

A LECTURE COMPANION

**"Nothing in biology makes sense
without teleology" by Michael Levin**

Michael Levin

Recorded on July 5, 2025

About this document

This document is a companion to the recorded lecture "*Nothing in biology makes sense without teleology*" by Michael Levin, recorded on July 5, 2025. You can watch the original lecture or listen in your favorite podcast feeds — all links are on the page [here](#).

This document pairs each slide with the aligned spoken transcript from the lecture. At the top of each slide, there is a "Watch at" timestamp. Clicking it will take you directly to that point in the lecture on YouTube.

Lecture description

This is a deliberately provocative talk (~1 hour 20 minutes) I gave on teleology (a pretty taboo subject in a lot of the life sciences) delivered at Caltech. I try to go step by step and show the philosophical background of how I think about goal-directedness in physically embodied agents (cells and tissues etc.) and then the data - classic examples and new discoveries that this framework allowed us to make.

Editing by <https://x.com/DNAMediaEditing>

Follow my work

[Twitter](#) • [Blog](#) • [The Levin Lab](#)

Transcript note

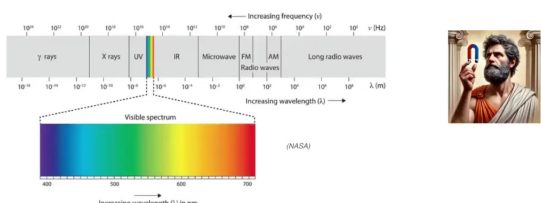
This transcript is generated automatically. While we strive for accuracy, occasional errors may occur. Please refer to the video for exact wording.

Want one for your lecture?

Want something like this for your own talk? Reach out to Adi at adi@aipodcast.ing.

The Power of Unification: from conceptual frameworks to technology

*"Nothing in Biology Makes Sense Except in the Light of Evolution" - 1973 essay by evolutionary biologist Theodosius Dobzhansky
Looking for unifying principles in biology*



- understand how seemingly-diverse phenomena are a continuum
- create technology to observe and detect parts of that continuum of which we were previously oblivious
- utilize this knowledge for a myriad of applications that improve quality of life

If you want to download any of the primary papers, the data, the software, everything is free and open on our website, which is here. This is my personal blog about what I think some of these things actually mean.

I modeled the title of this talk based on a very, very famous quote by Dobzhansky called "Nothing in biology makes sense in light of evolution." What he was referring to is a unifying principle, something that can be used to make sense of the biological world that we see. This idea of unification is really powerful. I often think about the example of the electromagnetic spectrum. Prior to having a good theory of electromagnetism, we thought that static electricity, lightning, magnets, light were all completely different things. They certainly seem like different things. We had different words for them. The vocabulary that we used helped us to see them as different. The unifying theory allowed us to see them all as one spectrum and to figure out that they have something very fundamental in common. It also allowed us to make new technology that enabled us to operate in parts of that spectrum that our evolutionary history does not allow. We're only sensitive naturally to a narrow space, but now we have technology to see all of it.

There's a large part of my efforts. Our lab is about 43 people, and a lot of the practical output of our lab is interventions for birth defects, regenerative medicine, cancer, bioengineering. We do some AI as well, and these are practical things for modern regenerative medicine. What underlies all this is my commitment to try to understand how minds exist in the physical world. What does it mean to be a mind that operates in the physical universe? What is embodiment? We try to bring some of those ideas into practical, both computational and wet lab applications that, to whatever extent they're useful and help us discover new things, tell us whether we're on the right track or not.

What I'm interested in is understanding very diverse kinds of intelligence on the same scale. We're talking about not just birds and primates, and maybe a whale or an

octopus, but also very weird colonial creatures and swarms and synthetic new life forms and AIs, whether software or robotic. Perhaps someday exobiological agents, such as alien life when we find it, and even agents that are actually patterns within media. I'm not the first person to try this. Rosenbluth, Weiner, and Bigelow were trying to make a kind of spectrum of categories going from passive matter through various cybernetic stages up through human-level, second-order metacognition.

I'm interested in understanding what all minds have in common, no matter their scale, origin story, or composition, for mainly two reasons. First, to help move forward practical applications in biomedicine, synthetic morphology, and so on. Second, an improvement in our ethical systems, because I think the ones that we have are based on ancient philosophical and linguistic categories that pick out no real distinctions in the world, and are going to get us into a lot of trouble as we start to share our world with increasingly diverse beings—cyborgs and hybrot and all kinds of things. Some of them are here now, some of them are coming, but all of them will need a very different kind of ethical system that isn't so obsessed with human-style brains.

Before I get going and show you a bunch of biology, here's an outline of my points for today. First of all, let's get some metaphysics out of the way. When I talk about teleology, I do not mean anything related to quibbles of linguistics or philosophy. I don't think you can sit back and make decisions about whether teleology does or doesn't exist. I think it's an empirical question. The way that we treat these things empirically is this. All of these frameworks, no matter how weird, I don't care how dimensional, whether it's vitalism or anything else, I'm game for anything, as long as it helps us generate new discoveries and new capabilities. I don't mean look backwards at something that already happened and give some kind of account of what it was, because you can always do that at the lowest level. You can always tell an atomic story about anything that happened. What I mean is, does the framework help you do new experiments and discover new things that you otherwise wouldn't have discovered? That's my criteria for today.

Second, when I say teleology, I don't mean human-level purpose, meaning that it's a second-order kind of thing where I know what my goal is and I know that I'm pursuing a goal. I don't necessarily mean that, but it's on the spectrum. What I mean by teleology is the kind of thing that cybernetics does for us. It gives us a non-mysterious science of systems that are consistent with both the laws of physics and the ideas in cognitive science about systems that have goals. Cybernetics tells us about how systems can have goals, and computer science tells us how to study the ingenuity of these systems to meet those goals under different circumstances. We have these sciences now and we no longer have to freak out about goals being some sort of magical thing that somehow science isn't going to be able to get a handle on.

Part of my argument: there are two main things I want to do today. First, I'm going to show you the data for a very simple argument. Goal-directedness and teleology, here I'm using interchangeably. The first uncontroversial point is that whatever goals are, they're the sort of thing that humans have. We all know we can represent goals in our

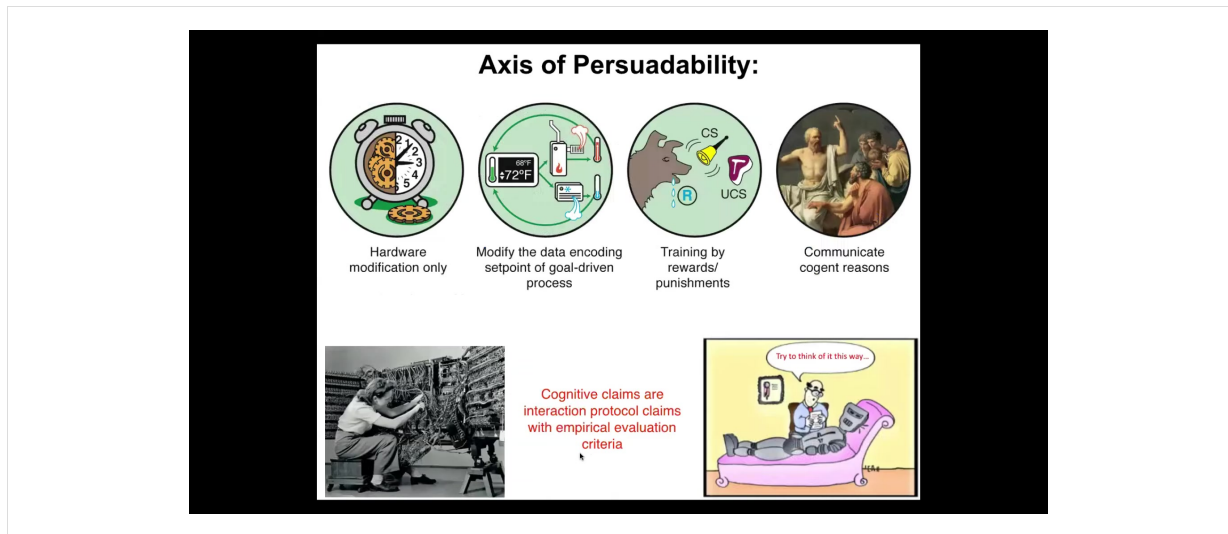
minds. We can undergo actions to pursue them. We have some degree of ingenuity to meet those goals. Sometimes we succeed, sometimes we fail. But if goals mean anything at all, we as humans are an example of it. That's point #1.

Point #2, we actually know something about the mechanisms that enable goal-directed behavior in brainy organisms like us. We know something about the electrophysiological networks in our brain that enable goal-directedness. But remarkably, those mechanisms are ancient. They are highly conserved. They are not unique to brains and nervous systems. There is a very deep and fascinating evolutionary homology, not just analogy, not just metaphorical thinking, but actual homology between cognition and morphogenesis. This is seen most clearly in developmental bioelectricity, which is what I'm about to show you.

What I'm going to go through is a bunch of examples of biological goal-directed problem solving, mostly in the bioelectric realm, but I'll also show you some other examples that are different. That's my argument, that we have a non-controversial system, AKA us, brainy animals that clearly have teleological behavior. What I'm going to point out is that the same mechanisms that underlie that are very widely distributed in the biological world and, I think, beyond that. That's a topic for another day. It's a very simple argument.

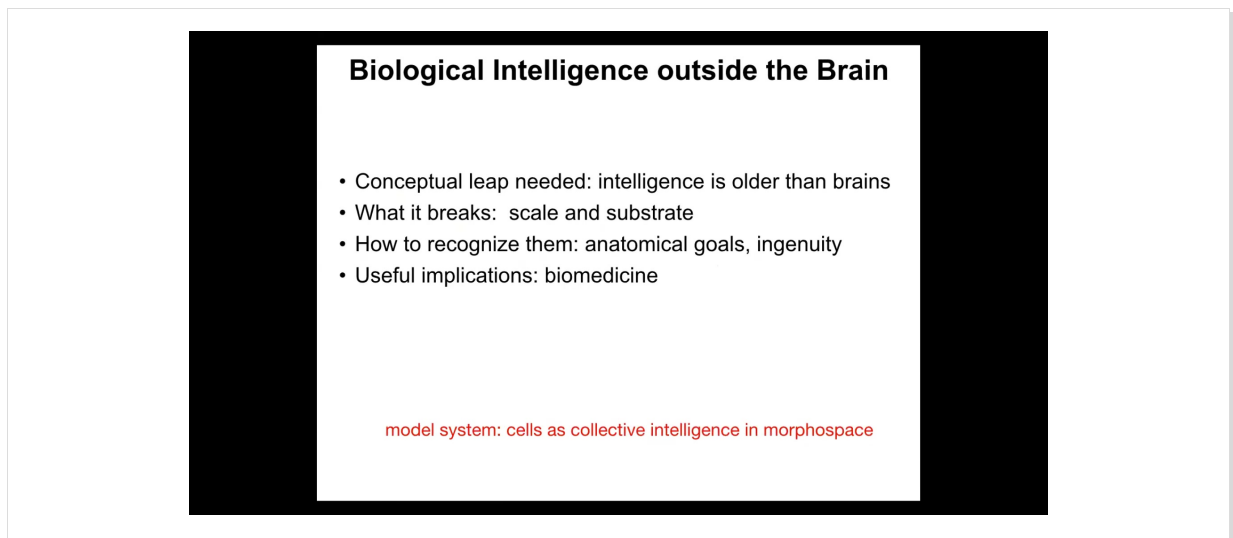
The other thing I want to say is that teleology, I think, is just the tip of the iceberg. If you're worked up about teleology, it's going to get much worse. What I think is really the unifying principle of biology isn't just teleology. It's cognition broadly considered.

Goal directedness is just the first step of a huge spectrum of cognitive skills, and that is the spectrum that I think underlies everything that we're interested in. Teleology is just the first step. This is part of a view that some people call cognition all the way down, and it's a reasonably extreme position in the emerging field of diverse intelligence, but that's where I am.



So what I like to claim is this. I think that cognitive terms, meaning when you say something has goals, that it is intelligent, that it has memories, or they can be trained, I suggest that those things are interaction protocol claims. I think that all of these are observer-relative interaction protocol claims. So what you're saying when you say that a certain system has or doesn't have some capacity is you're basically picking out a set of tools across a spectrum of very different kinds of tools: hardware modification in the case of mechanical clocks, resetting set points via cybernetics and control theory, as Rob just pointed out, behavioral science, psychoanalysis, and psychiatry. We have bags of tools, and what you're saying is, I'm going to try applying a particular set of tools, and then we can all see how well I did, because it's an empirical thing as to how that turns out.

So my claim, when I talk about this in more conventional audiences, is I'm going to take the tools of behavioral and cognitive neuroscience, both the conceptual tools and the bench tools, and I'm going to apply them far outside the brain. I'm going to show you that the tools actually don't distinguish, and if the tools don't distinguish... What is it that we are trying to draw a boundary between? The fact that this kind of thinking allows us to discover new capabilities by reusing those tools elsewhere, and the conventional model did not do that because it suggests that those tools should only be used in the brain, that should tell us something.



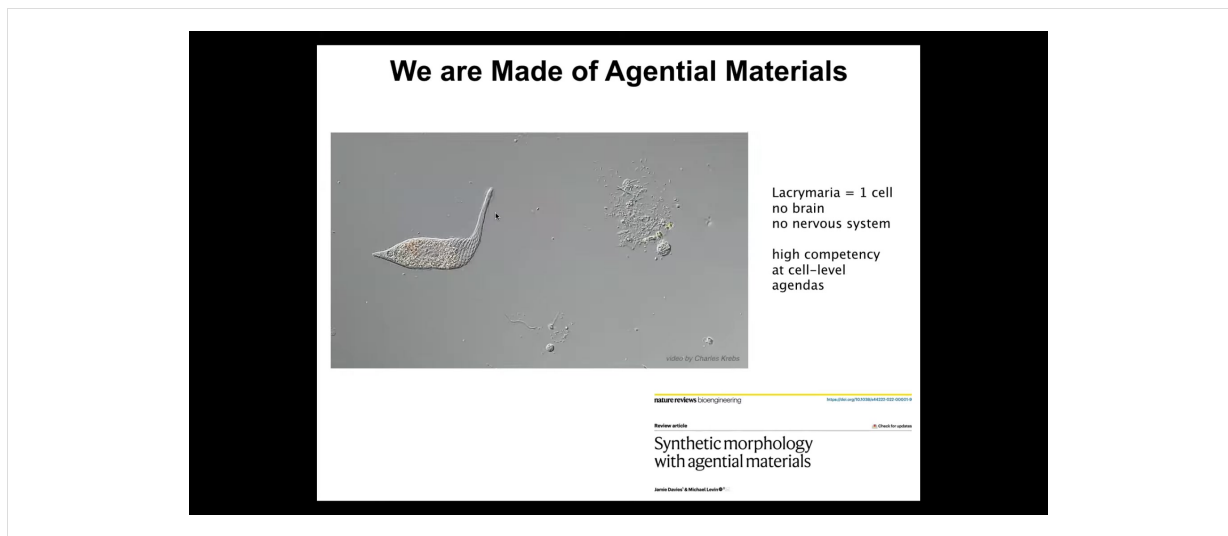
Biological Intelligence outside the Brain

- Conceptual leap needed: intelligence is older than brains
- What it breaks: scale and substrate
- How to recognize them: anatomical goals, ingenuity
- Useful implications: biomedicine

model system: cells as collective intelligence in morphospace

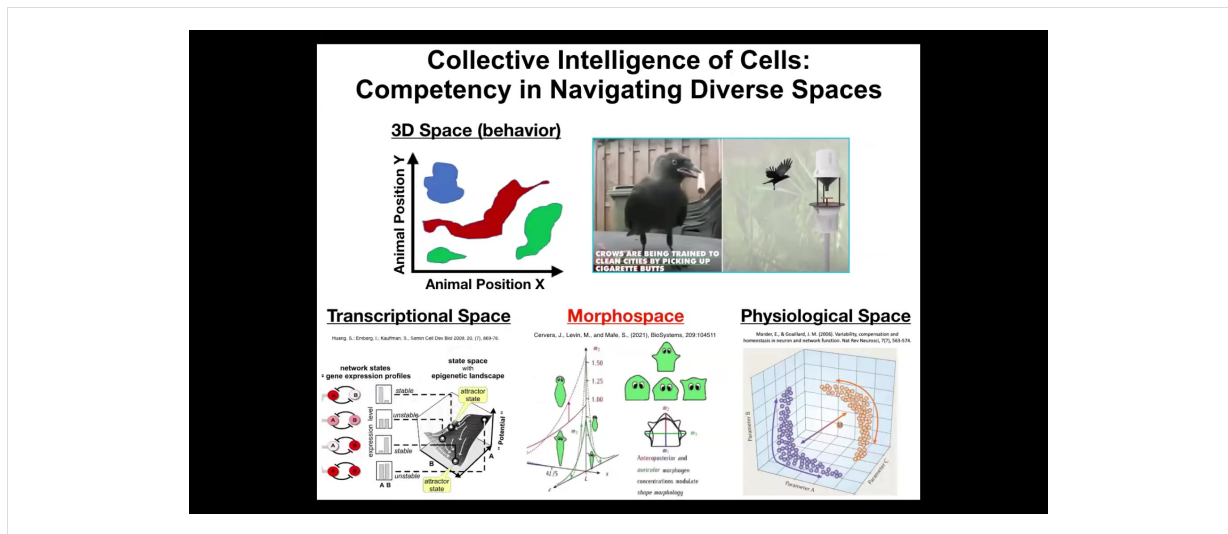
The model system that we use for a lot of this—years ago we started working on this model system, which is groups of cells as a collective intelligence. I’m going to define those terms momentarily.

