

More Neurons Did Not Help Everywhere

Part 3 in The Making of Elliot. A sweep across sensory populations showed that added capacity did not help uniformly. Some channels mattered far more than others, and combining strong channels could make the system worse again.

2026.03.23 21:49 DENNIS HEDEGREEN BUILD OPEN-NOTE

[HTTPS://HEDEGREENRESEARCH.COM/ARTICLES/MORE-NEURONS-DID-NOT-HELP-EVERYWHERE/](https://hedegreenresearch.com/articles/more-neurons-did-not-help-everywhere/)

After the first Elliot benchmarks, one pattern kept returning.

Improvement did not arrive cleanly.

More structure helped a little.

Randomness helped more.

But every gain came with a cost.

That raised the next question.

If each sense gets more neurons, does Elliot improve because the whole sensory layer becomes richer, or because a few specific channels carry most of the useful signal?

That was the next test.

The Setup

I kept the base version fixed at v3, the version with five neurons per sense and loop-breaking noise already active. Same world. Same benchmark structure. Same agent. The only thing that changed was sensory population placement.

First I ran a uniform sweep. Every sense was set to the same population size: 1, 2, 3, 4, and 5. That was the cleanest way to test the simple claim that more sensory capacity everywhere should help.

It did not.

The uniform sweep showed no clear monotonic improvement. Expanding all senses together did not produce a smooth rise in food-seeking. That made the problem more interesting, not less. If more neurons were helping, they were not helping evenly.

So the next pass changed one sense at a time while the others stayed fixed at baseline.

That is where the signal became clearer.

What Happened

Single-sense expansion did not produce equal gains across the sensory layer. Some channels mattered much more than others.

In the long confirmation batch, baseline performance sat at a mean food rate of 0.08044.

When only the west food channel was expanded to two neurons, mean food rate rose to 0.09212.

At four west-food neurons, it reached 0.09123.

At five, it was 0.08977.

That is already enough to reject the simplest version of the idea. Elliot did not need more neurons everywhere. He improved most when a specific input channel was given more capacity.

The east food channel also turned out to matter more than most of the others. On its own, east-food expansion could also push performance well above baseline.

But the more surprising result came after that.

When I combined two of the strongest-looking food channels instead of expanding only one, performance did not rise again. It fell back down.

The test case of Food_E=2 + Food_W=4 reached only 0.08214, very close to baseline and far below the stronger single-channel setups.

At the same time, cost rose sharply. Pain events increased. Wall events stayed elevated. More capacity in two useful channels did not combine into a better system. It produced interference.

What That Means

The result is not that more neurons are useless.

The result is that more neurons do not help in a simple global way.

That matters, because it cuts against an easy intuition. If a system improves when some capacity is added, it is tempting to assume that adding more capacity more broadly should help more. Elliot does not support that conclusion.

Instead the gains appear uneven and channel-dependent.

That leaves two live possibilities.

One is that certain input channels really do carry more useful information for this agent in this world.

The other is that the world itself is structurally biased, and that Elliot is exploiting asymmetry in the environment rather than discovering a more general sensory principle.

That distinction matters.

If the world is static in a way that privileges certain directions, then this is not yet a general neural finding. It is a finding about how a particular organism interacts with a particular world.

That is still useful. It just has to be named honestly.

Why It Matters

Elliot was never meant to be a performance demo.

He is useful because he can fail in public.

This is one of those failures, or at least one of those corrections. The system did not reward a broad increase in sensory richness. It rewarded narrow increases in specific channels, and then punished over-combination.

That is a better result than a generic win would have been.

It means the next question is sharper.

Not “how do I give Elliot more neurons?”

But:

What exactly is he learning from the world he is in?

Which improvements are real?

Which ones are artifacts of the grid?

What Comes Next

The next step is not to add complexity for its own sake.

The next step is to pressure the world itself.

If west-facing food signals keep winning, I need to know whether that is a real property of the agent or a byproduct of a static environment, reset geometry, and directional bias in the current grid.

That means the next Elliot work should move carefully.

First: close the current sensor line cleanly.

Then: vary the world.

Only after that does it make sense to push Elliot toward a deeper field-based model with gradients, resistance, and more primitive energy logic.

For now, one thing is clear.

More neurons helped.

They just did not help everywhere.

— Dennis Hedegreen

RELATION MEMORY

CONTINUES [When Randomness Helps](#) · Continues the Elliot parameter-sweep line by showing that more capacity did not help everywhere.

CORRECTS [The Making of Elliot](#) · Adds a correction layer to the early organism experiment: more sensory capacity is not automatically better.

CANONICAL URL: [HTTPS://HEDEGREENRESEARCH.COM/ARTICLES/MORE-NEURONS-DID-NOT-HELP-EVERYWHERE/](https://hedegreenresearch.com/articles/more-neurons-did-not-help-everywhere/)

PDF VIEW GENERATED: 2026-07-11T15:35:45.905Z