

EXCERPTED FROM

STEPHEN
WOLFRAM
A NEW
KIND OF
SCIENCE

SECTION 12.10

*Intelligence in the
Universe*

Intelligence in the Universe

Whether or not we as humans are the only examples of intelligence in the universe is one of the great unanswered questions of science.

Just how intelligence should be defined has never been quite clear. But in recent times it has usually been assumed that it has something to do with an ability to perform sophisticated computations.

And with traditional intuition it has always seemed perfectly reasonable that it should take a system as complicated as a human to exhibit such capabilities—and that the whole elaborate history of life on Earth should have been needed to generate such a system.

With the development of computer technology it became clear that many features of intelligence could be achieved in systems that are not biological. Yet our experience has still been that to build a computer requires sophisticated engineering that in a sense exists only because of human biological and cultural development.

But one of the central discoveries of this book is that in fact nothing so elaborate is needed to get sophisticated computation. And indeed the Principle of Computational Equivalence implies that a vast range of systems—even ones with very simple underlying rules—should be equivalent in the sophistication of the computations they perform.

So in as much as intelligence is associated with the ability to do sophisticated computations it should in no way require billions of years of biological evolution to produce—and indeed we should see it all over the place, in all sorts of systems, whether biological or otherwise.

And certainly some everyday turns of phrase might suggest that we do. For when we say that the weather has a mind of its own we are in effect attributing something like intelligence to the motion of a fluid. Yet surely, one might argue, there must be something fundamentally more to true intelligence of the kind that we as humans have.

So what then might this be?

Certainly one can identify all sorts of specific features of human intelligence: the ability to understand language, to do mathematics, solve puzzles, and so on. But the question is whether there are more

general features that somehow capture the essence of true intelligence, independent of the particular details of human intelligence.

Perhaps it could be the ability to learn and remember. Or the ability to adapt to a wide range of different and complex situations. Or the ability to handle abstract general representations of data.

At first, all of these might seem like reasonable indicators of true intelligence. But as soon as one tries to think about them independent of the particular example of human intelligence, it becomes much less clear. And indeed, from the discoveries in this book I am now quite certain that any of them can actually be achieved in systems that we would normally never think of as showing anything like intelligence.

Learning and memory, for example, can effectively occur in any system that has structures that form in response to input, and that can persist for a long time and affect the behavior of the system. And this can happen even in simple cellular automata—or, say, in a physical system like a fluid that carves out a long-term pattern in a solid surface.

Adaptation to all sorts of complex situations also occurs in a great many systems. It is well recognized when natural selection is present. But at some level it can also be thought of as occurring whenever a constraint ends up getting satisfied—even say that a fluid flowing around a complex object minimizes the energy it dissipates.

Handling abstraction is also in a sense rather common. Indeed, as soon as one thinks of a system as performing computations one can immediately view features of those computations as being like abstract representations of input to the system.

So given all of this is there any way to define a general notion of true intelligence? My guess is that ultimately there is not, and that in fact any workable definition of what we normally think of as intelligence will end up having to be tied to all sorts of seemingly rather specific details of human intelligence.

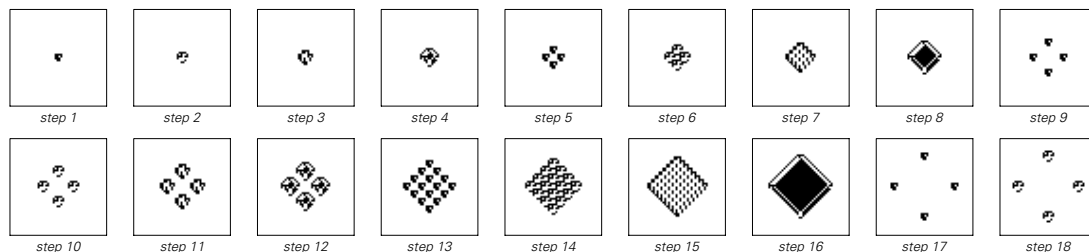
And as it turns out this is quite similar to what happens if one tries to define the seemingly much simpler notion of life.

There was a time when it was thought that practically any system that moves spontaneously and responds to stimuli must be

alive. But with the development of machines having even the most primitive sensors it became clear that this was not correct.

Work in the field of thermodynamics led to the idea that perhaps living systems could be defined by their ability to take disorganized material and spontaneously organize it—usually to incorporate it into their own structure. Yet all sorts of non-living systems—from crystals to flames—also do this. And Chapter 6 showed that self-organization is actually extremely common even among systems with simple rules.