And what this collective intelligence is doing is exerting behavior by navigating anatomical morphospace, which I will also define. In order to do that, this idea breaks a number of assumptions but has implications. Let’s look at that.



The first thing to point out is that we and our goal-directed capabilities are a multi-scale phenomenon. We're made of an agential material. We are not made of passive matter. The things that we are made of on multiple scales of organization themselves have agendas, learning properties, and other capabilities. Jamie Davies and I talk here about how you engineer with things like this. Engineering with agential materials is very different than engineering with passive matter or even active matter.

This is a multi-scale competency architecture, where through the scales, from the molecules up to the collectives and groups, these things are solving problems in diverse spaces. With different degrees of sophistication, what are these spaces?



We as humans, because of our own evolutionary history, we're obsessed with three-dimensional space. We're primed to notice intelligent behavior, goal-seeking behaviors of medium-sized objects moving at medium speeds, so birds and mammals.

But biology has been solving problems in other spaces for a really long time. There's the space of gene expression, high-dimensional transcriptional space, the space of physiological states, and anatomical morphospace, which is basically the morphospace of all possible configurations or shapes of a body. This is the one we're going to talk the most about today.

But my claim is that, just as animals navigate three-dimensional space to meet their goals, we have all kinds of structures in our body that are navigating with similar kinds of skills, navigating these other spaces that are very hard for us to notice, and until we have the right theory and the right instrumentation, which is what I think we're doing, just like with the idea of the electromagnetic spectrum. We're developing tools to port everything that works here, behavioral and cognitive science, to these other substrates, and that helps us visualize these things.

Teleology is a Subset of Intelligence Spectrum

“Intelligence is the ability to reach the same goal by different means.”

- William James

Hypothesis: morphogenesis is a collective intelligence, exerting its behavioral competencies in anatomical morphospace

what kind of collective intelligence do cellular swarms deploy?

- Different time scale (minutes, not milliseconds)
- Different problem space (morphospace, not 3D motion)
- but, same (homologous!) mechanisms - ion channels, GJs, NTs
- and, many of the same algorithms (tools of cog sci work great)

So I claim that teleology is a subset of the spectrum of intelligence. What I mean by intelligence is this. William James gave a nice, very cybernetic definition of intelligence. He didn't talk about brains. He didn't talk about what problem space was involved. He said it's the ability to reach the same goal by different means. This ability to reach the same goal, that's basically teleology, that there is a goal and that the system is going to exert effort to try to get there. But the different means is interesting. It means that every system is going to have some degree of the ability to get there when circumstances change, when you put barriers in its place and so on. That's a really important part of the toolkit here. If you want to study teleology, you cannot do it from pure observation. You have to put barriers between the system and whatever you think its goal is and observe the degree to which it can be clever by getting around whatever you just did to it.

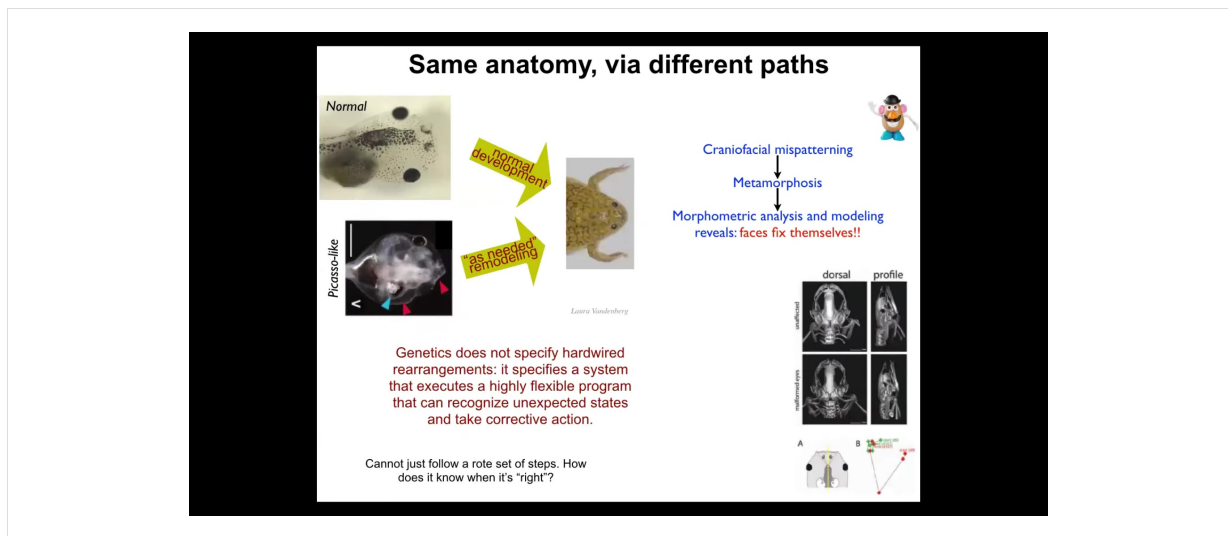
So our hypothesis is this. For this portion, I'm going to show you morphogenesis is a collective intelligence. We are all collective intelligences because we are made of parts. What it does is it exerts its behavioral competencies in anatomical space. In getting from being a single cell to whatever we end up being, it navigates that space. I'm going to show you some ways that it does that. In moving here from neuroscience, you have to go to a different time scale, so minutes instead of milliseconds, different problem space, morphospace instead of three-dimensional motion. But everything else is the same. The mechanisms are homologous. Many of the same algorithms, perceptual bistability, active inference, work well. I used to have my students do this by hand, and now we made a tool for people to do this. If you take any neuroscience paper and you do a find-and-replace—every time it says neuron, you say cell, and every time it says millisecond, you say minutes or hours—you've got yourself a developmental biology paper. It works amazingly well because there is this deep underlying symmetry between the two.

I'm going to show you six examples from our world of the components of intelligence. Let's start with goal-directedness in a simple way, anatomical homeostasis.

Here's a creature. This is an axolotl. It's an amphibian. This little guy regenerates its limbs, its eyes, its jaws, its ovaries, its spinal cords, portions of the brain and heart. The way that works is here's the leg. If you amputate, or if you let them bite off eac

And when they do, the leg they build ends up being a very nice frog leg, but it does not get there by the same path that the embryonic leg grows. It actually looks completely different, very plant-like. Development, we know, is very reliable, so the egg always gives rise to the kind of thing it's supposed to build. That's not why I'm calling it intelligent.

I can take this early embryo, cut it into pieces. I don't get half bodies. I get perfectly normal twins and triplets and quadruplets and so on. Because in anatomical space, the same system can get to where it's going, this ensemble of goal states representing a normal human target morphology. It can get there from here, but also from here, and from all kinds of different starting positions. Even avoiding local minima, it can get to where it needs to go. This process of homeostasis—half of it knows that the other half is missing and will rebuild it as much as it needs—is fundamental. I actually think development is basically just an extreme version of regeneration. Instead of regrowing your limb, you're regrowing the entire body from one cell. I think the whole thing is fundamentally just homeostasis.



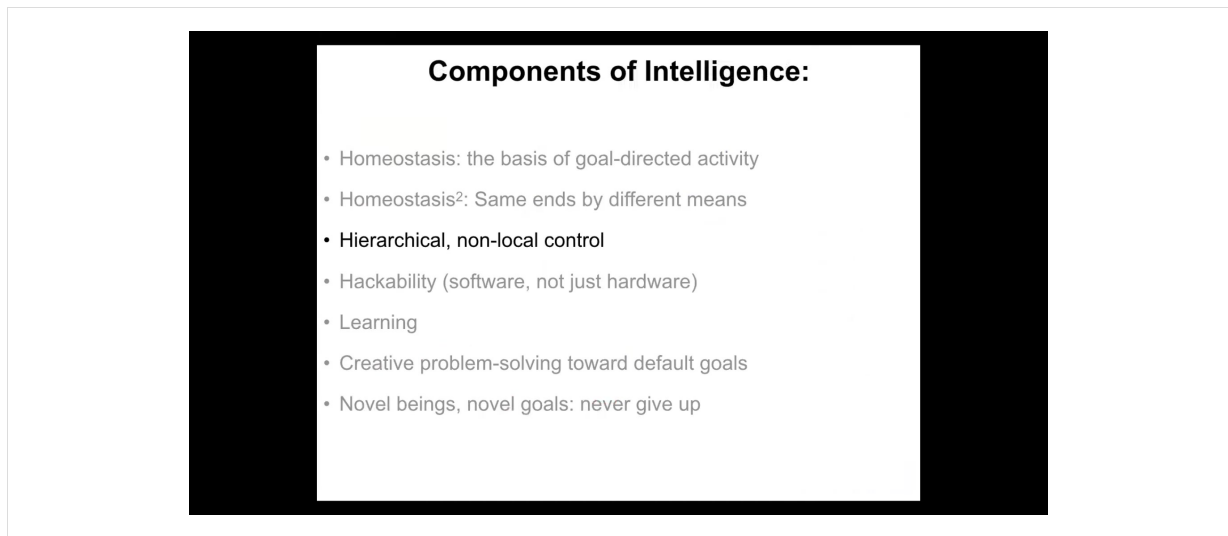
Here's another example. This is something that we discovered a few years ago.

Here's a tadpole. Here are the eyes, the nostrils, the brain, the mouth, the gut. This tadpole has to turn into a frog. All of these craniofacial organs have to move. They have to rearrange. The eyes have to move forward. The mouth has to pull out.

It used to be thought that this was hardwired, that the genetics encoded a set of precise movements for every structure so that you start out as a normal tadpole and end a normal frog. We decided to test this. You cannot just assume what the level of intelligence is.

What we did was we created so-called Picasso tadpoles. In Picasso tadpoles, everything is scrambled. The eyes are on the back of the head, the mouth is off to the side. It's like a Mr. Potato Head where you've shuffled around all the organs. They're just completely screwed up. What you find is that these animals make very normal frogs, because all of the different organs, starting in their crazy positions, will move through novel, unnatural paths to get to where they need to be. In fact, sometimes they go too far and have to double back.

The genetics does not give you a set of hardwired rearrangements. What it specifies is a system that can do an error minimization scheme. It knows what the set point is, and I'm going to show you how it knows in a minute. It is able to reduce that error by whatever actions it needs to take. These things are moving in novel paths. That, of course, raises the obvious question, then how does it know what a correct pattern is? I'm going to show you that momentarily.



Components of Intelligence:

- Homeostasis: the basis of goal-directed activity
- Homeostasis²: Same ends by different means
- Hierarchical, non-local control
- Hackability (software, not just hardware)
- Learning
- Creative problem-solving toward default goals
- Novel beings, novel goals: never give up

I've shown you two things. I've shown you basic homeostasis. I've shown you a more flexible kind of homeostasis. Now I'm going to show you non-local control.

The key thing to realize about this is that it is not just about damage. The point is not about injury response and restoring what you have. It's a much deeper type of teleology that's built into all of the aspects of this.

Here's the experiment. This wasn't us. This was these guys in the 50s. You take a tail and you graft it in the middle of the flank. Somewhere between the forelimb and the hindlimb, you surgically graft it to the flank.

What happens over time is that this thing turns into a limb. In place, it actively remodels from being a tail to being a limb.

Notice something interesting. If your tail-tip cells are sitting at the end of a tail, locally, there's nothing wrong. You are exactly where you're supposed to be, at the end of a tail. Yet you turn into fingers.

Why do you turn into fingers? Because the key aspect of these kinds of problem-solving, goal-seeking systems is that in a collective intelligence where the large-scale system is solving a goal in its own space, it has to control its own parts and deform their option landscape so that they do things that match the higher-level goals. No individual cell knows what a finger is or what a tail is or how many you're supposed to have, or where they're supposed to be.

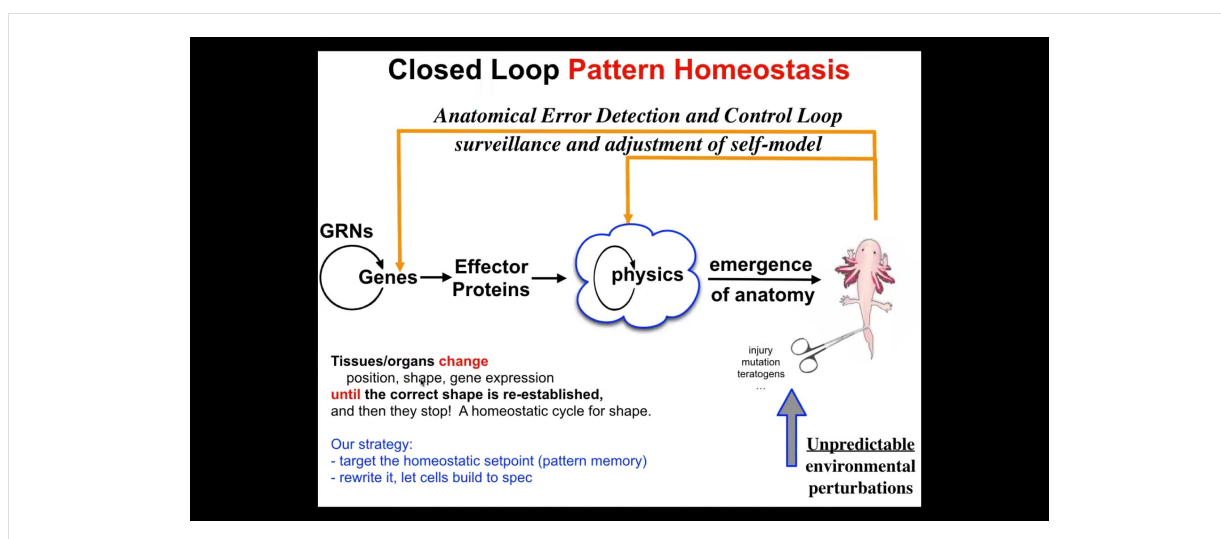
But the collective has a goal. The goal is a reasonably normal-looking axolotl. In order to meet that goal, it has to propagate downward all of the control signals to its various parts.

Notice the parallel of this in cognition. When I'm talking to you, I am not worried about having to adjust all your synaptic proteins so that you're able to remember what I said. You're going to do that yourself.

The beauty of these kinds of multi-scale problem-solving systems is that they take care of, and basically hide, the underlying molecular events, and they take care of whatever needs to happen in order for the higher-level system to meet its goal.

That's why these tail-tip cells start to become fingers. It's not just about repairing damage; it's about harnessing local order to a global plan. The goal belongs to the large-scale collective; it does not belong to any of the cells, but to the whole, and it's a very large-scale goal.

Slide 9 of 42 · Watch at [22:40](#)



So now, notice how all of this differs greatly from the mainstream approach to these problems that you find in your developmental biology textbook. The standard approach is this feed-forward, open-loop thing consisting of complexity and emergence. We have gene regulatory networks. These genes turn each other on and off. Some of them become effector proteins that diffuse or they're enzymes. And then there's all this local physics that happens and then this complex shape emerges. What that word means is simply that when you do a lot of repeated operations on local components, out might come something complex. That's absolutely true, and we have lots of model systems for the cellular automata and fractals. We know that by recurrent simple rules, you can get complexity.

But that isn't the whole story, that isn't how biology does a lot of these things, because there's a problem here. If that's what it is, if you want to make changes up here, let's say you wanted a three-fold symmetry instead of bilateral symmetry, what this model is telling you is that you have to make changes down here. And you don't know what changes to make down here in the large majority of changes you might want to make. This is not reversible. It's an inverse problem that is not solvable. In fact, what I'm adding to this is this idea that under this scheme, there is no goal. There is no teleology because there is no goal. There's only feed forward, turn the cranks, local rules, out will come something complex or it won't, but there is no goal. The system is not going towards anything.

I'm proposing a different model where this happens, but something else happens that's very important: there actually is a represented goal in the system, and I'm about to show you a picture of it. There is data around it: there is a represented goal. What the system will do if you deviate from the set point is turn on, both at the level of physics and genetics, all kinds of mechanisms that will try to reduce the delta and get you back to that goal.

