



Recursion Schemes for Mathematicians

An introduction to induction in programming

10 September 2020

Recursion schemes are a neat construction that allows one to better structure and reason about recursive functions. Instead of viewing recursion as functions that call themselves, these schemes allow us to take a higher-order point of view. Most of recursive definitions may be abstracted as some operator that takes a function which knows how to deal with individual parts of a structure and turns it into a function defined for the entire structure. The canonical example for this perspective is turning the binary function $+$, which adds two numbers, into a new function that sums all elements of a list with arbitrary length¹. Another use case is, when writing an evaluator for an expression tree, to only tell your program how to deal with each node and let a recursion scheme turn it into a full evaluator.

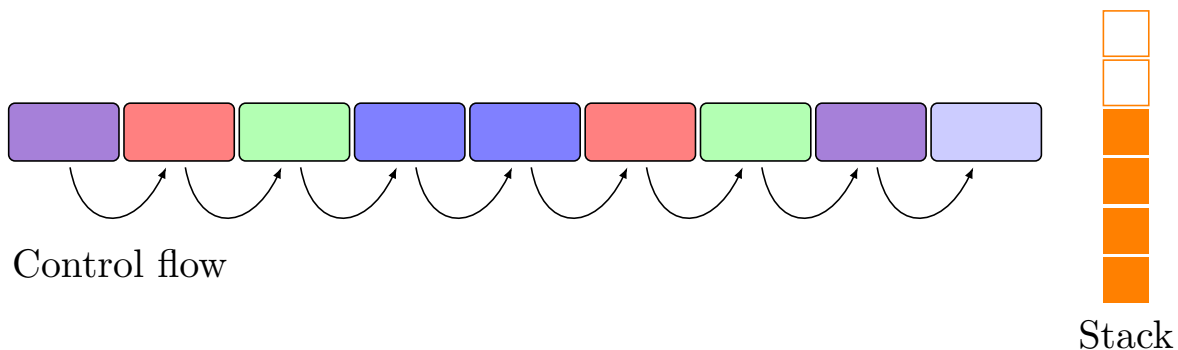
In this post, we will take a look at the three simplest varieties of recursion schemes: catamorphisms, anamorphisms, and hylomorphisms; all of them deeply linked with structural induction. As we will see, they respectively encapsulate the notions of folding a data structure, constructing a data structure and using a data structure as an intermediate step.

The first time I heard about recursion schemes was after stumbling with the paper [Functional Programming with Bananas, Lenses, Envelopes and Barbed Wire](#) by Erik Meijer, Maarten Fokkinga, and Ross Paterson. It is a real gem of functional wisdom but the authors use a notation called [Squiggol](#), which I had a really hard time trying to grasp. Fortunately, I later found [Patrick Thomson's excellent series of posts](#), which explain recursion schemes using Haskell code. After absorbing this content, I tried to explain the idea to some friends who don't know Haskell but know some category theory. This post is an enriched version of that presentation with more cohesion and much less of my bad drawings.

In the spirit of the original presentation, I've decided to make this post completely language-agnostic. I also tried to keep any type theory or lambda calculus to a minimum in order to make it easier to grasp. My hope is that anyone with some experience in programming and mathematics will be able to read this, instead of confining this post only to people that know a certain programming language. Of course, in doing that I'm risking that no one besides me will actually understand this post. But, very well, I will do my best.

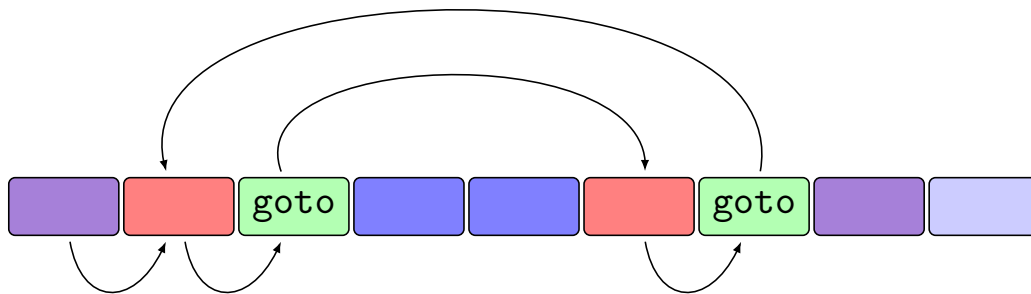
A not so historical introduction

We begin our journey in the prehistory of programming; that time when computers were still mostly made of coconuts and mammoth fur. At the time, the way to program was quite unstructured. The code of a program was essentially given by a sequence of commands whereas the data was simply piled up in a stack.