For a while it was thought that perhaps life might be defined by its ability for self-reproduction. But in the 1950s abstract computational systems were constructed that also had this ability. Yet it seemed that they needed highly complex rules—not unlike those found in actual living cells. But in fact no such complexity is really necessary. And as one might now expect from the intuition in this book, even systems like the one below with remarkably simple rules can still manage to show self-reproduction—despite the fact that they bear almost no other resemblance to ordinary living systems.



A two-dimensional cellular automaton that exhibits an almost trivial form of self-reproduction, in which multiple copies of any initial pattern appear every time the number of steps of evolution doubles. The rule used is additive, and takes a cell to be black whenever an odd number of its neighbors were black on the step before (outer totalistic code 204). The same basic self-reproduction phenomenon occurs in elementary rule 90, as well as in essentially any other additive rule, in any number of dimensions.

If one looks at typical living systems one of their most obvious features is great apparent complexity. And for a long time it has been thought that such complexity must somehow be unique to living systems—perhaps requiring billions of years of biological evolution to develop. But what I have shown in this book is that this is not the case, and that in fact a vast range of systems—including ones with very

simple underlying rules—can generate at least as much complexity as we see in the components of typical living systems.

Yet despite all this, we do not in our everyday experience typically have much difficulty telling living systems from non-living ones. But the reason for this is that all living systems on Earth share an immense number of detailed structural and chemical features—reflecting their long common history of biological evolution.

So what about extraterrestrial life? To be able to recognize this we would need some kind of general definition of life, independent of the details of life on Earth. But just as in the case of intelligence, I believe that no reasonable definition of this kind can actually be given.

Indeed, following the discoveries in this book I have come to the conclusion that almost any general feature that one might think of as characterizing life will actually occur even in many systems with very simple rules. And I have little doubt that all sorts of such systems can be identified both terrestrially and extraterrestrially—and certainly require nothing like the elaborate history of life on Earth to produce.

But most likely we would not consider these systems even close to being real examples of life. And in fact I expect that in the end the only way we would unquestionably view a system as being an example of life is if we found that it shared many specific details with life on Earth—probably down, say, to being made of gelatinous materials and having components analogous to proteins, enzymes, cell membranes and so on—and perhaps even down to being based on specific chemical substances like water, sugars, ATP and DNA.

So what then of extraterrestrial intelligence? To what extent would it have to show the same details as human intelligence—and perhaps even the same kinds of knowledge—for us to recognize it as a valid example of intelligence?

Already just among humans it can in practice be somewhat difficult to recognize intelligence in the absence of shared education and culture. Indeed, in young children it remains almost completely unclear at what stage different aspects of intelligence become active.

And when it comes to other animals things become even more difficult. If one specifically tries to train an animal to solve

mathematical puzzles or to communicate using human language then it is usually possible to recognize what intelligence it shows.

But if one just observes the normal activities of the animal it can be remarkably difficult to tell whether they involve intelligence. And so as a typical example it remains quite unclear whether there is intelligence associated with the songs of either birds or whales.

To us these songs may sound quite musical—and indeed they even seem to show some of the same principles of organization as human music. But do they really require intelligence to generate?

Particularly for birds it has increasingly been possible to trace the detailed processes by which songs are produced. And it seems that at least some of their elaborate elements are just direct consequences of the complex patterns of air flow that occur in the vocal tracts of birds.

But there is definitely also input from the brain of the bird. Yet within the brain some of the neural pathways responsible are known. And one might think that if all such pathways could be found then this would immediately show that no intelligence was involved.

Certainly if the pathways could somehow be seen to support only simple computations then this would be a reasonable conclusion. But just using definite pathways—or definite underlying rules—does not in any way preclude intelligence. And in fact if one looks inside a human brain—say in the process of generating speech—one will no doubt also see definite pathways and definite rules in use.

So how then can we judge whether something like a bird song, or a whale song—or, for that matter, an extraterrestrial signal—is a reflection of intelligence? The fundamental criterion we tend to use is whether it has a meaning—or whether it communicates anything.

Everyday experience shows us that it can often be very hard to tell. For even if we just hear a human language that we do not know it can be almost impossible for us to recognize whether what is being said is meaningful or not. And the same is true if we pick up data of any kind that is encoded in a format we do not know.

We might start by trying to use our powers of perception and analysis to find regularities in the data. And if we found too many regularities we might conclude that the data could not represent

enough information to communicate anything significant—and indeed perhaps this is the case for at least some highly repetitive bird songs.