So if that were the case, that would be a huge and important thing for, for example, regenerative medicine, because it would mean that we don't need to solve this terrible inverse problem. We don't need to mess with the hardware necessarily. We just need to re-specify the goal states and let the system do what it does best, which is try to meet the goal.

And so this is what we've been doing now for about 25 years: to find, if this is going to be a goal-seeking system like your thermostat, it has to represent the set point somewhere. It has to remember what the correct pattern is. That means I should be able to find it, I should be able to decode it, and I should be able to rewrite it. Those are the three things. If I can't do those things, then we have not shown that the system has goals.

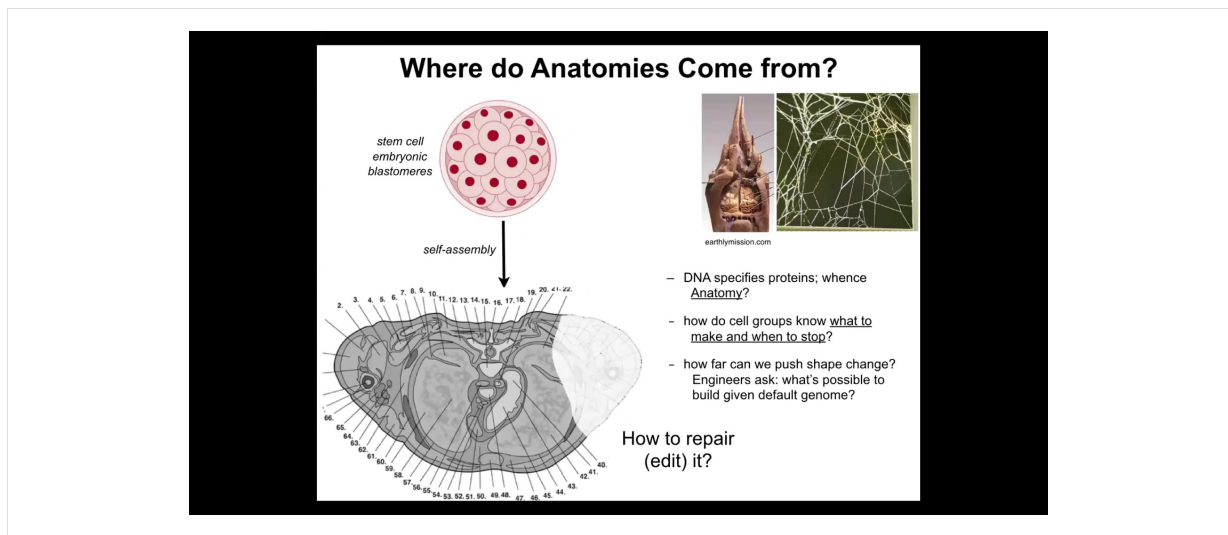
I want to show you some of this data. The idea, again, is that morphogenesis is a collective intelligence and that its behavioral competencies occur in anatomical space. Those are the kind of problems that it solves.

Here's what I'm going to claim that bioelectric networks can do for us. These claims parallel what the deal is in neuroscience. I'm only going to show you a couple of these things. The first thing these networks do is they store memories. If you're going to have a homeostatic system that seeks goals, you have to know what your goal is. You have to somehow store and represent the set point. That's the first thing they do. They do the computations needed to perform error minimization, to actually navigate the space. Then you need to integrate all kinds of long-range information to assess the current state.

What does my body look like? Do I have a tail in the middle or do I have legs or nothing? Or what does it look like to compare against the set point to initiate any kind of remodeling?

The final thing is they have to serve as a kind of cognitive glue. The reason that you and I are not just a pile of neurons and that we know things that our neurons don't know is that electrophysiology serves as a cognitive glue binding individual cells toward larger-scale systems whose goals project into new spaces, spaces that the individual parts don't have any access to. Bioelectricity does this in your body like it does in the brain.

Slide 10 of 42 · Watch at [27:12](#)



So let's talk about something that you might think is a solved problem. Oftentimes, people think that developmental biology has already figured out everything. The question is, where do anatomies come from in the first place? Here's a cross-section through a human torso. You see incredible order. Everything is in the right place, the right size, next to the right thing, amazing.

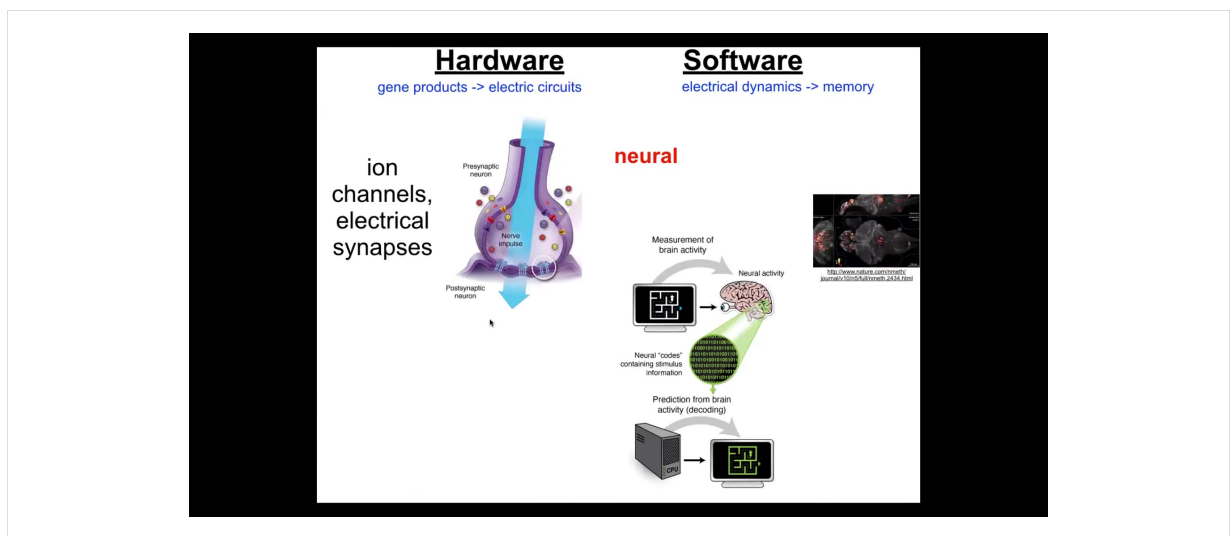
But we all just started life as a ball of embryonic blastomeres. Where was this pattern encoded? And then people often will say, well, it's in the DNA, it's in the genome, but we can read genomes. We know that what the genome encodes is proteins. It doesn't directly specify any of this stuff any more than the genome of termites specifies the structure of the mound or the genome of a spider will tell you exactly what the web is going to look like. So the genome gives you the hardware. The genome tells you what computational hardware all the cells are going to have. But we need to understand the software. How do the cells know what to build? How do they know when to stop?

As engineers, we'd like to know, could we ask them to build something completely different? The same cell, same DNA, could we ask them to build something different?

If we get to it, I'll show you at the end that the answer is yes, we can build some crazy things. And for regenerative medicine, we'd like to know how do we tell them to repair things that are missing.

This kind of goal-directed behavior in the brain is handled by electrical circuits, which consist of cells that have ion channels.

Slide 11 of 42 · Watch at [28:36](#)



These ion channels allow the cell to reach some voltage potential. They have these electrical synapses that allow the voltage to propagate, or not, to neighboring cells. The enormous network of these things executes a kind of electrophysiological dynamics that neuroscientists think are the bearers of cognition.

So, here's a living zebrafish brain that these guys imaged while the fish was doing a cognitive task. And there's this idea of neural decoding such that if we were to record this information from a fish or from a human, if we knew how to decode it, we would be able to read the thoughts, the memories, the goals, the preferences.

One example of goal directedness that pretty much everybody agrees on is humans. This is the mechanism that maintains it. This mechanism allows you to store memories as goals. It allows you to control your muscles to try to act on those goals. It does all the computations needed to reduce the delta between where you are now and wherever you're trying to be. So this is the mechanism of well-accepted teleology.

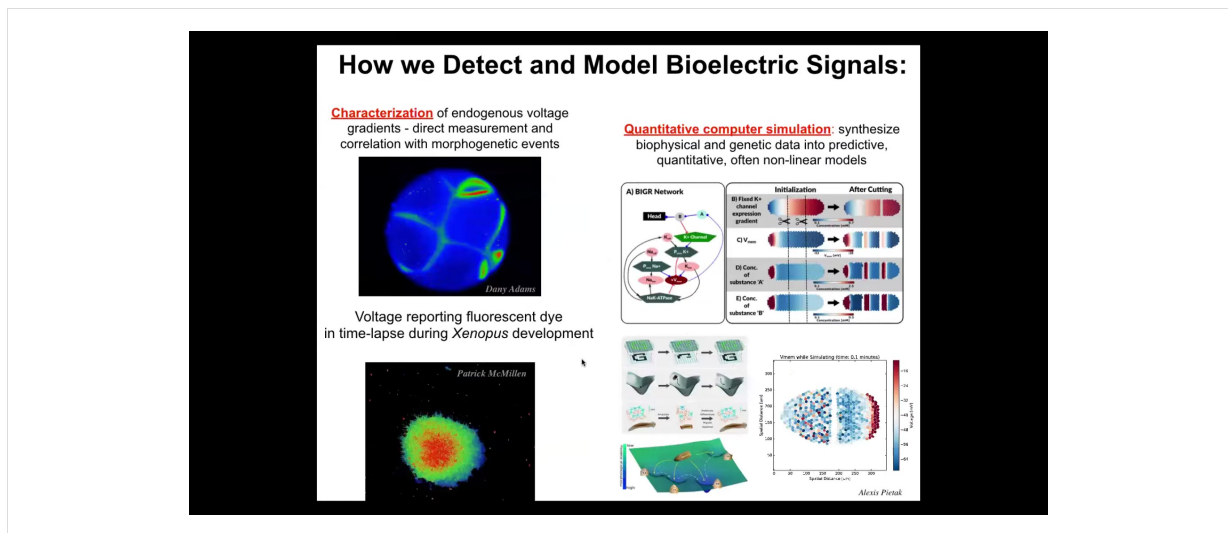
However, it turns out evolution caught on to this way before neurons and brains evolved. Every cell in your body has ion channels. Most cells have these electrical

synapses. This idea of using electricity to coordinate and bind individuals into a larger-scale unit was discovered around the time of bacterial biofilms. Gurl's work shows how bioelectricity was used in bacterial networks to organize microbes into a colony, a primitive multicellular creature.

What we decided to do was to steal everything from computational neuroscience and use the same techniques to try to decode what the cognitive content of the body might be. We know what the brain thinks about. The brain thinks about moving you through three-dimensional space to reach your goals. But what do the bioelectrical networks of your body, which exist both evolutionarily and developmentally long before you have nerves, think about?

It turns out that they think about shape. They think about moving your body through the anatomical morphospace as opposed to 3D space. I think that is where the brain learned all of its cool tricks: by pivoting this much more ancient problem-solving teleological system from morphospace into 3D space and then speeding up things so that you can run around and eat other things.

Slide 12 of 42 · Watch at [31:27](#)

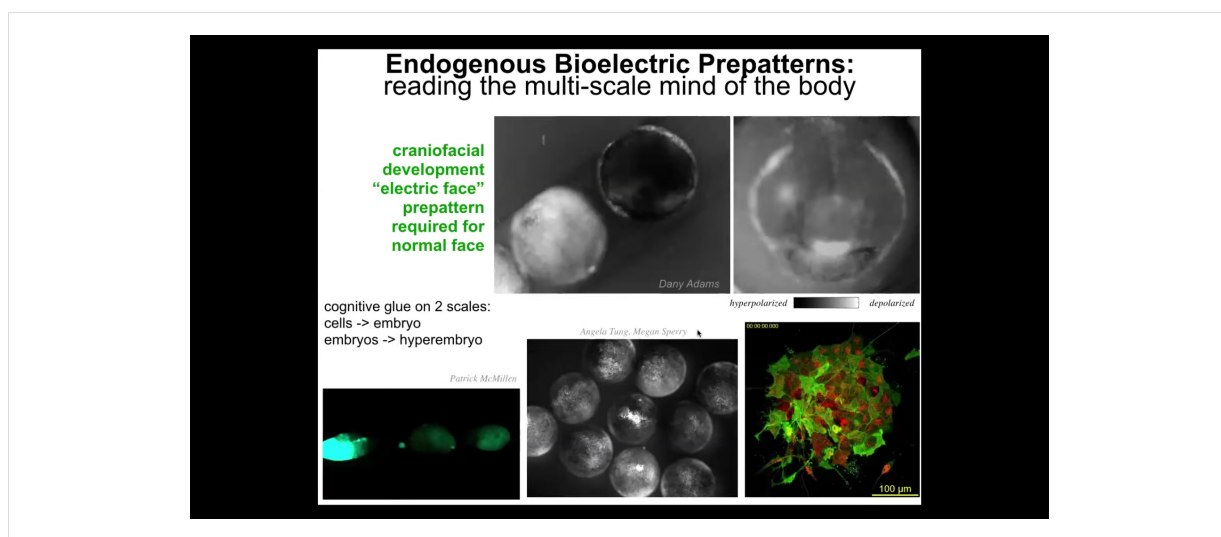


And so we developed a bunch of tools, and I'll show you, the first thing we had to develop was a way to read, just like neuroscientists read in the brain. We had to read the electrophysiology of the body. We developed tools — voltage-sensitive fluorescent dyes and reporter proteins. Here's a frog embryo where you can see all the electrical conversations that a cell is having with each other in real time. This isn't a model, this

is real data. Cells in culture, you can see that the colors represent voltage. You can see the different voltages as the cells are making decisions about whether to be part of this morphogenetic mass or to leave.

We do a lot of multi-scale simulations, starting with the ion channels that produce the gradients into tissue-level dynamics that help you explain all kinds of interesting facts about regeneration and so on. Then we go further and we try to merge this with ideas in connectionist machine learning and dynamical systems theory so that we can try to understand memories as attractors and pattern completion and all of these things, these proto-cognitive things going on in electrical networks.

Slide 13 of 42 · Watch at [32:43](#)



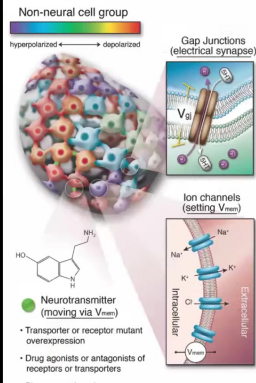
I'm going to show you a couple of real examples of this. This is an early frog embryo putting its face together. The grayscale are the different voltage values. There's a lot going on here, but look at this one frame from that video. We call this the electric face. Long before the genes actually turn on to regionalize that face, here's where the eye is going to go, here's where the mouth is going to go, here are the placodes out here. This is the subtle information structure that tells the face, and the other eye comes in slightly later, it tells the face where all the craniofacial organs are supposed to go. And if you move any of these bioelectrical patterns, the gene expression will follow and the anatomy will follow. This is how you can make a scrambled Picasso tadpole by changing the pattern that normally tells these things where they should start out.

I will point out that not only is bioelectricity a cognitive glue on the scale of a single embryo, meaning that it binds individual cells towards embryo-level outcomes, it also

works on the next level. These are multiple embryos. There are three of them. If you poke this guy right here, there's an injury wave within seconds. This one finds out about it, and then this one—same thing happens here. It turns out that bioelectrical communication among embryos within a group actually helps them resist teratogens, deviations from their set point, much better than singletons. So the collective: it's at least two layers of collective intelligence, cells into a single embryo, a single embryo into a collective that can solve problems that individual cells have trouble doing. We have tons of this kind of data, watching and trying to decode what these patterns represent, just like neuroscientists try to decode from the brain.

Slide 14 of 42 · Watch at [34:35](#)

Rewriting Somatic Pattern Memories



Non-neural cell group

hyperpolarized ← → depolarized

Gap junctions (electrical synapse)

Ion channels (setting V_{mem})

Neurotransmitter (moving via V_{mem})

- Transporter or receptor mutant overexpression
- Drug agonists or antagonists of receptors or transporters
- Photo-uncaging of neurotransmitter

Tools we developed
(no applied fields!)

- Dominant negative Connexin protein
- GJC drug blocker
- Cx mutant with altered gating or permeability

Synaptic plasticity

- Dominant ion channel over-expression (depolarizing or hyperpolarizing, light-gated, drug-gated)
- Drug blocker of native channel
- Drug opener of native channel

Intrinsic plasticity

But more importantly than just watching them, you have to be able to rewrite them. We have to detect the memory, and then we have to rewrite the set point. That's really critical.

