018), each pedestrian aligns their heading (or speed) with a weighted average of the heading (or speed) of their neighbors, whose weights decay exponentially with distance. In our new visual model (Dachner, Wirth, Richmond & Warren, PRSB 2022), a pedestrian’s heading (or speed) is a function of the mean optical expansion and mean angular velocity of neighbors, which trade off with eccentricity, and the influence of far neighbors is reduced by visual occlusion. We compare the two models in multi-agent simulations to test the conditions under which they converged to collective motion. $N=50$, 100, or 200 agents continuously interacted on a torus, with synchronous updating and 100 runs per condition. Their initial positions on a grid were randomly jittered, and initial conditions were parametrically varied: interpersonal distance (IPD=1, 2, 4, 8, 10m), heading range (45°, 90°, 135°, 180°, 270°, 360°), base speed (1.0, 2.0, 3.0 m/s), and speed range (base ± 0.2 , ± 0.4 , ± 0.6 , ± 0.8 m/s). Performance measures included the mean normalized velocity (order parameter, range 0-1) and the relaxation time. Both models converge to collective motion, but the visual model is more robust over a wider range of initial conditions. The visual model thus outperforms the omniscient model, in addition to being more cognitively plausible.

This work is licensed under a [Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License](https://creativecommons.org/licenses/by-nc-nd/4.0/).