But what if we could find no particular regularities? Our everyday experience with human language might make us think that the data could then have no meaning. But there is nothing to say that it might not be a perfectly meaningful message—even one in human language—that just happens to have been encrypted or compressed to a point where it shows no detectable regularities.

And indeed it is sobering to notice that if one just listens even to bird songs and whale songs there is little that fundamentally seems to distinguish them from what can be generated by all sorts of processes in nature—say the motion of chimes blowing in the wind or of plasma in the Earth's magnetosphere.

One might imagine that one could find out whether a meaningful message had been communicated in a particular case by looking for correlations it induces between the actions of sender and receiver. But it is extremely common in all sorts of natural systems to see effects that propagate from one element to another. And when it comes even to whale songs it turns out that no clear correlations have ever in the end been identified between senders and receivers.

But what if one were to notice some event happen to the sender? If one were somehow to see a representation of this in what the sender produced, would it not be evidence for meaningful communication?

Once again, it need not be. For there are a great many cases in which systems generate signals that reflect what happens to them. And so, for example, a drum that is struck in a particular pattern will produce a sound that reflects—and in effect represents—that pattern.

Yet on the other hand even among humans different training or culture can lead to vastly different responses to a given event. And for animals there is the added problem of emphasis on different forms of perception. For presumably dogs can sense the detailed pattern of smell in their environment, and dolphins the detailed pattern of fluid motion around them. Yet we as humans would almost certainly not recognize descriptions presented in such terms.

So if we cannot identify intelligence by looking for meaningful communication, can we perhaps at least tell for a given object whether intelligence has been involved in producing it?

For certainly our everyday experience is that it is usually quite easy to tell whether something is an artifact created by humans.

But a large part of the reason for this is just that most artifacts we encounter in practice have specific elements that look rather similar. Yet presumably this is for the most part just a reflection of the historical development of engineering—and of the fact that the same basic geometrical and other forms have ended up being used over and over again.

So are there then more general ways to recognize artifacts?

A fairly good way in practice to guess whether something is an artifact is just to look and see whether it appears simple. For although there are exceptions—like crystals, bubbles and animal horns—the majority of objects that exist in nature have irregular and often very intricate forms that seem much more complex than typical artifacts.

And indeed this fact has often been taken to show that objects in nature must have been created by a deity whose capabilities go beyond human intelligence. For traditional intuition suggests that if one sees more complexity it must always in a sense have more complex origins.

But one of the main discoveries of this book is that in fact great complexity can arise even in systems with extremely simple underlying rules, so that in the end nothing with rules even as elaborate as human intelligence—let alone beyond it—is needed to explain the kind of complexity we see in nature.

But the question then remains why when human intelligence is involved it tends to create artifacts that look much simpler than objects that just appear in nature. And I believe the basic answer to this has to do with the fact that when we as humans set up artifacts we usually need to be able to foresee what they will do—for otherwise we have no way to tell whether they will achieve the purposes we want.

Yet nature presumably operates under no such constraint. And in fact I have argued that among systems that appear in nature a great many exhibit computational irreducibility—so that in a sense it becomes irreducibly difficult to foresee what they will do.

Yet at least with its traditional methodology engineering tends to rely on computational reducibility. For typically it operates by building systems up in such a way that the behavior of each element can always readily be predicted by something like a simple mathematical formula.

And the result of this is that most systems created by engineering are forced in some sense to seem simple—in mechanical cases for example typically being based only on simple repetitive motion.

But is simplicity a necessary feature of artifacts? Or might artifacts created by extraterrestrial intelligence—or by future human technology—seem to show no signs of simplicity?

As soon as we say that a system achieves a definite purpose this means that we can summarize at least some part of what the system does just by describing this purpose. So if we have a simple description of the purpose it follows that we must be able to give a simple summary of at least some part of what the system does.

But does this then mean that the whole behavior of the system must be simple? Traditional engineering might tend to make one think so. For typically our experience is that if we are able to get a particular kind of system to generate a particular outcome at all, then normally the behavior involved in doing so is quite simple.

But one of the results of this book is that in general things need not work like this. And so for example at the end of Chapter 5 we saw several systems in which a simple constraint of achieving a particular outcome could in effect only be satisfied with fairly complex behavior.

And as I will discuss in the next section I believe that in the effort to optimize things it is almost inevitable that even to achieve comparatively simple purposes more advanced forms of technology will make use of systems that have more and more complex behavior.

So this means that there is in the end no reason to think that artifacts with simple purposes will necessarily look simple.