We took tools from the neuroscientists. We do not use any electrodes. There are no magnets, no frequencies, no electromagnetic fields, no waves. What we are doing is using the same interface that cells use to hack each other, which is that we control the gap junctions, the electrical synapses, so we can control the topology of the network in cells and tissues. We can also control the actual voltage by targeting the ion channels, whether it be through drugs or optogenetics. We can lay down light masks and put specific patterns of voltage and see. We have some other tricks with neurotransmitters.

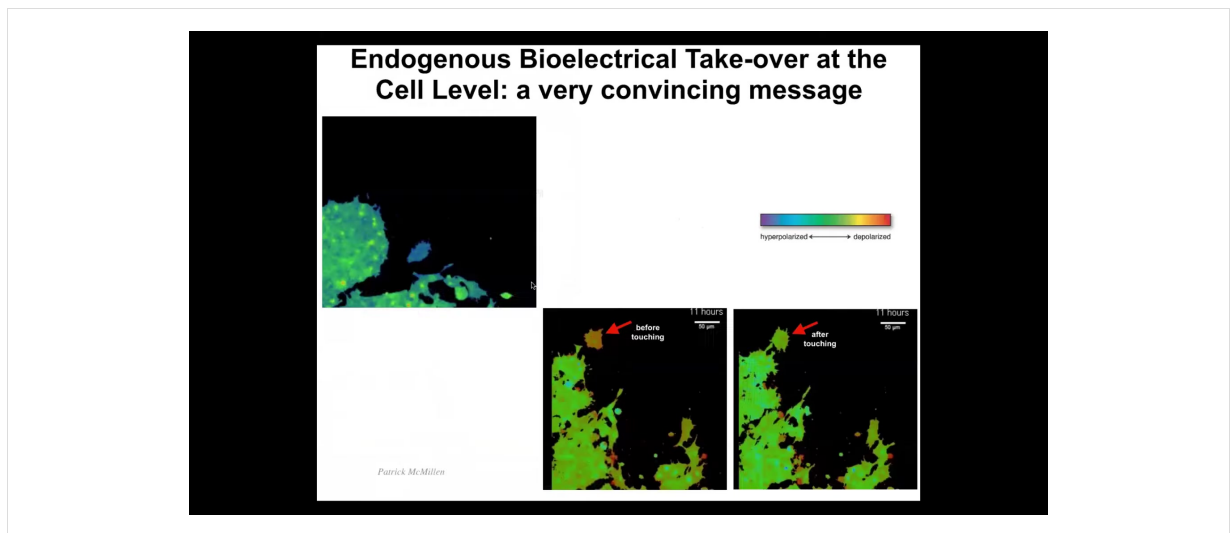
Same exact stuff that the neuroscientists do, except we do all of this outside the nervous system.

Fun story that Rob might appreciate. When I was a postdoc, I wanted to do this, and I started writing to people, to neuroscientists, asking them for plasmids and coding all these different channels. At first I was dumb enough to tell people why I was doing it. I said, I'm going to express these things in non-neural cells to guide decision-making and morphogenesis. At one point, my boss, my postdoc mentor, Mark, came to me and he said, I just got a letter from some neuroscientist, he's warning me that my postdoc is unhinged because he's writing me these crazy requests. He warned the guy that I might be dangerous.

This is the level of surprise that this engendered in neuroscientists, this idea that these channels are going to do anything else outside the nervous system, other than maybe some kind of supportive housekeeping. It was considered a completely crazy idea, but what I did is I would write to them, then I stopped telling them why, and I would just write to people and ask for these different channels that I was then going to use to control cell voltage in these settings.

Before I show you what happens when we do it, I want to show you a natural example.

Slide 15 of 42 · Watch at [36:43](#)

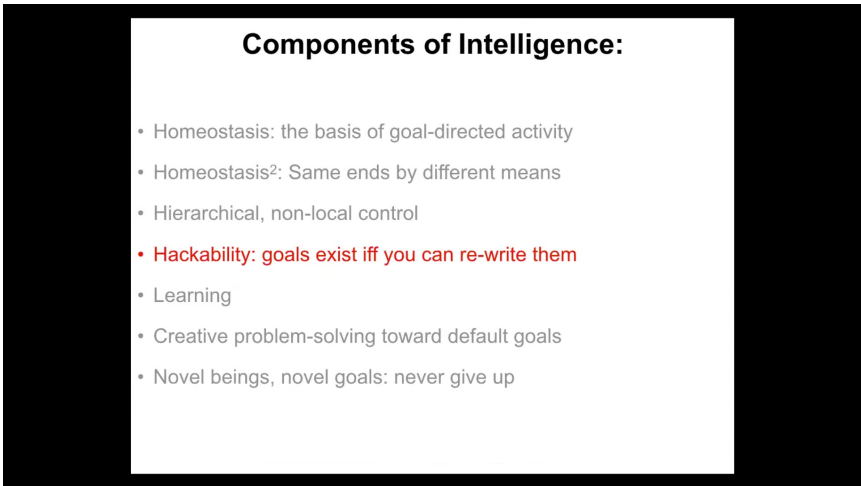


These are a natural example of cells hacking each other. These are cells in 2D culture. The colors are voltage indicators. What you see is something very interesting. This cell has a different voltage than all these other cells. In the next frame, it touches this little projection. As soon as it touches, boom, these guys have rewritten its electrical state.

Watch. Here it comes. It's different. It's moving around. This one's going to reach out and touch it. Bang. There. It's now one of them. Here it comes. It's going to participate in this event.

The ability of cells to control this kind of electrical state. This is what it naturally looks like at the single cell level. What I'm going to show you now is the use of this interface, and that's what I think bioelectrics is. I think the reason it's interesting is because it's this amazing interface to the teleology.

Slide 16 of 42 · Watch at [37:39](#)

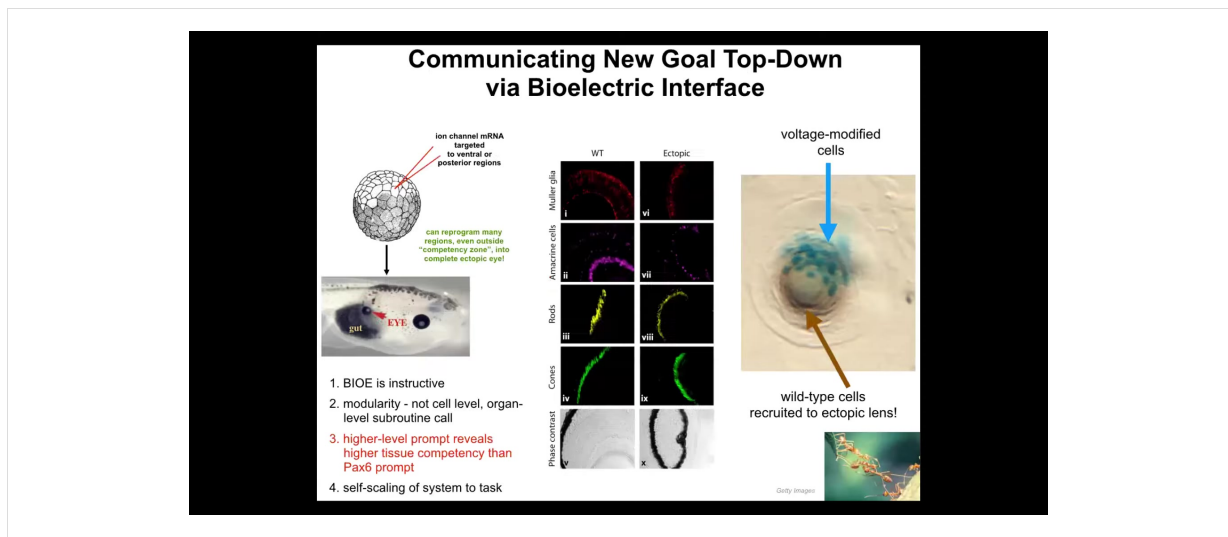


Components of Intelligence:

- Homeostasis: the basis of goal-directed activity
- Homeostasis?: Same ends by different means
- Hierarchical, non-local control
- **Hackability: goals exist iff you can re-write them**
- Learning
- Creative problem-solving toward default goals
- Novel beings, novel goals: never give up

Here's my claim: goals exist if and only if you can rewrite them. If I can't rewrite, if I can't change the goal and have the exact same hardware execute a totally novel goal, then I can't say that there's a goal.

This is back to my claim that this is not a philosophical debate. This is practical. This is my criterion: can I hack these things specifically by rewriting goal states?



I'll show you an example. Here's an early frog embryo. What we're going to do is inject some ion channel RNA that we chose. There are many RNAs, many ion channels you could choose as long as they set the voltage to the correct state. This is not specific to the molecular biology. It doesn't care which channel, doesn't even care which ion it is. As long as you get the voltage right, it works. We're going to inject that RNA into some cells that will become various parts of the body, for example, in this case, the gut. So I inject the RNA. Sure enough, these cells will form an eye. I already showed you that there's a very particular voltage pattern that induces eyes. Here's the eye. If you section them, they have lens, retina, optic nerve. They have all the same stuff they're supposed to have.

Now, notice a few really interesting things. First of all, it tells you that this bioelectrical signal is instructive. We didn't just screw something up. We actually called up a coherent entire organ. So it's instructive for what happens. It's not permissive if it's instructive.

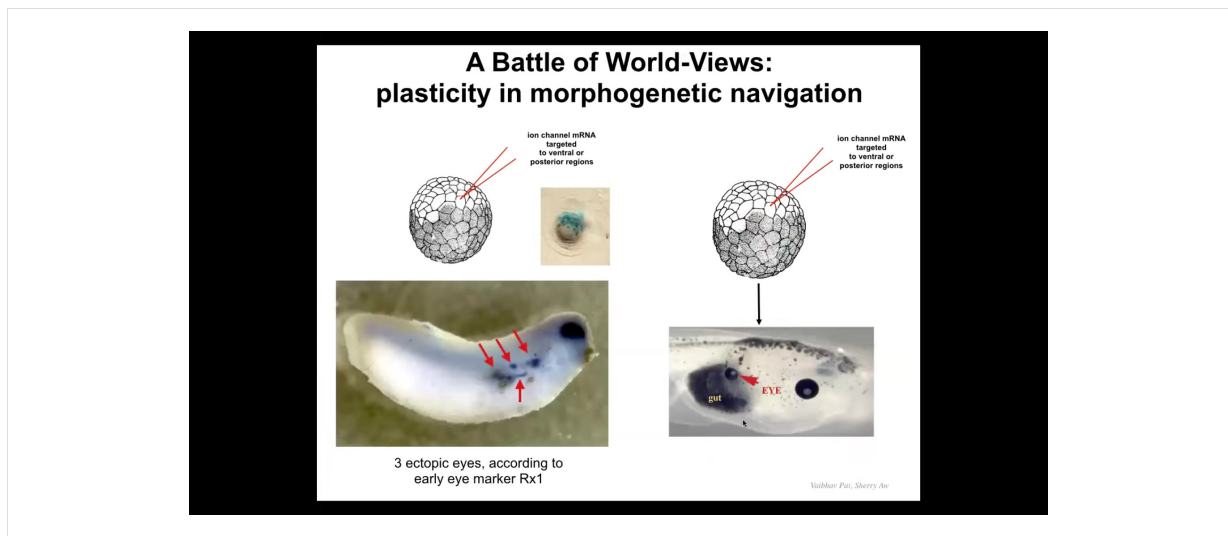
Second thing, it's incredibly modular. We don't have any idea how to build an eye. We don't know how to tell the cells which stem cells go where and what layers go where. We don't have any ability to do that. It's a very complicated process. We've discovered a very simple message, a high-level trigger, that when the cells receive it, they will organize all the downstream stuff themselves to build that eye. This is what I was saying before about multi-scale cognitive systems. Just like when I talk to you, I don't worry about the molecular biology of your brain. Same thing here. We gave it a voltage message that said build an eye and everything else happened after that. All the genes turned on and off. The system takes care of all that. So it's very modular.

Next thing I'll tell you is that if you look at the developmental biology textbook, it actually says that in vertebrates, only the anterior neurectoderm, so this stuff up here,

is competent to form eyes. The other tissue can't do it. The reason they say that is because up until then, and this was around 2010 or 2012, everybody had been probing these things with what they call the master eye gene or PAX6. And sure enough, if you only use PAX6, these are the only cells that are willing to make an eye. But the competency problem wasn't the problem with the cells. It was a problem with us, the scientists, because we weren't using the right prompt. With any intelligent system, if you don't see its intelligence, it's a two-way IQ test. It might not be the system. It might be you. We just weren't using the right prompt. If you use the right prompt, pretty much any region of the tadpole can form an eye. That's a humility warning for us.

The final thing I'll point out is something cool. If you take a cross-section through this eye, for example, the lens, the blue cells are the ones that we injected. What about all this other stuff? All these other cells participating in this formation, we never touched them, we never injected them. They got recruited by the cells that we did inject, and basically these cells can tell there's not enough of them to finish the eye, and they do a secondary recruitment to get these other cells to play along. Again, competency of the material: we didn't teach them to do that; they already knew that. There are many other collective intelligences that do this. Ants and termites, when they encounter something that's too heavy for them to lift, call up the rest of their buddies to help them. The system self-scales to the task.

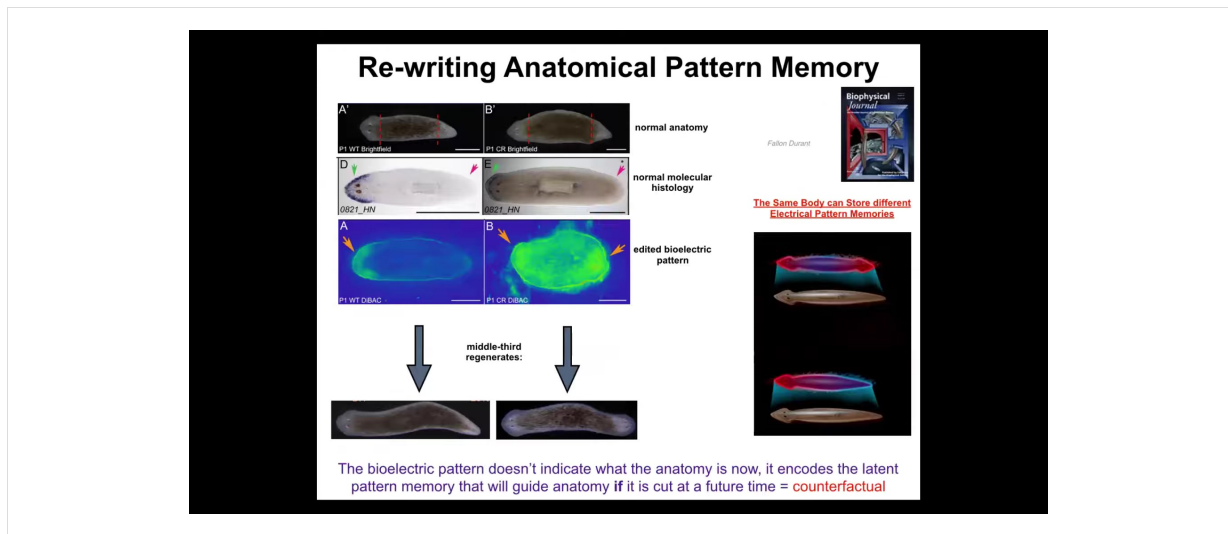
That's one thing. There's something quite interesting here.



When we first did this, we noticed that if I inject my channel, I can get embryos with, so this is an expression of a gene called Rix1, and Rix1 is an early eye marker. So you see, okay, here's the native eye, but here, one, two, three, four, there's a bunch of these. You might think, okay, amazing, this embryo is going to have 5 extra eyes. But you look the first day and there's five or six, you look the next day there's only three, and you look the next day there might be one. What's happening to this?

It turns out that there's a battle of chemical signals, but I actually think the more interesting aspect of this is that it's a battle of goals, it's a battle of worldviews. Because the cells we injected are saying to their neighbors, we should definitely be an eye, help me build an eye. But there's a cancer suppression mechanism that functions this way. If you're a cell and your neighbor shows up with some weird voltage, you try to normalize it. It's a continuous cancer suppression mechanism. So you use your gap junctions to connect and you wipe out their crazy voltage. So what happens here is that the patterns fight. They compete. And there's a pattern that says, no, you should be skinned. And these guys are saying, no, we're going to be eye. And sometimes the eye wins and sometimes the skin wins.

This reminds us that for regenerative medicine applications, you need to be able to speak the language. You need to be able to give these kinds of prompts, but you have to make the prompts convincing. This is not like rewiring a mechanical clock where you micromanage everything and put it where it goes. You are dealing with a decision-making system where you have to get the buy-in of the cells by re-specifying their set point. And sometimes it'll take and sometimes it won't unless we really understand how to make it stick.