But the program does not have to execute the instructions in order. Oh no, it would be rather boring and inexpressive... In this programming paradigm, there is a special command called `goto` which allows one to jump to any labeled

part of the program. In fact, this is the only command responsible for the program's control flow.



The `goto` is a pretty powerful construction, but it also has its caveats. So much for its use to be *considered harmful*. Although any kind of program flow can be constructed using only `goto`, it may become too confuse. The `goto` statement makes it too easy to write spaghetti code, which is practically impossible to debug.

After prehistory, we arrive into the societal period of imperative programming. In here, the programming languages become *structured*. Data is no long viewed as simply a stack of memory but classified into types. There are some primitive types as well as structures to combine them into more complex types. In a language like C, for example, there are `struct` types, `union` types, and array types.

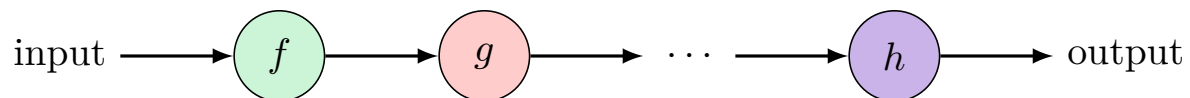
The changes also arrived to the control flow and an effort was made to tame the once wild `goto`. The programmers from that time analysed its most common use cases and created new statements fulfilling each of these. You are probably acquainted to them as `for` loops, `while` loops, `switch/case` statements, function and subroutine declarations and `if/else` statements; just to name a few.

Both in regards to data and control flow, the movement we encounter in here consists of substituting general purpose structures that concretely represent the instructions we give to the processor by specific structures having more abstract roles in the code. In terms of computational expressiveness, nothing changes. What the processor sees is the same in both unstructured and structured programming. The benefit lies in the programmer's expressiveness. With more specific statements, it becomes easier to write larger, more complex programs as well as properly debug them. As an example, we may notice that it is possible to guarantee that a `for` loop always terminates if it is iterating over a

block of code that is known to terminate. No strange thing can happen. On the other side, the general character of the `goto` allows us to simulate the behavior of a `for` but there is no guarantee that a program with a `goto` statement must terminate.

Until now, we were looking at a programming paradigm called *imperative programming*, but the world of programming languages is not so uniform. Other paradigms exist. And the one that mostly fits what we will see today is called *functional programming*.

While the imperative paradigm views the code as a sequence of commands the programmer gives the computer to execute (hence the name, *imperative*), the functional paradigm views the code as a composition of functions in the mathematical sense. A function takes a value as input, does some processing to it, calling other functions for this, and returns another value as output.



If every program only consisted of applying a finite amount of previously defined functions, the language's expressiveness would be rather limited. To overcome this, we need some form of control flow, which is achieved via recursion.

A *recursive function* is a function that, in order to process its input into the output, may call itself in an intermediate step. Probably the most famous recursive function is the factorial, defined as

$$\text{fat}(n) = \begin{cases} 1, & n = 0 \\ n \cdot \text{fat}(n - 1), & \text{otherwise.} \end{cases}$$

The expressiveness gained from recursion is essentially the same as the one gained from `goto` in imperative languages. That is, the control flow given by function composition together with recursion allows one to do anything imaginable with the programming language. However, all this expressiveness comes with its caveats. It is too easy to write a functional spaghetti code if recursion is used indiscriminately. Because of all its power, recursive code lacks in safety. It would even be fair to say that, like `goto`, it is too unstructured.

The idea of structured control flow really caught on in the imperative world. One hardly sees a wild `goto` in the middle of a modern piece of code. In fact, many languages don't even allow it. In the functional world, on the other side,

the trendy for taming recursion never really caught on. Despite many functional languages organizing their data using types, control is still done using crude recursion.