And so if we are just presented with something, how then can we tell if it has a purpose? Even with things that we know were created by humans it can already be difficult. And so, for example, there are many archeological structures—such as Stonehenge—where it is at best unclear which features were intended to be purposeful.

And even in present-day situations, if we are exposed to objects or activities outside the areas of human endeavor with which we happen to be familiar, it can be very hard for us to tell which features are immediately purposeful, and which are unintentional—or have, say, primarily ornamental or ceremonial functions.

Indeed, even if we are told a purpose we will often not recognize it. And the only way we will normally become convinced of its validity is by understanding how some whole chain of consequences can lead to purposes that happen to fit into our own specific personal context.

So given this how then can we ever expect in general to recognize the presence of purpose—say as a sign of extraterrestrial intelligence?

And as an example if we were to see a cellular automaton how would we be able to tell whether it was created for a purpose?

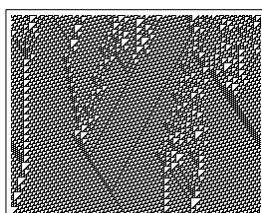
Of the cellular automata in this book—especially in Chapter 11—a few were specifically constructed to achieve particular purposes. But the vast majority originally just arose as part of my investigation of what happens with the simplest possible underlying rules.

And at first I did not think of most of them as achieving any particular purposes at all. But gradually as I built up the whole context of the science in this book I realized that many of them could in fact be thought of as achieving very definite purposes.

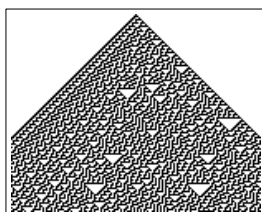
Systems like rule 110 shown on the left have a kind of local coherence in their behavior that reminds one of the operation of traditional engineering systems—or of purposeful human activity. But the same is not true of systems like rule 30. For although one can see that such systems have a lot going on, one tends to assume that somehow none of it is coherent enough to achieve any definite purpose.

Yet in the context of the ideas in this book, a system like rule 30 can be viewed as achieving the purpose of performing a fairly sophisticated computation. And indeed we know that this computation is useful in practice for generating sequences that appear random.

But of course it is not necessary for us to talk about purpose when we describe the behavior of rule 30. We can perfectly well instead talk only about mechanism, and about the way in which the underlying rules for the cellular automaton lead to the behavior we see.



rule 110



rule 30

Cellular automata whose behavior does and does not give the impression of being purposeful.

And indeed this is true of any system. But as a practical matter we often end up describing what systems do in terms of purpose when this seems to us simpler than describing it in terms of mechanism.

And so for example if we can identify some simple constraint that a system always tries to satisfy it is not uncommon for us to talk of this as being the purpose of the system. And in fact we do this even in cases like minimization of energy in physical systems or natural selection for fitness in biological systems where nothing that we ordinarily think of as intelligence is involved.

So the fact that we may be able to interpret a system as achieving some purpose does not necessarily mean that the system was really created with that purpose in mind. And indeed just looking at the system we will never ultimately be able to tell for sure that it was.

But we can still often manage to guess. And given a particular supposed purpose one potential criterion to use is that the system in a sense not appear to do too much that is extraneous to that purpose.

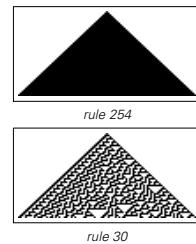
And so, for example, in looking at the pictures on the right it would normally seem much more plausible that rule 254 might have been set up for the purpose of generating a uniformly expanding pattern than that rule 30 might have been. For while rule 30 does generate such a pattern, it also does a lot else that appears irrelevant to this purpose.

So what this might suggest is that perhaps one could tell that a system was set up for a given purpose if the system turns out to be in a sense the minimal one that achieves that purpose.

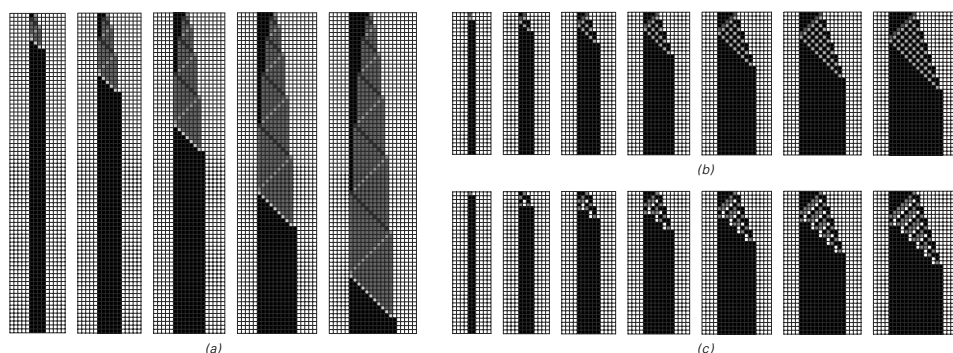
But an immediate issue is that in traditional engineering we normally do not come even close to getting systems that are minimal. Yet it seems reasonable to suppose that as technology becomes more advanced it should become more common that the systems it sets up for a given purpose are ones that are minimal.