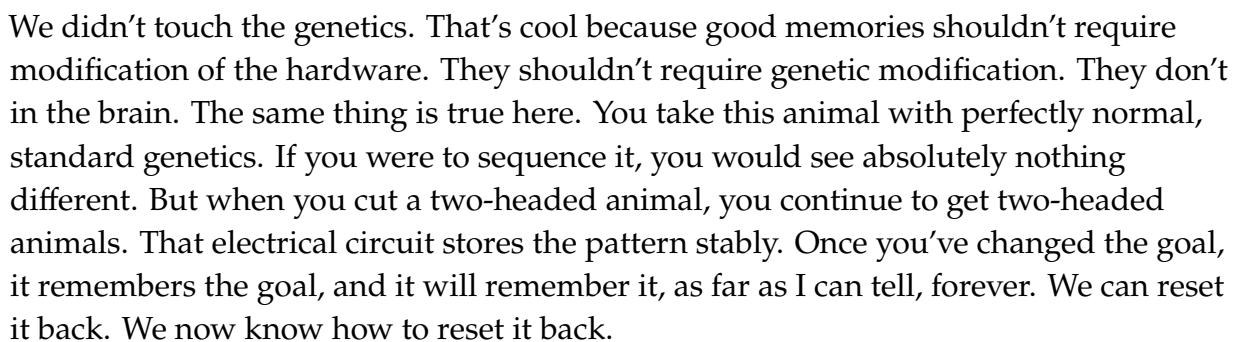
So I want to show you another wild example of rewriting this anatomical pattern memory. These are planaria. These are flatworms. One cool thing about these flatworms is that if you cut them into pieces, every piece will regenerate very reliably. Here's a worm, here's a head, here's a tail. I cut off the head, cut off the tail. This middle fragment here will give me one perfectly normal worm. You can ask the question, how does this thing know where the head goes? It's got two wounds. Back here, the cells back here need to make a tail, but the cells over here need to make a head. So they're very close to each other. You can't tell from the local position whether you should be a head or a tail. You don't know because it could be either one, depending on what the rest of the fragment is doing. So it's actually quite interesting, how do you know how many heads you're supposed to have? It turns out this thing has a particular voltage pattern that says one head, one tail. We ended up figuring out how to change that voltage pattern.

This is an old picture, and it's a mess. We targeted the ion channels to produce a pattern that says two heads instead of one. When you do that, nothing happens at first. Here's a normal animal. It is normal according to the molecular biology markers. In other words, the head marker is on in the head. It is not on in the tail. These cells do not think they are a head. Everything sits quietly until you injure the animal. Once you cut them, that's when this latent memory becomes a recalled memory, and that's when they build a two-headed worm. This is not Photoshop, this is not AI, these are real animals.

This is quite cool. This bioelectric map is not a map of this two-headed animal. It is a map of a perfectly anatomically normal one-headed animal, and it is a latent memory that only becomes activated if I injure the thing, but then it serves as the new set point. The new memory of what a correct planarian looks like, the new goal, becomes active

So, two kinds of interesting things. A normal planarian body can store at least one of two internal representations of a future goal state of what it's going to do if it gets injured. That's a counterfactual memory. I think it's actually an evolutionary precursor to our brain's amazing ability to time travel, to have thoughts about things that are not happening right at this moment, either past or future. I think that's what it is.

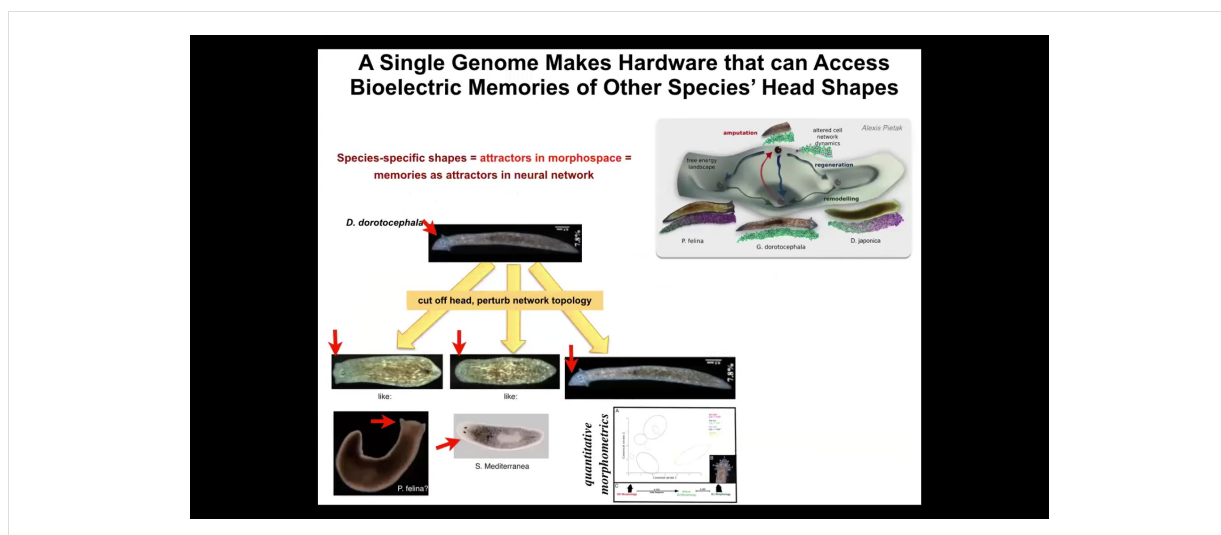
Slide 20 of 42 · Watch at [46:25](#)



25

what the genome does, namely specify the number of heads you're supposed to have," and clearly that's not the case. I bring these videos, although now with AI, anything's possible. This is what I've been showing people. No, those animals absolutely exist because the number of heads, like many other interesting things, is not encoded in the genome. They are actually an active memory that is maintained, in this case, by a bioelectrical circuit.

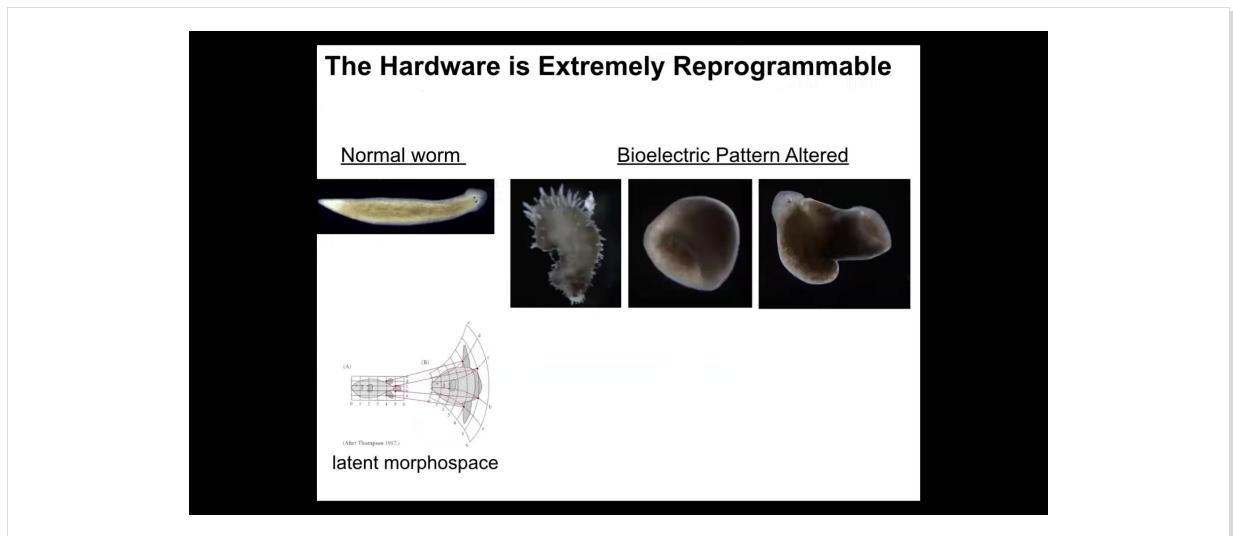
Slide 21 of 42 · Watch at [47:34](#)



It is not just about head number. It is also about head shape. We can take this planarian with a nice triangular head and cut off the head, perturb the bioelectrical network that guides the movement of the cells in anatomical space.

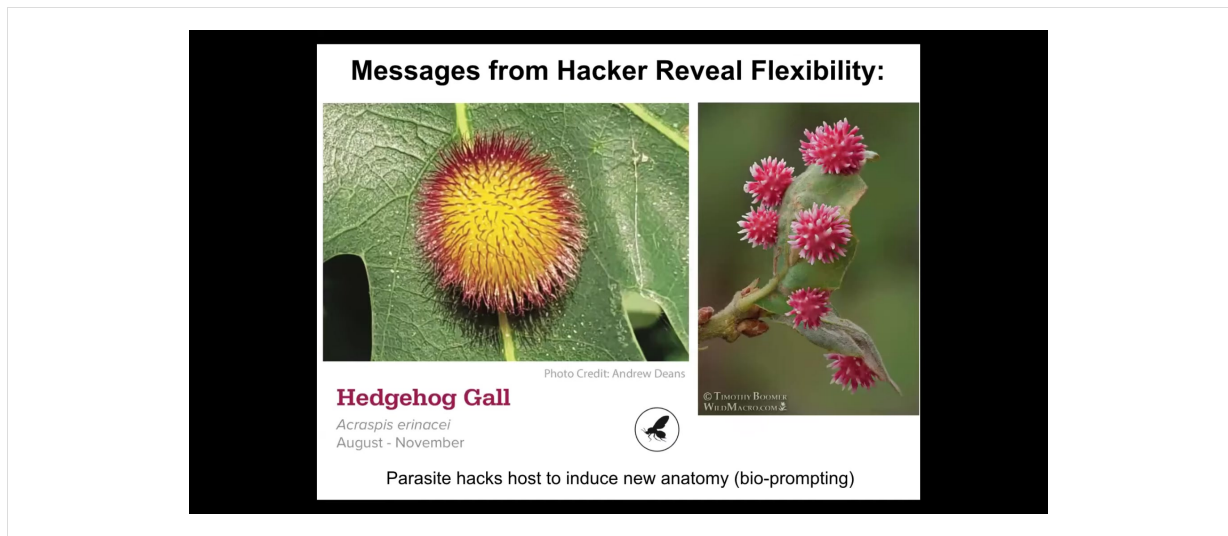
Normally you're supposed to go to this area, this attractor where duratocephalus live, but that same hardware is perfectly happy to go to other attractors. It can build a flat head like a P. falina. It can build a round head like an S. Mediterranean. There is 100 to 150 million years evolutionary distance between these guys and this guy.

The insides of the brain of the head also change. They have brains like these other species. They have a distribution of stem cells like these other species. Again, no genetic change, but if you change the pattern memory, they will use it to navigate to a different region of that amorphous space where typically other species hang out.

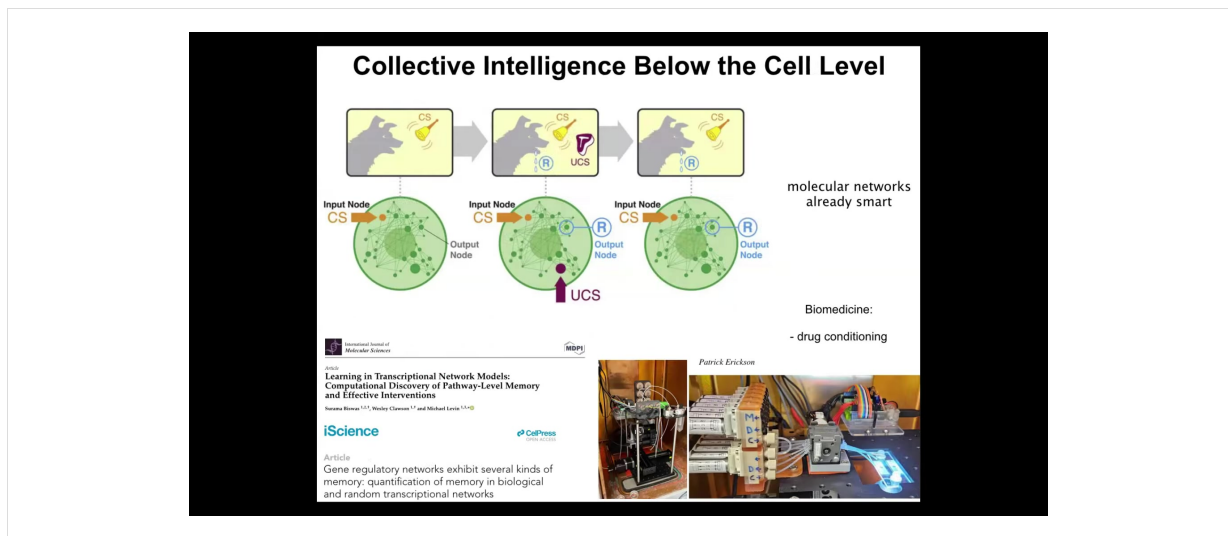


You can even go further and find out that actually the hardware is really plastic. You can make things that don't look like planaria at all. You can make these weird spiky things and make round, a totally different symmetry type, and these composite things. This latent morphospace is massive. This is one of my favorite examples of this kind of thing with plants.

These are acorns. You would look at this and you would say, we know what the oak genome encodes. It encodes this, it encodes this kind of leaf.



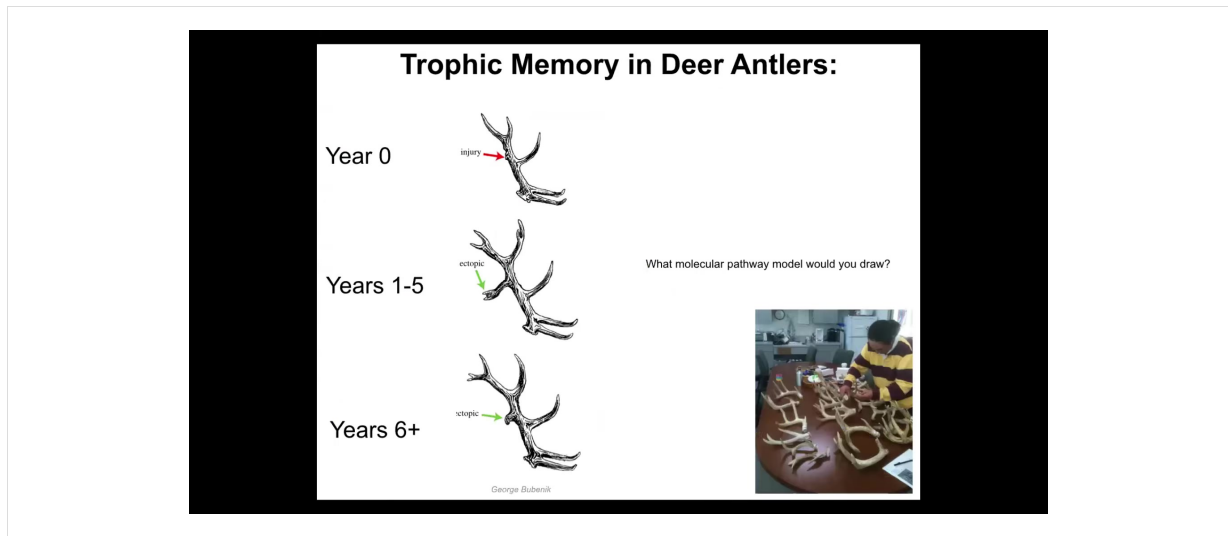
Along comes a non-human bioengineer, this little wasp, and because these circuits are so reprogrammable, the wasp can put down signals that will cause the plant cells to build something like this is called a gall. And you can see, instead of this normal flat green thing, it makes these round, spiky, incredible other constructions. We would have had no idea that this is even possible, except that this thing found an example. And one of the things we're doing in our lab is trying to automate this process to crack the morphogenetic code so that we can call up structures at will, so we can program them.



I want to talk about two examples of learning. Previously, I showed you that we directly rewrote or edited the memory that guides goal-directed activity. We specifically changed those memories, but I want to look at the way the system has the ability to change its own memory. In other words, learning. Learning is the ability of systems to write their own memory n-grams.

The first one is at the molecular level. We found that even gene regulatory networks alone—never mind the cell, the nucleus, any of that stuff—just the gene regulatory network itself, a set of genes that turn each other on and off, and any kind of pathway that has elements that turn each other on and off, are able to show up to six different kinds of learning. Association or Pavlovian conditioning, and also simpler things like habituation and sensitization and counting to small numbers. Even the molecular networks inside cells are able to do this, and we are building devices in our lab to take advantage of this for biomedical purposes like drug conditioning.

My claim is that the material is a multi-scale, agential material: even below the cell level, you already have some examples of learning, which can then be harnessed by the higher levels of organization and underlie more complex forms of goal-directedness.