Recursion schemes are ways to organize and structure different kinds of recursive functions. Instead of writing a recursive function in terms of itself, we define higher order functions that receive an ordinary function as argument, do some recursive magic on it, and return a recursive analogue of that function as output. It is similar to how a `while` loop takes a boolean statement and a block of code, and turns them into a repeating block of code. If it seems too confusing, just keep on. What I mean in here will become clearer after we [construct catamorphisms](#).

Before we end this motivation and proceed to the actual construction of recursion schemes, there is an intuition that I believe useful to have in mind. In imperative programming, there is a close relationship between how we structure data and how we structure the control flow. Working with arrays almost asks the programmer to design programs with `for` loops. Similarly, it is natural to deal with `union` types using `switch` statements (also called `case` or `cond` in some languages). Recursion schemes will arise as an analogous to this idea in the functional programming setting. So far so good for motivation, let's dive into some math.

Algebraic Data Types

The simplest form to write a type is by enumerating its elements such as

$$\mathbf{Bool} := \mathbf{true} \mid \mathbf{false}.$$

This way, we define the type of boolean values. That is, a type with exactly two terms called `true` and `false`. In general any finite type can be written just by enumerating its elements. As another example, there is the type of musical notes

$$\mathbf{Notes} := \mathbf{do} \mid \mathbf{re} \mid \mathbf{mi} \mid \mathbf{fa} \mid \mathbf{sol} \mid \mathbf{la} \mid \mathbf{si}.$$

Nevertheless, in practice we also want to deal with types that are more complex than finite collections. The solution to this is assuming that the programming language comes with some built-in types such as integers, characters, and

floating-point numbers, together with structures that allow us to compose these types.

One common compound type consists of a structure capable of storing data of more than one type at the same time. As an example, let's say we are in a spatial war against many alien species. Your job is to catalogue how many battleships each army has. One way to store this data is with a type containing a string for the species which the army belongs together with a natural number for how many battleships they have. We will write this type as

$$\mathbf{Army} := \text{ships } \mathbf{String} \times \mathbb{N}.$$

Here, $\text{ships}: \mathbf{String} \times \mathbb{N} \rightarrow \mathbf{Army}$ is called a *constructor* for the type **Army**. Constructors are (possibly multivariate) functions that receive terms as arguments and return a term of a compound type.

After you finish your catalogue, orders arrive for you to develop a new laser cannon to be used in the war. This weapon should scan the sky to find the enemy armies' positions (modeled in a type **Pos** and shoot them. But beware! There are also allied bases around, encapsulated in the type **Base**. Since friendly fire is really far from the ideal, our target type should have two different constructors, one for allies and another for enemies:

$$\begin{aligned} \mathbf{Target} := & \text{ally } \mathbf{Position} \times \mathbf{Base} \\ & | \text{enemy } \mathbf{Position} \times \mathbf{Army}. \end{aligned}$$

In here we extended our notation for enumerating types to also accept constructors. Different constructors always produce different terms of the compound type, thus, we may view the functions $\text{ally}: \mathbf{Position} \times \mathbf{Base} \rightarrow \mathbf{Target}$ and $\text{enemy}: \mathbf{Position} \times \mathbf{Army} \rightarrow \mathbf{Target}$ as tags representing from which type our **Target** was constructed. This works as if we are storing an element of type **Army** together with a tag **enemy** on the type **Target**. So the definition of target is saying that its terms are of the form $\text{ally}(p, x)$ or $\text{enemy}(p, x)$, just like we previously enumerated the terms of finite types.

This point of view allows us to define functions on compound types by enumerating what it does on the different constructors, a method called *pattern matching*. For example, the function $\text{not}: \mathbf{Bool} \rightarrow \mathbf{Bool}$ is defined as

$$\begin{aligned} \text{not true} &= \text{false} \\ \text{not false} &= \text{true}. \end{aligned}$$

While our ally-aware cannon shooting function may be defined as something like

$$\begin{aligned}\text{shoot}(\text{enemy}(p, y)) &= \text{laser_blast}(p) \\ \text{shoot}(\text{ally}(p, x)) &= \text{wave_your_hand}(p).\end{aligned}$$

Taking advantage of the type system via pattern matching is a nice way to make it obvious that our code does what it should. In this case, the types don't allow us to obliterate our friends.