So as an example of all this consider cellular automata that achieve the purpose of doubling the width of the pattern given in their input. Case (a) in the picture on the next page is a cellular automaton one might construct for this purpose by using ideas from traditional engineering.

But while this cellular automaton seems to have little extraneous going on, it operates in a slow and sequential way, and its underlying



If the purpose is to generate a uniformly expanding pattern it seems more plausible that the top cellular automaton should have been the one created for this purpose.



Examples of cellular automata that can be viewed as achieving the purpose of doubling the width of the pattern given in their input. Rule (a) involves 6 colors, and works sequentially, much as a typical traditional engineering system might. Rule (b) involves 4 colors, and works in parallel. Rule (c) was found by a large search, and involves only 3 colors. It takes the fewest steps of any 3-color rule to generate its result. Its rule number is 5407067979.

rules turn out to be far from minimal. For case (b) gets its results much more quickly—in effect by operating in parallel—and its rules involve four possible colors rather than six.

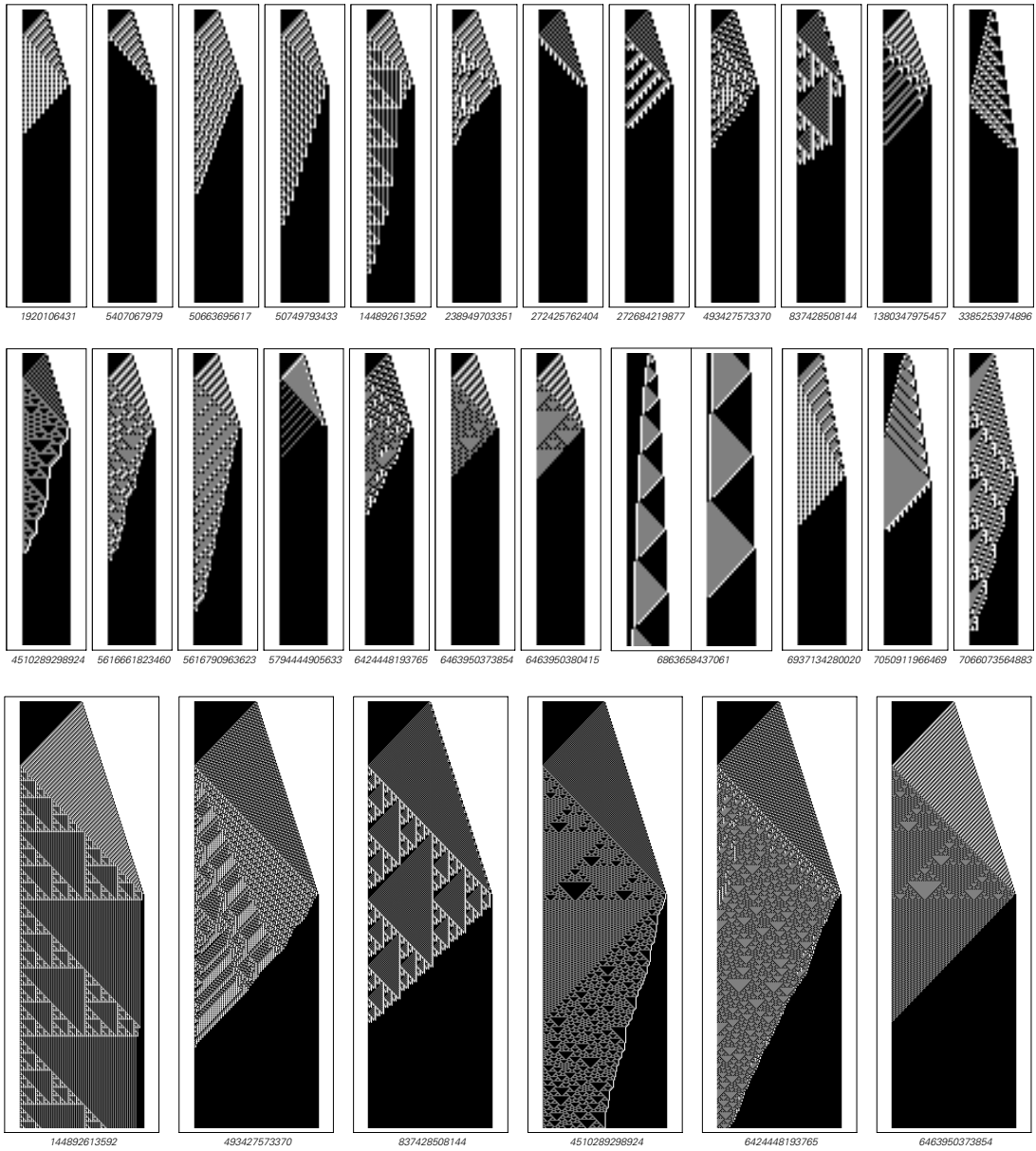
But is case (b) really the minimal cellular automaton that achieves the purpose of doubling its input? Just thinking about it, one might not be able to come up with anything better. But if one in effect explicitly searches all 8 trillion or so rules that involve less than four colors, it turns out that one can find 4277 three-color rules that work.