To contrast that, a very large kind of scale example of learning, and this is called trophic memory in deer antlers.

A father and son team did these experiments for about 30 years in Canada in a herd of deer. An inconvenient model system; it's amazing that this kind of data set will probably never be gotten again. What they found is that every year the basic pattern of the antlers remains constant and the whole thing is shed and then they regrow a new rack. If you have a bit of damage at one particular location that breaks through the skin and etches into the bone, there's a little callus in the heels and that's that for that year. Eventually this whole thing falls off. Next year they grow a brand new tine, a brand new antler rack rather, that has an ectopic tine at that location. There's great specificity here. This goes on for about five years. Every year it has this ectopic branch, and then eventually it gets smaller and smaller and disappears.

Think about what this means and how you might try to draw a molecular pathway model, the kind of models that are popular in the cell biology textbooks where there's some arrows and some things binding, some other things and so on. What possible pathway could you draw to explain that this thing somehow conveys to the cells, presumably the cells at the base of the scalp, where the damage occurred? Those cells store that information for months until it gets ready to build a new antler rack. Then it propagates that information back up to guide growth toward a specific outcome, a very specific outcome, next year.

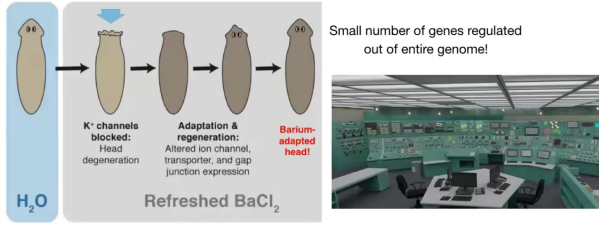
That ability to form memories of location in three-dimensional space in the structure is remarkable. We don't even have the beginnings of models in molecular biology of how you would do this, but in neuroscience, you can imagine models around bioelectric networks remembering locations in space.

This is another example. This guy, apparently I was the only one who was writing about these things in the modern literature. When he retired and needed to get rid of all the antlers, he sent them to me. I have 13 boxes of these things; each is labeled Lenny 1987, 1988, and so on, in terms of these multi-year experiments—just an unbelievably remarkable biological object that again shows how you really need concepts from these other fields related to cognitive science in order to be able to understand these kinds of phenomena.

The next thing I want to show after learning is creative problem solving. The system has a default goal, presumably conveyed onto it by evolution, that sets up a default memory pattern that's being used as the set point for the homeostatic activity. We're going to look at different ways to get to that goal under novel circumstances.

Slide 26 of 42 · Watch at [54:28](#)

Solving Novel Physiological Problems by Navigation of Transcriptional Space



Small number of genes regulated out of entire genome!

- planarian heads degenerate after exposure to barium
- planaria eventually adapt and regenerate heads that tolerate barium
- a relatively few transcripts were altered to produce barium tolerance
- how did the system choose exactly the right genes to modulate, to deal with this evolutionarily-novel challenge?

The first example that I want to show you has to do with planaria. We discovered this a few years ago. Here is a flatworm. You put the flatworms in a solution of barium. Barium is a non-specific potassium channel blocker. It blocks all the potassium channels. The heads are very unhappy. It's got neurons and so on. They're very unhappy about not being able to pass potassium. The whole head explodes. Basically, it has blown off within 24 to 48 hours — just gone.

Then something amazing happens: if you leave them in the barium, they eventually regenerate a new head, and the new head is totally fine in the barium. It's somehow barium-adapted or insensitive. We did an experiment to try to figure out how this was

possible. We did RNA-seq to look at all the genes that were expressed in barium-adapted heads versus normal heads, and we did a subtraction and asked what was the difference. It's really only a handful of genes that are different.

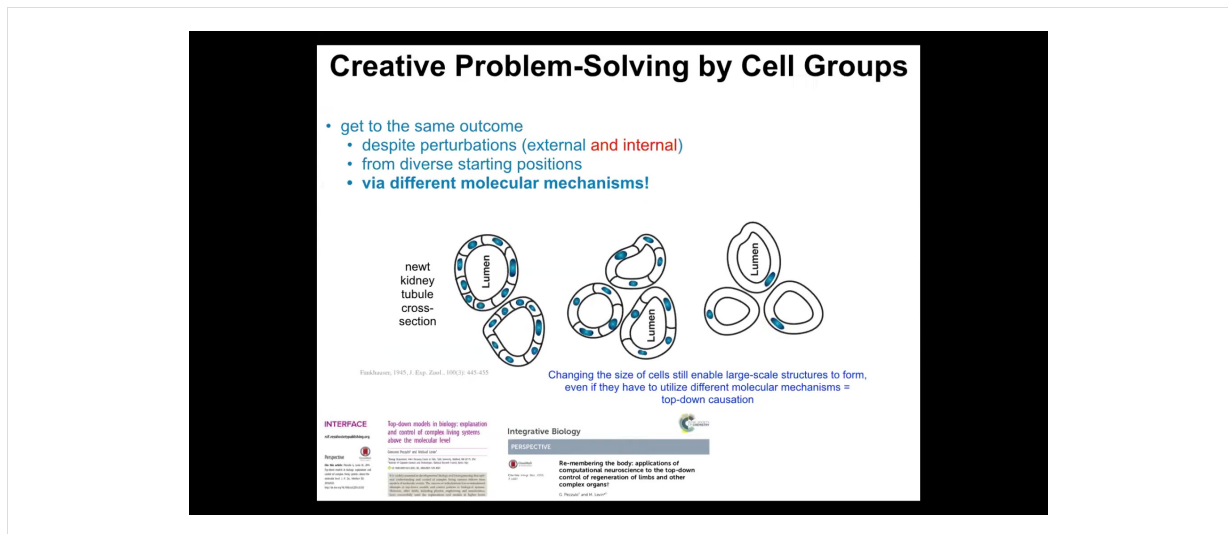
I'm always amazed by this example, and I visualize it like being inside a nuclear reactor control room and the thing's melting down, and there's a million different buttons and knobs, and figuring out which ones I need to flip in order to solve my problem.

Think about what's happening here. They're under a physiological stressor, one for which evolution did not really prepare them. Planaria don't see barium in the wild. They've got this novel physiological stressor, and they've got tens of thousands of potential actions they can take, those being genes that they could turn on and off. Those are effectors in transcriptional space. They're looking for a solution, a behavior in transcriptional space, meaning turning certain genes on and off, to solve a physiological stressor.

Very rapidly, they find it. They don't have enough time for any kind of random search or even hill-climbing, because it happens pretty fast and their cells don't turn over that fast. It's not like bacteria where everybody can try a different strategy and then most of them die and then somebody survives. Planaria cells don't turn over like that. They really need to find a solution in place.

If this is the most recent data, if we do this in a collection of worms kept in separate dishes, they all find the same solution. Apparently there aren't, or at least not easily found, multiple solutions to this problem. They all roughly turn on and off the same handful of genes.

How is it that they navigate that transcriptional space to find a solution to a problem they probably have not seen before? Certainly these worms haven't seen them, but maybe the lineage has seen something similar. For example, maybe epileptic seizures may have some similarity, but certainly not exactly this. We have no idea how it finds so rapidly the effectors that it needs. I consider this a really interesting example of problem solving, of resolving a novel stressor using the genetic affordances you have. I don't know if they have some sort of internal self-model through which the cells can figure out what actions they should take to resolve specific stressors. I'm not sure. This is a big open question.



Another example that I really like is this: it is a morphogenetic one.

This is a cross-section of a kidney tubule. It's got a lumen and a bunch of cells making the actual tubule. What you see is that there's maybe 8 to 10 cells working together to form this kind of structure normally.

Now you can make polyploid newts, which is basically taking the fertilized egg and performing a procedure so that it ends up with more copies of chromosomes. So it has extra copies of the genetic material. The cells get bigger to accommodate the bigger nucleus, but the newt stays the same size. This is Fankhauser's work. That's amazing. When you look to see how that could be possible, you see what's going on. It uses fewer yet larger cells. The number of cells adjusts to the size in order to give you that same structure.

Something amazing happens when you make highly polyploid newts. I think these are 6N or something like that. The cells get truly gigantic. Then just one cell will wrap around itself to still give you a lumen inside.

There are a few amazing things about this. First, notice the adjustments in the adaptation. They adjust the cell number to the new cell size. When needed, they pick a completely different molecular mechanism to do so. In one case they're using cell-to-cell communication and tubulogenesis. But here it's cytoskeletal bending. It's just one cell.

There's an amazing example of top-down causation. The collective is working towards an anatomical goal, and it's able to not only instruct the cells what to do, but it's able to pick different molecular affordances in its toolkit to get the job done.

This is a standard example used in IQ tests, where they give you a set of objects and ask you to use these objects to solve a particular problem.

Imagine what it's like to be a newt coming into the world. Never mind the environment and the fact that you can't really know what's going to happen with your environment. You can't even trust your own parts. You don't know how many copies of your genome you're going to have. You don't know how big your cells are going to be. You don't know how many cells you're going to have.

Other work showed a kind of invariance to the number of cells that you start with. You have to get the job done with the tools you have, regardless of what's going on, not only in the environment, but in your own parts. Even your own parts are unreliable.

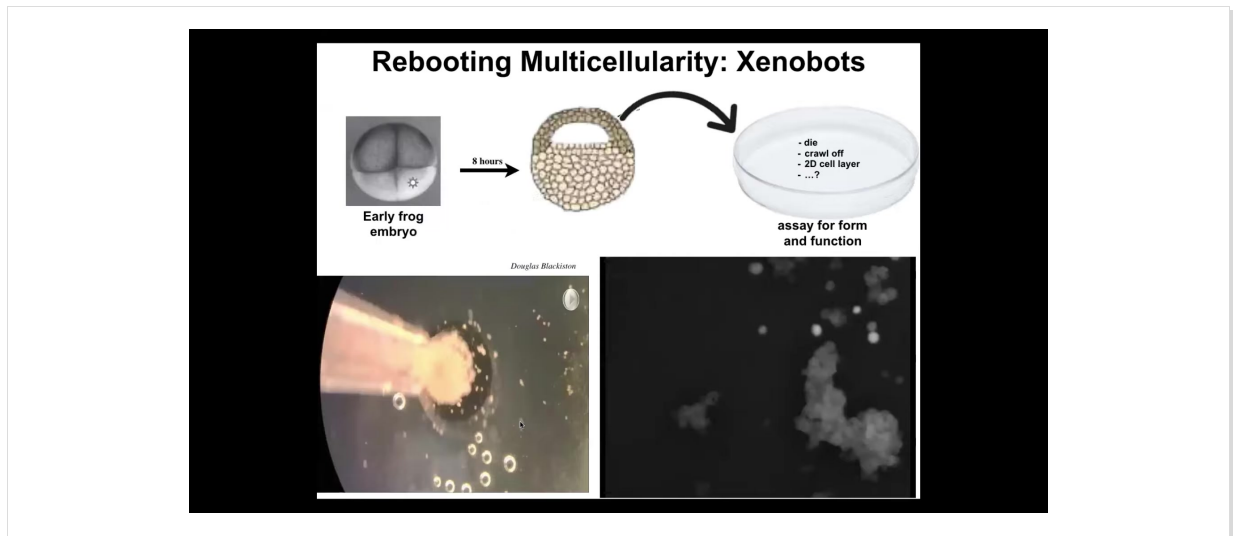
Biology as an unreliable medium turns out to be very important for the development of intelligence.

This is not only goal-directed, meaning that under all of these circumstances, the system makes the same large-scale shape, but it uses different ways to get there when circumstances change in a radical way.

The other thing about life and intelligence is that beyond the basics of achieving a single goal and then stopping, there's the next level. That was basic homeostasis. I showed you a kind of second-order homeostasis where you're able to reach the same end by different paths. If you can't reach your same goals, or if you're a novel creature that's never had the benefit of an evolutionary history that selects for specific goal states, what goals will you have? Do you have the ability to form novel goals and find new forms and behaviors that you can adopt?

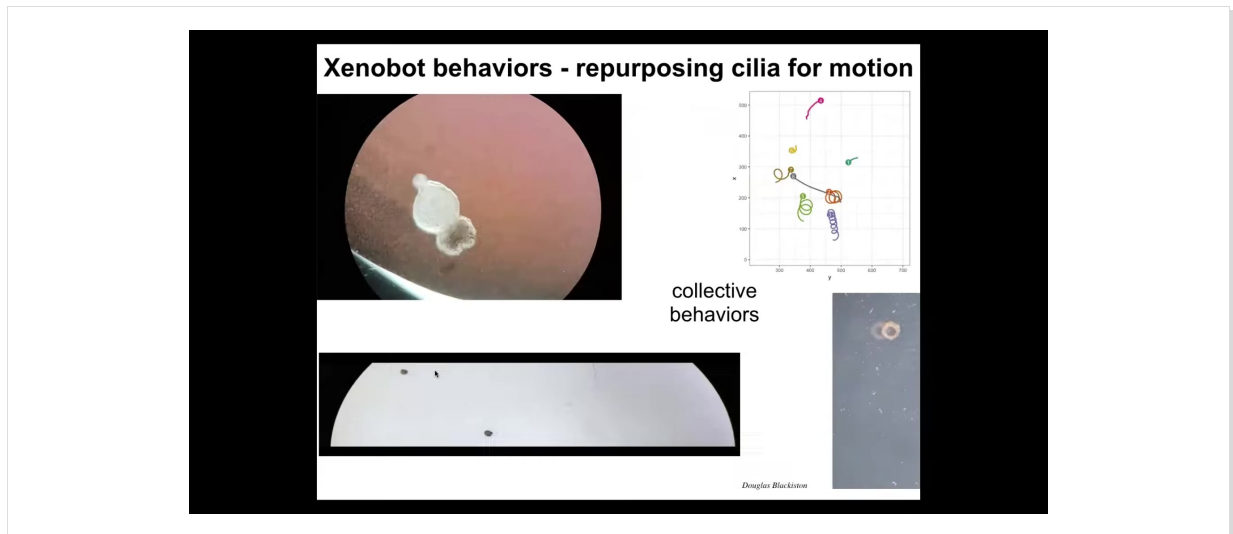
I organize this in terms of progressively more sophisticated cognitive capacities.

Even beyond creative problem solving toward default goals, let's look at some novel beings and this idea that life is good at picking new problems, not just solving old problems.

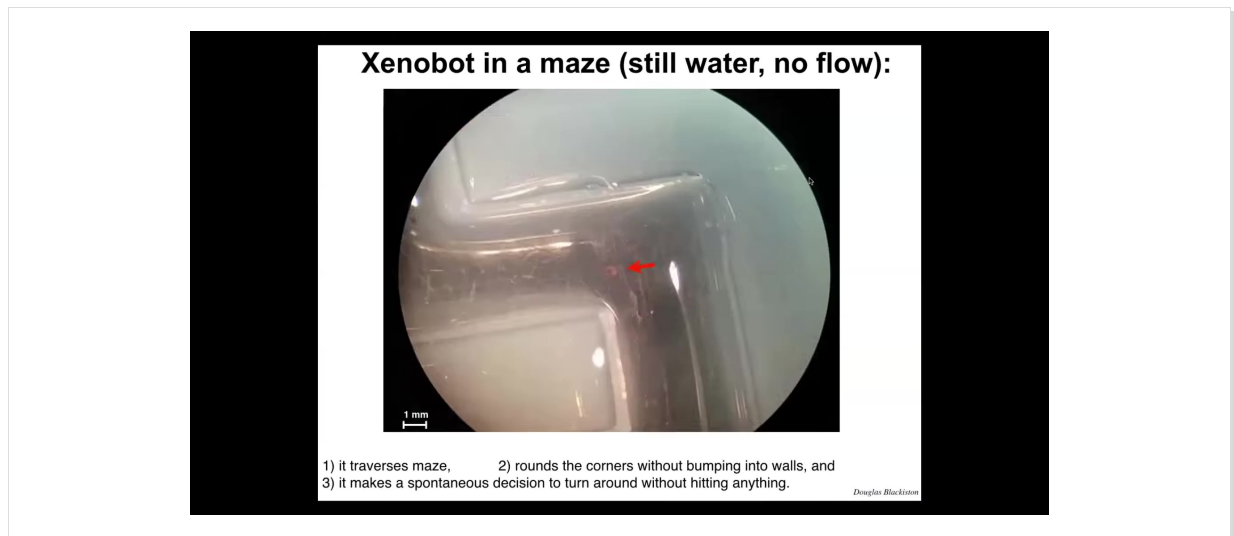


The first thing that I'm going to show you is xenobots. Xenobots are self-organized little multicellular creatures, they're biobots as well, composed of frog epithelial cells. In the early frog embryo, we take these cells, we dissociate them, we put them by themselves here. They could have died. They could have crawled off. They could have made a flat monolayer cell culture. Instead, what they do is they coalesce together and they make this. The way they combine is actually amazing.