Although these compound types are nice to organize our data, the true power of types comes from taking advantage of two other constructions: *function types* and *fixed point types*. The function type between two types A and B , represents all the functions receiving an input of type A and returning an argument of type B . To mimic the usual notation for type signatures, this function type is denoted as $A \rightarrow B$. A type system with function types allows us to define higher-order functions. For example, given a function $\phi: A \rightarrow B$, the composition operator K_ϕ defined as

$$K_\phi f = f \circ \phi$$

has type signature $K_\phi: (B \rightarrow C) \rightarrow (A \rightarrow C)$. As you can see, it takes a function and turns it into another one. This was a simple example but be sure that many more will soon come, since higher-order functions are at the heart of recursion schemes.

Now we arrive at the last kind of compound type we are going to discuss today: fixed point types. These are the star of today's show, so pay attention. We may define a compound data type that is parameterized by some variable (a function between types, if you prefer) such as

$$\text{maybe } A := \text{nothing} \mid \text{just } A.$$

Given a type A , a term of type $\text{maybe } A$ is either **nothing** or the constructor **just** applied to a term of type A . Thus, **maybe** receives a type and augments it with a special point².

Given a type operator F , a fixed point of F is a type X satisfying

$$X \simeq FX.$$

In here, we use the isomorphism sign $\simeq$ instead of equality because there is some boilerplate involved regarding tags. Although this is a simple construction, the concept is extremely powerful, being directly connected to recursion.

As an example let's explore the fixed point of **maybe**. It is a type N satisfying

$$N \simeq \mathbf{maybe} N = \mathbf{nothing} \mid \mathbf{just} N.$$

This means that any term of N is either **nothing** or **just** a term of N . Since we know **nothing**, we can construct a new term **just(nothing)**, then **just(just(nothing))**, and proceed successively in this way. Thus for any natural number n , applying **just** to **nothing** n times defines a unique term of N . Moreover, the definition of N says that all of its terms are of this form, meaning that N is isomorphic to the natural numbers³.

If you find this last result strange, remember that the natural numbers are inductively defined as being either zero or the successor of another natural number. Using our notation for types, this is equivalent to saying

$$\mathbb{N} = \mathbf{zero} \mid \mathbf{succ} \mathbb{N}.$$

If we alter the tag names, this tells us that $\mathbb{N}$ is a fixed point of **maybe** as expected.

An example with lists

The natural numbers are surely the most famous example of an inductive type. Nevertheless, almost no one thing of them like that while programming. Treating **3** as **succ(succ(succ(zero)))** would be cumbersome, to say the least. Any language comes with a (possibly signed) integer type already bundled with the usual arithmetic operations defined for it. Thus, let's discuss a little bit about another inductive type that is also pretty famous but has more of a inductive data structure flavour to it: the *linked list*.

Let A be your favorite type. Intuitively, a list over A is a finite-length sequence of elements of A . The simplest case possible is an empty list. To represent a non-empty list, we notice that any non-empty list with n elements may be “factorized” into its first element and another list with the remaining $n - 1$ elements.



Thus, by defining a type operator P as

$$PX := \text{nil} \mid \text{cons } A \times X,$$

the previous discussion shows that the type of lists over A , hereby denoted $L(A)$, is a fixed point of P ,

$$L(A) \simeq \text{nil} \mid \text{cons } A \times L(A).$$

Here, the constructor **nil** takes the role of the empty list while the constructor **cons** represents a pair containing an element of A and another list.

This inductive definition of lists allows us to recursively define operations over it. For example, let's construct a function that sums all the elements in a list of integers. Its signature is **sum**: $L(\mathbb{Z}) \rightarrow \mathbb{Z}$. If the list is empty, the sum returns zero. Otherwise, it takes the first element and adds it to the sum of what remains,

$$\begin{aligned} \text{sum}(\text{nil}) &= 0, \\ \text{sum}(\text{cons}(x, l)) &= x + \text{sum}(l). \end{aligned}$$

This last definition conceals an extremely useful pattern. It is the programming analogue of a proof by induction over the list. The empty list is the base case, which we set to zero. On a list of length $n > 0$, we assume that we've already solved the computation for the sublist of length $n - 1$ and then add it to the remaining element. Just like an inductive proof, see?