The pictures on the facing page show a few typical examples.

Each uses at least a slightly different scheme, but all achieve the same purpose of doubling their input. Yet often they operate in ways that seem considerably more complex than most familiar artifacts. And indeed some of the examples might look to us more like systems that just occur in nature than like artifacts.

But the point is that with sufficiently advanced technology one might expect that doubling of input would be implemented using a rule that is in some sense optimal. Different criteria for optimality could lead to different rules, but usually they will be rules like those on the facing page—and sometimes rules with quite complex behavior.

But now the question is if one were just to encounter such a rule, would one be able to guess that it was created for a purpose? After all,



Examples of rules with three colors that achieve the purpose of doubling the width of the pattern given in their input. These examples are taken from the 4277 found in effect by searching exhaustively all 7,625,597,484,987 possible rules with three colors. In most cases the number of steps to generate the final pattern increases roughly linearly with the width of the input—although in the case of the fourth-to-last rule on the second row it is $2(n^2 - n + 1)$ for width n .

there are all sorts of features in the behavior of these rules that could in principle represent a possible purpose. But what is special about rules like those on the previous page is that they are the minimal ones that exhibit the particular feature of doubling their input.

And in general if one sees some feature in the behavior of a system then finding out that the rule for the system is the minimal or optimal one for producing that feature may make it seem more likely that at least with sufficiently advanced technology the system might have specifically been created for the purpose of exhibiting that feature.

Computational irreducibility implies that it can be arbitrarily difficult to find minimal or optimal rules. Yet given any procedure for trying to do this it is certainly always possible that the procedure could just occur in nature without any purpose or intelligence being involved.

And in fact one might consider this not all that unlikely for the kind of fairly straightforward exhaustive searches that I ended up using to find the cellular automaton rules in the pictures on the previous page.

So what does all this mean for extraterrestrial intelligence?

Extrapolating from our own development we might expect that given sufficiently advanced technology it would be almost inevitable for artifacts to be constructed on an astronomical scale—perhaps for example giant machines with objects like stars as components.

Yet we do not believe that we have ever seen any such artifacts.

But how do we know for sure? For certainly our astronomical observations have revealed all sorts of phenomena for which we do not yet have any very satisfactory explanations. And indeed until just a few centuries ago most such unexplained phenomena would routinely have been attributed to some kind of divine intelligence.

But in more recent times it has become almost universally assumed that they must instead be the result of physical processes in which nothing like intelligence is involved.

Yet what the discoveries in this book have shown is that even such physical processes can often correspond to computations that are at least as sophisticated as any that we as humans perform.

But what we believe is that somehow none of the phenomena we see have any sense of purpose analogous to typical human artifacts.

Occasionally we do see evidence of simple geometrical shapes like those familiar from human artifacts—or visible on the Earth from space. But normally our explanations for these end up being short enough that they seem to leave no room for anything like intelligence. And when we see elaborate patterns, say in nebulae or galaxies, we assume that these can have no purpose—even though they may remind us to some extent of human art.

So if we do not recognize any objects that seem to be artifacts, what about signals that might correspond to messages?

If we looked at the Earth from far away the most obvious signs of human intelligence would probably be found in radio signals.

And in fact in the past it was often assumed that just to generate radio signals at all must require intelligence and technology. So when complex radio signals not of human origin were discovered in the early 1900s it was at first thought that they must be coming from extraterrestrial intelligence. But it was eventually realized that in fact the signals were just produced by effects in the Earth's magnetosphere.