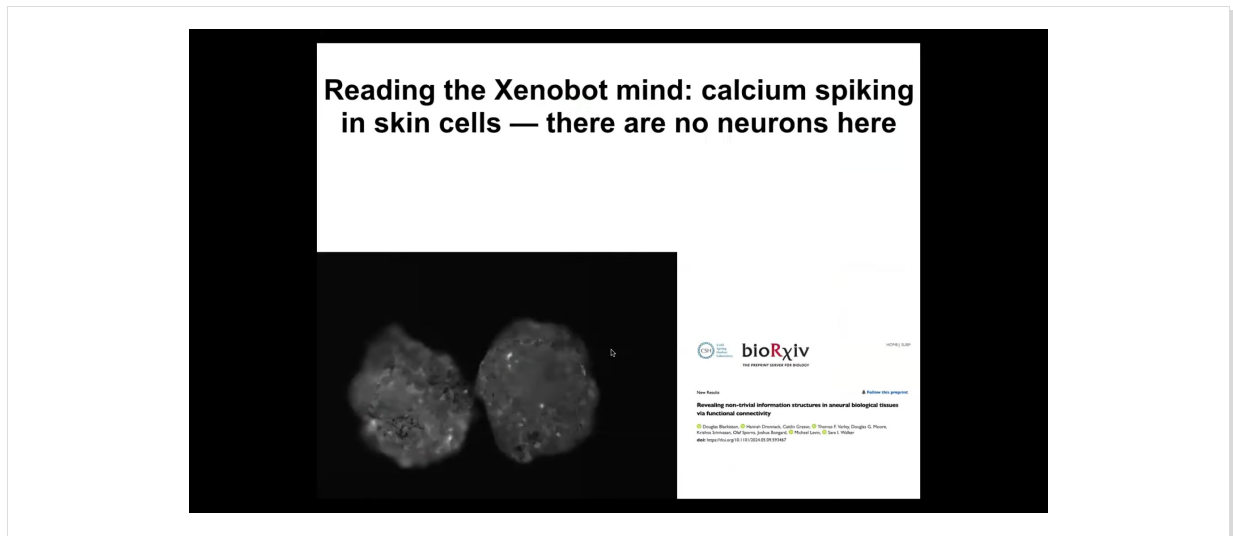
Take a look here. Each one of these circles is a single cell. Here they are. Look at this little clump. It's fun that it looks like a little horsey. It wanders over here using some motility mechanism we don't understand. It's sniffing around and then eventually you get this little calcium flash of a communication event. All of these things, if you could stare at this all day, there's all kinds of amazing patterns, but eventually they all coalesce into this thing.



It has some interesting behaviors. The outer surface is covered with cilia, and it organizes the cilia so that they can swim. Here it is swimming through some particles. It can go in circles. It can patrol back and forth like this. You can make them into weird shapes like these donuts. They have collective behaviors, or group behaviors; these two are interacting. These are sitting there doing nothing. This one's going on a long journey.



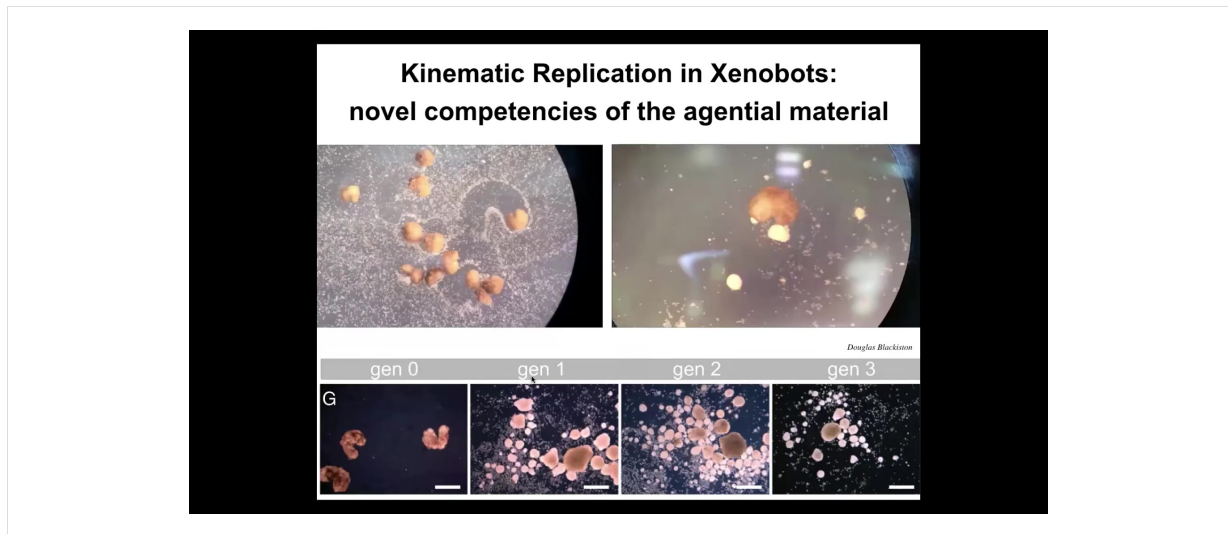
Here's one. Traversing a maze. There's no flow here. There are no gradients. Watch. It takes the corner without bumping into the opposite wall. At this point, it spontaneously turns around and goes back where it came from. They have all these interesting behaviors after they self-organize.



Interestingly, they have some very lively calcium signaling.

We have been in collaboration with Josh Bongard's lab and Sarah Walker. People such as Thomas Varley have been analyzing the calcium flashing and using metrics that are used to analyze brains. Connectivity — that is, information structures — but also, in more recent work, some amazing studies of causal emergence and other metrics that people use on brain activity to try to determine if you're looking at a human, somebody, a pile of neurons, somebody with locked-in syndrome, or somebody who's anesthetized, and so on.

The idea is to use these metrics from neuroscience to understand the collectivity of the system. We've been using that, and it's really interesting. Stay tuned for that, even though there's no neurons here. These are all epithelial cells, and yet those same metrics are revealing some very interesting information structures here.

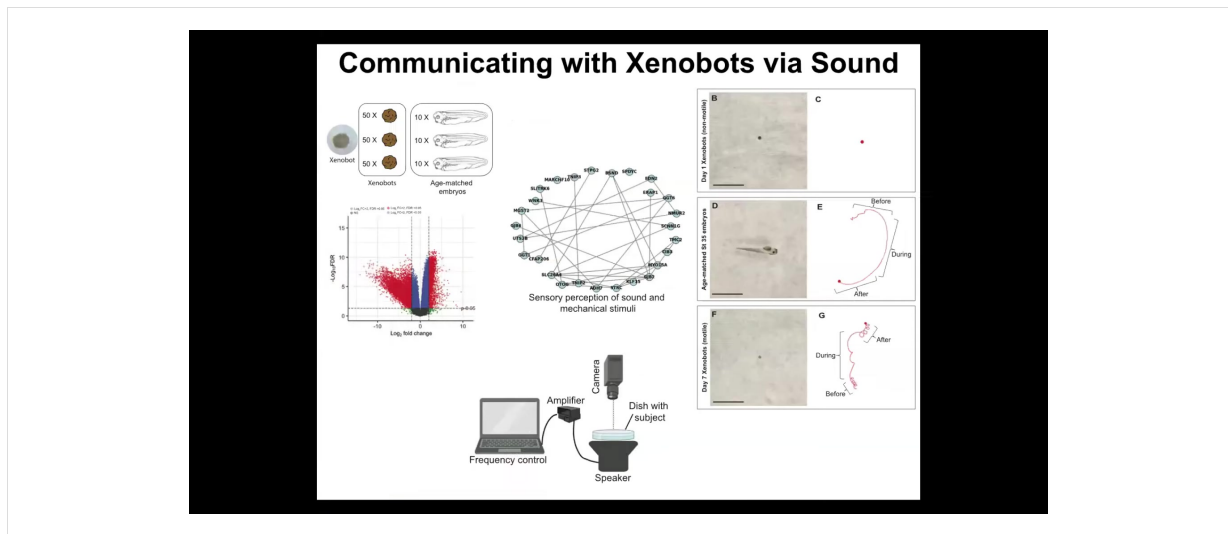


Now, one of the remarkable things about these Xenobots is that, while they can't reproduce in the normal frog-like fashion, they do something quite different. They solve it a different way. If you provide them with loose epithelial cells here, they move around, both individually and as a group, and they push all of them into little balls, and then they polish the little balls like this. Because they're dealing with an agential material themselves, meaning cells, these balls mature to be the next generation of Xenobots. They run around and make the next generation, and these make the next generation, and so on.

So we call this kinematic self-replication. It is von Neumann's dream of a robot that goes around and collects materials from the environment and makes copies of itself. As far as we know, no other creature on Earth does. We don't know of any other examples that reproduce by kinematic self-replication. There have never been any Xenobots. There has never been selection to be a good Xenobot. There has never been selection for kinematic self-replication. This is just something that they do.

I think it's really important that there's a research program to help us get better at predicting these kinds of things, to not be surprised and simply label them as emergent when we do get surprised, but to really understand what are the competencies of the material that we're working with that are not well explained or made predictable by knowing the history of selection under which the genome arose. Because, of course, with everything else I showed you, this is a wild-type genome. So when you ask, what does a *Xenopus laevis*, which is the frog we're working with, what does that genome encode? Well, it encodes tadpoles and frogs, sure, but it also encodes this, and who knows how many other things that have very different lifestyles. And you get all of that as a software byproduct in the biological hardware that genome encodes.

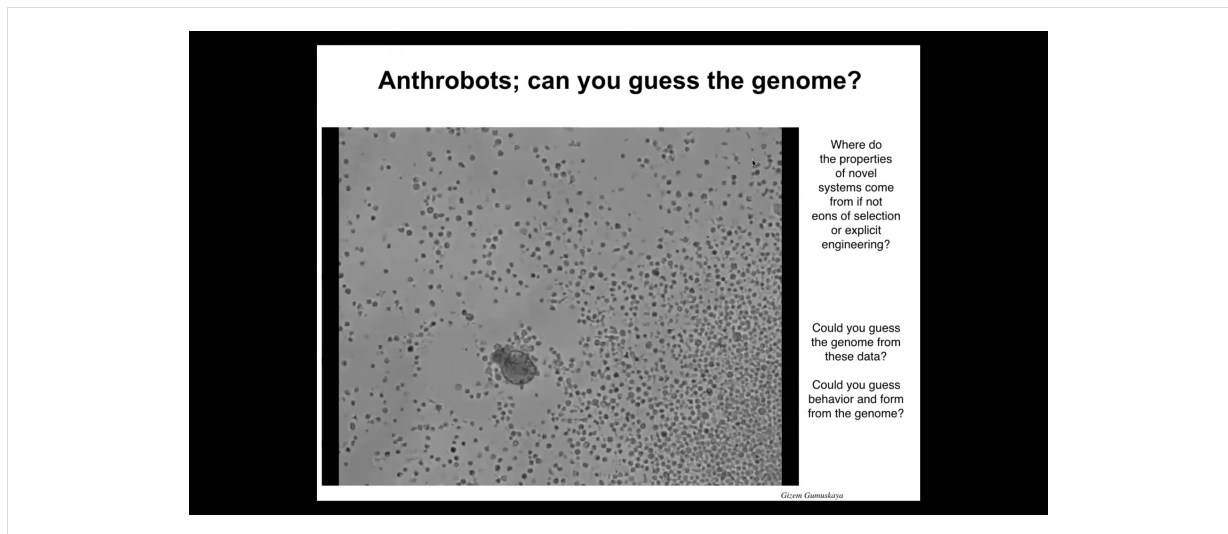
And so we asked one other question recently: these creatures with their new lifestyle have the same genome, but what's their transcriptome look like?



What kind of genes do they express? When we looked at differential gene expression between xenobots and age-matched frog embryos, we found hundreds of genes that xenobots expressed that frog embryos don't express. Going back in the other direction doesn't make any sense because they're missing a whole bunch of genes; they don't have all the tissues that these guys have. We ignored that. We just looked at what new genes they express that normal embryos express at a very low level or not at all. There are hundreds of them. There are many interesting things. One of the things we found in this drastically remodeled transcriptome is a cluster of genes that are normally related to hearing and the perception of sound.

We asked, is it possible that they can hear? We put a speaker under the dish. When you play sounds to them with a speaker underneath the dish, they change their movement. During the sound, they move in a completely different way, and then they go back to almost the same movement after it stops. They can, in fact, react to sound.

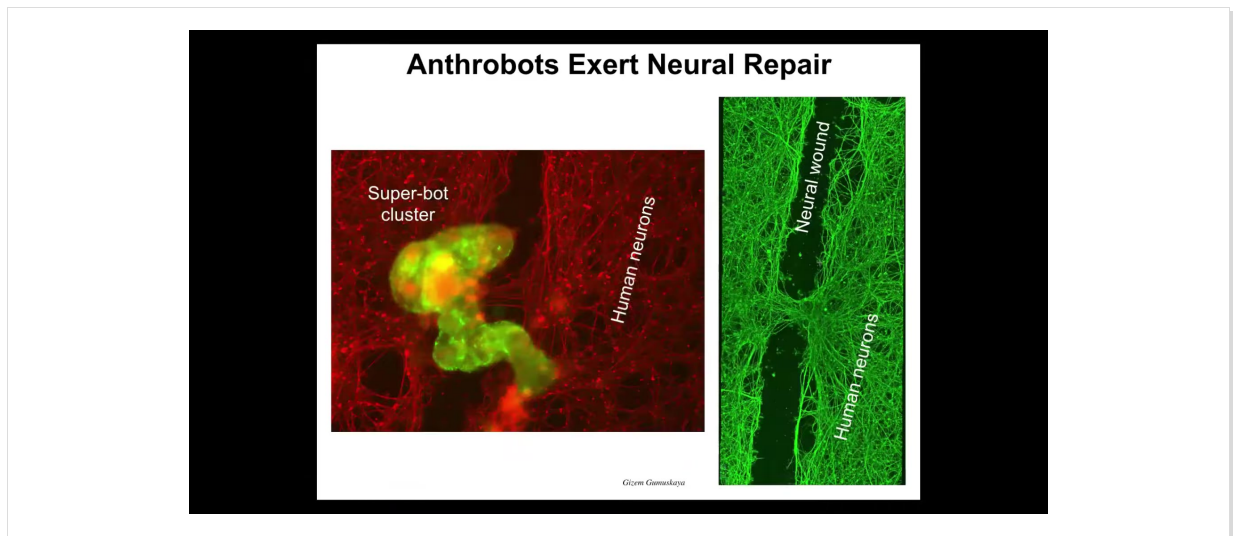
They have a very different transcriptome that is just a function of their new lifestyle. They can dip into their genome and change transcription based on their new lifestyle, not requiring any change to the genome. We don't know what the purpose here is. For example, it is known that coral larvae react to sounds made by different kinds of coral reefs, by the surf hitting the coral reefs. They use this to guide their decisions of where they're going to go. I don't know if that's the kind of thing these guys are trying to do, but this is just the beginning of understanding these transcriptional changes and why these lifestyles are driving specific genes to be expressed in a coherent manner.



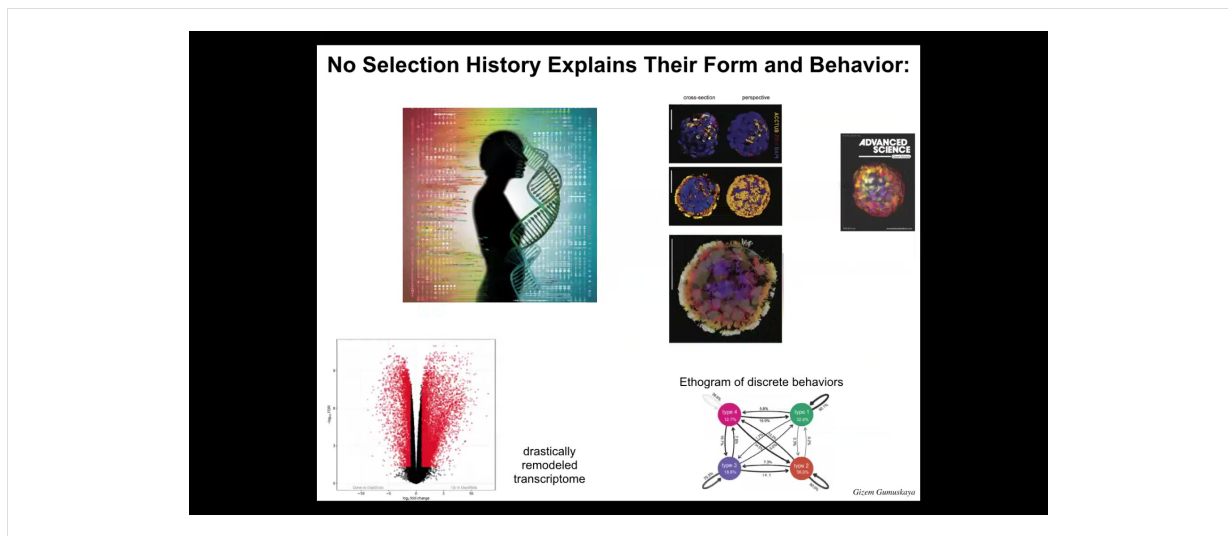
The next thing I want to show you is another biobot called an anthrobot. Now, when you look at this, it looks like something that would have come from the bottom of a pond somewhere. It looks like a primitive organism. If you try to guess the genome based on those criteria, you would be very wrong. The genome here is 100% normal *Homo sapiens*. These are made of adult human cells, unedited; they're made from tracheal epithelia. They self-organize into these motile little creatures. It's swimming here in a group of other cells. We try to do this; it seems to try to do the same thing that Xenobots do, which is to collect these things into a ball, but it doesn't seem to do kinematic replication successfully the way that Xenobots can.

For these and the Xenobots, I'm not making any claims about their specific capabilities in cognition, because we're just now starting to study their memories, their ability to learn, their preferences, the problems they can solve. We don't know that yet.

Morphologically, what's happened is that these cells have formed a new kind of creature that doesn't look anything like any stage of normal human development. It has a completely different lifestyle. It doesn't look like anything that you would expect from human cells. It has a number of interesting properties. It can self-repair. What you see here, after mechanical damage, the anthrobot can repair itself, and it will repair itself toward the new pattern. The kind of homeostasis you see in limbs, in amphibians, and in planarian fragments: this thing eventually repairs itself to the new anthrobot shape.



They don't just repair themselves, they also repair other things in their environment. If you grow a bunch of human neurons on a Petri dish, you take a scalpel, you put a big scratch through it—here's a neural wound. The anthrobots, which are here labeled in green, will collect; we call it a superbots cluster. There's maybe a dozen of them. After about four days they're knitting across the gap. They're healing the injury. Your tracheal epithelial cells, which sit there quietly in your airway for long periods of time dealing with mucus and air particles, have the ability to comprise a novel, self-motile little creature that also knows how to heal neural wounds.



And so, again, this is what it looks like. These little guys have cilia covering their surface. That is how they move. There are several different types. They have four different behaviors. They have discrete behaviors, and we can draw an ethogram of transition probabilities between their behaviors. And even more than zenawatts, they have thousands, about 9,000 transcripts, so half the genome, that they express differently than their tissue of origin. So again, no genetic change, but in a new environment, they dip into their genetically provided toolkit and express thousands of genes differently. We haven't even begun to scratch the surface of what they're actually capable of.

So that plasticity, this idea that you're a bunch of tracheal epithelial cells in a matrix, you're not going to be able to make an entire human, but you can make something else. And you can make something else and then repair to that goal and express novel genes and exert other types of effects on your environment that we're not very good at guessing or predicting up front. So again, no history of selection to be a good anthropot. There's never been any anthropot. So I think that plasticity is part of the spectrum of intelligence that we're all talking about. So I'm going to start to wrap up here, but I just want to mention a couple of other topics that we could talk about.

More: intelligence ratchet and evolution

- Hierarchical control = search the nicer space of inducements, rewards, signals, incentives, etc. not the chaotic space of microstates
- Competencies of modules smooths the evolutionary landscape and buffers negative effects of mutation, giving the process patience (improves evolvability)
- Competency hides genetic information from selection, resulting in intelligence ratchet: further improvements by raising the IQ of the agential material, vs. micromanaging the hardware (genome)



Intelligence climb due to:

- unreliable mechanisms,
- leaky abstraction layers,
- multi-scale competency of agential material,
- self-modifying hardware,
- poly-computing and ubiquitous hacking,
- observer-dependence of meaning of information

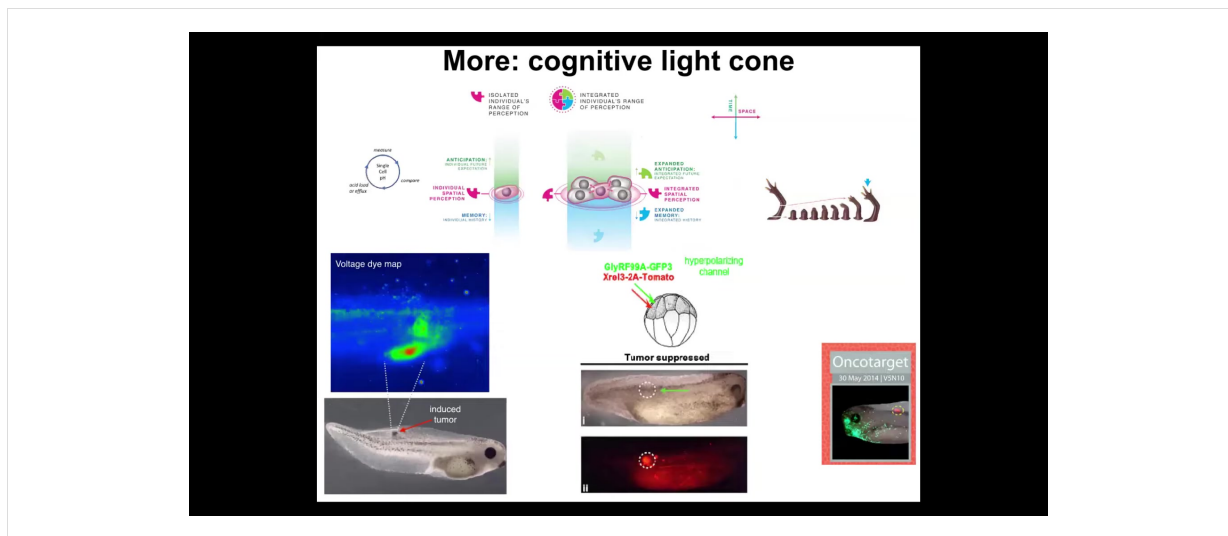
Cellular Competency during Development Alters Evolutionary Dynamics in an Artificial Embryogeny Model

Lakshmi Shrivastava¹ and Michael Levin^{1,2*}

Darwin's agential materials: evolutionary implications of multiscale competency in developmental biology

Michael Levin^{1,2*}

The first is the intelligence ratchet. The competency of the material, the fact that the material is an agential material that does teleological types of problem solving and also many other things makes evolution go completely differently. You can see that in these two papers, we did a lot of simulations and some analysis of what happens when you evolve over an agential material. It ends up making this interesting ratchet that just turns, that keeps cranking up the different types of intelligence. There are implications for evolution.



We could also talk about a story of the cognitive light cone. We define the cognitive light cone as the size of the biggest goals that a system could pursue. Not how far it can sense or act, but the size of the biggest goal it can pursue. You can see what's happening here with single cells having a tiny cognitive light cone that only cares about this region of space and a little bit of anticipation potential, a little bit of memory. But when you gather into collectives, you can have these grandiose goals. These things have the goal of building this limb. Again, it's a goal because if you try to prevent them by cutting it off, it will keep rebuilding it.

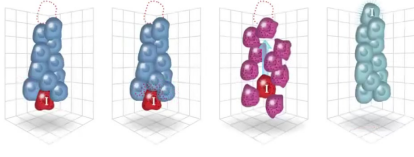
This is part of the story around the scaling of goals in teleology. You can be a teleological system, but you can have little tiny humble goals, or you can have these very large grandiose construction projects as goals. That way of thinking about it has some great implications for cancer therapeutics. We're pursuing ways to specifically understand cancer as a dissociative identity disorder of this collective intelligence. When cells disconnect, you can use that as a cancer detection modality. You can forcibly reconnect them despite their oncoproteins and prevent or reverse tumorigenesis. We did this in frog, and now it's in human cancer spheroids.

More: stress sharing as goal incentive scaling

Tell me what stresses you out, and I can estimate your cognition



$\text{stress} \propto \text{error}$



STRESS TEMPERATURE

LOW HIGH



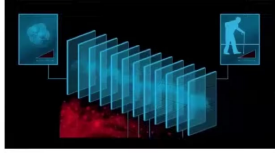
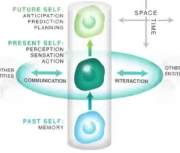
Stress sharing as cognitive glue for collective intelligence: A computational model of stress as a mechanism for morphogenesis

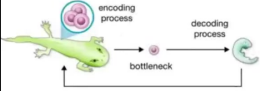
bioRxiv preprint doi: <https://doi.org/10.1101/151111>; this version posted May 10, 2017. The copyright holder for this preprint (which was not certified by peer review) is the author/funder, who has granted bioRxiv a license to display the preprint in perpetuity. It is made available under aCC-BY-NC-ND 4.0 International license.

Export of stress molecules (which have no ownership metadata) result in increased plasticity of tissue - your problem becomes everyone's problem. This implements cooperation toward group goals *without altruism*. Sharing of stress info enables larger-scale goals.

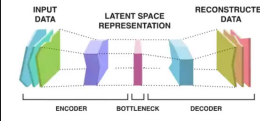
And then the next thing I'll mention is what it is that drives this homeostatic loop. So one thing that you can imagine is that stress and the sharing of stress among cells is part of the engine that keeps cells in this homeostatic cycle. So being away from your goals is stressful. And there are some interesting collective dynamics of how stress scales skills to enable larger systems to pursue larger goals.

More: Memory Engrams and DNA - same paradox



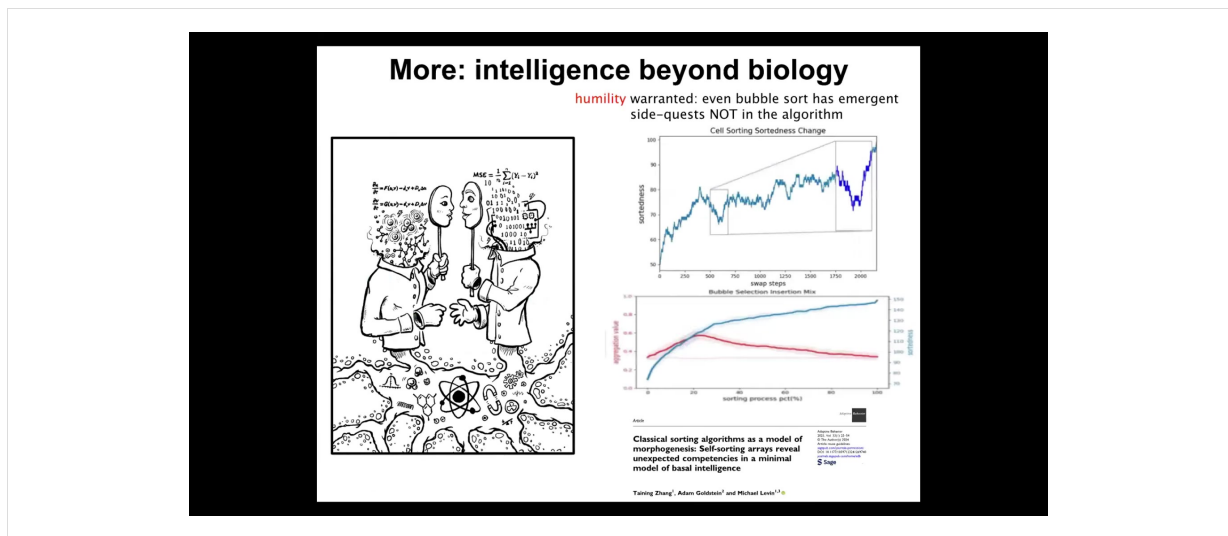
Biology assumes the hardware is unreliable
 Environment and your own parts will change, **don't over-train**
 Little allegiance to past Self's meaning of engrams
 Re-interpret on-the-fly - present/future is all that matters
 Engram is highly compressed - creative remembering, not deduction
 Undependable hardware and confabulation are a feature, not a bug
 Consciousness = active, creative story-telling from own memory traces





Project for:
**Self-Improving Memory: A Perspective on Memories as
 Agential, Dynamically Reinterpreting Cognitive Glue**
 Michael Lewis ©

The final part of this story is that the same issue with an unreliable medium affects both cognition and evolution. In other words, both memory engrams and genetic information have to be interpreted. They cannot be taken literally because of this bottleneck where information is squeezed down into this thin representation and latent space. This is learning. From here, this is the now moment where information has to be creatively reinterpreted. This, I think, is what provides this amazing plasticity. It's because from the very beginning of life, the material could not take any information literally. It had to take them as prompts or suggestions. Evolution cranked on this part, which is the creative, confabulatory aspect of life that takes its information, whether those be your memories or those be your genetic material, and does its best to tell not the same story that was told before necessarily, but the most adaptive story it can tell going forward. All of that is described here. Again, I think this is part of that intelligence ratchet.



The actual final thing I'll say is this: what we've been looking at today during this talk is different aspects of intelligence in the living material.

Unexpected cognitive competencies, like goal directedness, like different kinds of problem solving, are not just the province of complex biological systems. It turns out that even extremely minimal deterministic simple systems. In this paper we studied sorting algorithms, things like bubble sort that people have been studying for probably 60 years or more. We found that even in that system, if you look at it the right way, it's actually doing some very interesting side quests that are not anything that the algorithm is telling them to do.

Even simple systems can provide unexpected navigation of various spaces. This means that we need to be very careful about assuming that large complexity is necessary for some of these things or that "passive matter" isn't doing some of the things that life is doing in this regard. I think we are really bad at noticing cognition in unexpected places. For this reason, I like looking for it in these very surprising, very minimal systems. I think we really need to get better at noticing it in molecular networks and other kinds of substrates.

Thank you to:

Post-docs and staff scientists in the Levin lab:
Douglas Blackiston - brain-body interface plasticity, synthetic living bodies
Laura Vandenberg - craniofacial remodeling
Vaibhav Patel - voltage gradients in eye/brain induction and repair
Patrick Erickson - cellular learning
Patrick McMillen - bioelectric imaging and transitions to multicellularity
Nesher Orvedo - biomechanics of planarian pattern memory

Graduate Students:
Gizem Gumuskaya, Nikolay Davey - Anthrobot
Fallen Durant - biomechanics of planarian pattern memory
Sherry Aw - bioelectricity of eye induction
Angela Tung - cross-embryo communication
Lakshmi Shreecha - computational models of stress sharing & evolution on agential materials
Adam Goldstein, Taining Zhang - emergent competencies of algorithms

Undergraduate Students:
Maya Emmons-Bell - planarian head plasticity, barium adaptation
Pranjal Srivastava, Ben G. Cooper, Hannah Lesser, Ben Semegran, Andrew Bender, and Douglas Hazard - Anthrobot
 + many other undergraduate students working in our lab over the years

Technical support:
Rakela Colon, Jayati Mandal - lab management
Erin Switzer - vertebrate animal husbandry
Emma Lederer - Xenobot behavior
Joan Lemire - molecular biology

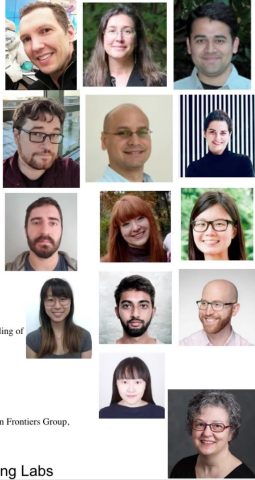
Collaborators:
Allen Center members:
Dany Adams - bioelectric face prepattern
Alexis Pietak, Marcel Blatner, Salvador Mafe, Javier Cervera - computational modeling of
Joshua Bengard - Xenobot simulations and AI
Giovanni Pezzulo - cognitive science models of pattern regulation
Sebastian Risi - open-ended evolution
Simon Garnier - computational analysis of Anthrobot form and function
Chris Fields - physics of sentience and sentience of physics

Model systems: tadpoles, planaria, zebrafish, slime molds, human cells, and chick embryos

Funding support: TWCF, JTF, Goy Foundation, Emerald Gate, AFOSR, ARO, DARPA, Paul G. Allen Frontiers Group, NIH, NSF

Illustrations: Jeremy Guay @ Peregrine Creative

Disclosures: MorphoCeuticals, Fauna Systems, Astonishing Labs



I'll stop here and thank all the people who did the work. The postdocs and PhD students and our many, many amazing collaborators. I have to do a disclosure. There are three companies that have spun out of our work that provide some funding for some of this research and various other funders over the years to whom I am extremely grateful. Lots of thanks go to the actual model systems which do all the hard work as we try to figure this stuff out. Thank you very much.

Thank you for reading.

More lectures

You can find more of my lectures [here](#).

Follow my work

[Twitter](#) • [Blog](#) • [The Levin Lab](#)

Want one for your lecture?

Want something like this for your own talk? Reach out to Adi at adi@aipodcast.ing.