And then again in the 1960s when the intense and highly regular signals of pulsars were discovered it was briefly thought that they too must come from extraterrestrial intelligence. But it was soon realized that these signals could actually be produced just by ordinary physical processes in the magnetospheres of rapidly rotating neutron stars.

So what might a real signal from extraterrestrial intelligence be like? Human radio signals currently tend to be characterized by the presence of sharply defined carrier frequencies, corresponding in effect to almost perfect small-scale repetition. But such regularity greatly reduces the rate at which information can be transmitted. And as technology advances less and less regularity needs to be present.

But in practice essentially all serious searches for extraterrestrial intelligence made so far have been based on using radio telescopes to look for signals with sharply defined frequencies. And indeed no such signals have been found. But as we saw in Chapter 10 even signals that are nested rather than purely repetitive cannot reliably be recognized just by looking for peaks in frequency spectra.

And there is certainly in general no lack of radio signals that we receive from around our galaxy and beyond. But the point is that these signals typically seem to us quite random. And normally this has made us assume that they must in effect just be some kind of radio noise that is being produced by one of several simple physical processes.

But could it be that some of these signals instead come from extraterrestrial intelligence—and are in fact meaningful messages?

Ongoing communications between extraterrestrials seem likely to be localized to regions of space where they are needed, and therefore presumably not accessible to us. And even if some signals involved in such communications are broadcast, my guess is that they will exhibit essentially no detectable regularities. For any such regularity represents in a sense a redundancy or inefficiency that can be removed by the sender and receiver both using appropriate data compression.

But if there are beacons that are intended to be noticed even if one does not already know that they are there, then the signals these produce must necessarily have recognizable distinguishing features, and thus regularities that can be detected, at least by their potential users.

So perhaps the problem is just that the methods of perception and analysis that we as humans have are not powerful enough. And perhaps if we could only find the appropriate new method it would suddenly be clear that some of what we thought was random radio noise is actually the output of beacons set up by extraterrestrial intelligence.

For as we saw in Chapter 10 most of the methods of perception and analysis that we currently use can in general do little more than recognize repetition—and sometimes nesting. Yet in the course of this book we have seen a great many examples where data that appears to us quite random can in fact be produced by very simple underlying rules.

And although I somewhat doubt it, one could certainly imagine that if one were to show data like the center column of rule 30 or the digit sequence of π to an extraterrestrial then they would immediately be able to deduce simple rules that can produce these.

But even if at some point we were to find that some of the seemingly random radio noise that we detect can be generated by simple rules, what would this mean about extraterrestrial intelligence?

In many respects, the simpler the rules, the more likely it might seem that they could be associated with ordinary physical processes, without anything like intelligence being involved.

Yet as we discussed above, if one could actually determine that the rules used in a given case were the simplest possible, then this might suggest that they were somehow set up on purpose. But in practice if one just receives a signal one normally has no way to tell which of all possible rules for producing it were in fact used.

So is there then any kind of signal that could be sent that would unambiguously communicate the presence of intelligence?

In the past, one might have thought that it would be enough for the production of the signal to involve sophisticated computation. But the discoveries in this book have made it clear that in fact such computation is quite common in all sorts of systems that do not show anything that we would normally consider intelligence.

And indeed it seems likely that for example an ordinary physical process like fluid turbulence in the gas around a star should rather quickly do more computation than has by most measures ever been done throughout the whole course of human intellectual history.

In discussions of extraterrestrial intelligence it is often claimed that mathematical constructs—such as the sequence of primes—somehow serve as universal signs of intelligence.

But from the results in this book it is clear that this is not correct.

For while in the past it might have seemed that the only way to generate primes was by using intelligence, we now know that the rather straightforward computations required can actually be carried out by a vast range of different systems—with no apparent need for intelligence.

One might nevertheless imagine that any sufficiently advanced intelligence would somehow at least consider the primes significant.

But here again I do not believe that this is correct. For very little even of current human technology depends on ideas about primes. And I am also fairly sure that not much can be deduced from the fact that primes happen to be popular in present-day human mathematics.

For despite its reputation for generality I argued at length in the previous section that the whole field of mathematics that we as

humans have historically developed ultimately covers only a tiny fraction of what is possible—notably leaving out the vast majority of systems that I have studied in this book.

And if one identifies a feature—such as repetition or nesting—that is common to many possible systems, then it becomes inevitable that this feature will appear not only when intelligence or mathematics is involved, but also in all sorts of systems that just occur in nature.

So what about trying to set up a signal that gives evidence of somehow having been created for a purpose? I argued above that if the rules for a system are as simple as they can be, then this may suggest the presence of purpose. But such a criterion relies on seeing not only a signal but also the mechanism by which the signal was produced.

So what about a signal on its own? One might imagine that one could set something up—say the solution to a difficult mathematical problem—that was somehow easy to describe in terms of a constraint or purpose, but difficult to explain in terms of an explicit mechanism.

But in a sense such a thing cannot exist. For given a constraint, it is always in principle simple to set up an exhaustive search that provides a mechanism for finding what satisfies the constraint.

However, this may still take a lot of computational effort. But we cannot use that alone as a criterion. For as we have seen, many systems that just occur in nature actually end up doing more computation than typical systems that we explicitly set up for a purpose.

So even if we cannot find an abstract way to give evidence of purpose or intelligence, what about using the practical fact that both the sender and receiver of a signal exist in the same physical universe? Can one perhaps use a signal that is a representation of actual data in, say, astronomy, physics or chemistry?

As I discussed earlier, the more direct the representation the more easily an ordinary physical process can be expected to generate it, and the less there will be any indication of intelligence—just as, for example, something like a photograph can be produced essentially just by projecting light, while a diagram or a painting requires more.

But as soon as there is interpretation of data, it can become very difficult to recognize the results. For different forms of perception and

different experiences and contexts can cause vastly different features to be emphasized. And thus, for example, the fact that we can readily recognize pictures of animals in cave paintings made by Stone Age humans depends greatly on the fact that our visual system still picks out the same specific features.

But what about more abstract art?

Although one has the feeling that this involves more human input, it rapidly becomes extremely difficult to tell what has been created on purpose. And so, for example, if one sees a splash of paint it is almost impossible to know without detailed cultural background and context whether it is intended to be purposeful art.

So what does all this mean about extraterrestrial intelligence?

My main conclusion is rather similar to my conclusion about artificial intelligence in Chapter 10: that the basic issue is not finding systems that perform sophisticated enough computations, but rather finding ones whose details happen to be similar enough to us as humans that we recognize what they do as showing intelligence.

And there is perhaps some analogy to recognizing the capability for sophisticated computation in the first place. For while this is undoubtedly very common say in cellular automata, the most immediate suggestions of it are in class 4 systems like rule 110 that in effect happen to do their computations in a way that looks at least somewhat similar to the way we as humans are used to doing them.

So should we expect that somehow recognizable extraterrestrial intelligence will occur at a level of a few percent—like class 4 systems?

There is clearly more to the phenomenon of intelligence than this. But if we require something that follows too many of the details of us as humans there is already evidence that it does not exist. For if such intelligence had ever arisen in the past, then extrapolating from our own history we would expect that some of it would long ago have colonized our galaxy—at least with signals, if not with physical objects.

But I suspect that if we generalize even quite modestly our definition of intelligence then there will be examples that can be found—at least with sufficiently powerful methods of perception and analysis. Yet it seems likely that they will behave in some ways that are

as bizarrely different from human intelligence as many of the simple programs in this book are different from the systems that have traditionally been studied in human mathematics and science.

Implications for Technology

My main purpose in this book has been to build a new kind of basic science. But I expect that in time what I have done will also have many implications for technology. No doubt there will be all sorts of specific applications of particular results and ideas. But in the long run probably the most important consequence will be to introduce a vast new range of systems and processes that can be used for technology.

And indeed one of the things that emerges from this book is that traditional engineering has actually considered only a tiny and quite unrepresentative fraction of all the kinds of systems and processes that are in principle possible.

Presumably the reason—as I have mentioned several times in this book—is that its whole methodology has tended to be based on setting up systems whose behavior is simple enough that almost every aspect of them can always readily be predicted. But doing this has immediately excluded many of the systems that I have studied in this book—or for that matter that occur in nature. And no doubt this is why systems created by engineering have in the past usually ended up looking so much simpler than typical systems in nature.

And with traditional intuition it has normally been assumed that the only way to create systems that show a higher degree of complexity is somehow to build this complexity into their underlying rules.

But one of the central discoveries of this book is that this is not the case, and that in fact it is perfectly possible for systems even with extremely simple underlying rules to produce behavior that has immense complexity—and that looks like what one sees in nature.

And I believe that if one uses such systems it is almost inevitable that a vast amount of new technology will become possible.

There are some places where just the abstract ability to produce complexity from simple rules is already important. One example discussed in Chapter 10 is cryptography. Other examples include all