As it stands, you would be right to guess that constructing functions in this way is so pervasive in programming as proofs by induction are in mathematics. A simple variation of **sum** would be a function that multiplies a list of real numbers, **prod**: $L(\mathbb{R}) \rightarrow \mathbb{R}$. To construct it, all we need to do is take the definition of **sum**, replace 0 by 1, replace addition by multiplication, and it is all done!

$$\begin{aligned} \text{prod}(\text{nil}) &= 1, \\ \text{prod}(\text{cons}(x, l)) &= x \cdot \text{prod}(l). \end{aligned}$$

As expected, this pattern appears every time we want to somehow combine the elements of a list into a single value. It is so common that people have abstracted it on a function called **reduce**⁴. Its type signature is

$$\text{reduce}: B \times (A \times B \rightarrow B) \rightarrow (L(A) \rightarrow B).$$

Pretty scary, right? Let's break it down to see how it is just good ol' induction. When reducing a list of type $L(A)$ into a value of type B , there are two situations we may encounter. In the base case the list is empty, and we must specify

a value of type B to be returned. In the other steps, we consider that we already know the solution for the sublist of length $n - 1$, which yields a value of type B , and then need a rule to combine this value of the remaining term of type A into another term of type B . Thus, **reduce** is an instance of what we call a *higher-order function*. It is a machine that takes an initial value and a binary function and outputs another function, now defined on lists. We call it higher-order because it doesn't work with ordinary data, no, it transforms simple functions into recursive ones! Considering what it does, its definition is actually rather simple. The result $h = \text{reduce}(v, g)$ is the function that does

$$\begin{aligned} h(\text{nil}) &= v, \\ h(\text{cons}(x, l)) &= g(x, h(l)). \end{aligned}$$

Take a moment to absorb this definition, there really is a lot encapsulated on this. First substitute v by 0 and g by $+$ to see that it becomes **sum**. Then, substitute v by 1 and g by $\cdot$ to see that it becomes **prod**. Finally, congratulate yourself because you just understood your first recursion scheme!

If the way **reduce** was introduced made you think that it is only used to collapse a list into a primitive type, you should know that it is much more ubiquitous than that. Two of the most used list-manipulating functions in functional programming are **map** and **filter**. And, guess what, both are implementable in terms of **reduce**. The first, **map**, takes a function $f: A \rightarrow B$ and turns it into another function that applies f elementwisely to a list. Its type signature for **map** is therefore

$$\text{map}: (A \rightarrow B) \rightarrow (L(A) \rightarrow L(B)).$$

On the empty list, **map**(f) does nothing since it has no elements. On a **cons** node, it should apply f to the element stored on it and then proceed to apply f to the list's tail. Thus, the definition of **map** in terms of **reduce** is

$$\begin{aligned} g(x, l) &= \text{cons}(f(x), l), \\ \text{map}(f) &= \text{reduce}(\text{nil}, g). \end{aligned}$$

If it is hard to wrap your head around this last definition, try doing it step by step with a simple function such as x^2 and a small list such as $(1, 2, 3, 4)$. I promise you it will be enlightening.

The **filter** function takes a predicate, represented as $p: A \rightarrow \mathbf{Bool}$, and outputs a function that filters a list. That is removes all elements of the list for which p is false. Its type signature is thus

$$\mathbf{filter}: (A \rightarrow \mathbf{Bool}) \rightarrow (L(A) \rightarrow L(A)).$$

Since the empty list has no elements, there is nothing to do in this case. In a $\mathbf{cons}(x, l)$, we should return it exactly if $p(x)$ is true and return just l otherwise. We can assembly this in terms of **reduce** as

$$h(x, l) = \begin{cases} \mathbf{cons}(x, l), & \text{if } p(x) = \mathbf{true} \\ l, & \text{if } p(x) = \mathbf{false} \end{cases}$$
$$\mathbf{reduce}(f) = \mathbf{reduce}(\mathbf{nil}, h).$$

I hope you liked these examples because the next step in our journey is generalizing **reduce** to any inductive data type! To achieve this, we must pass through the land of category theory.

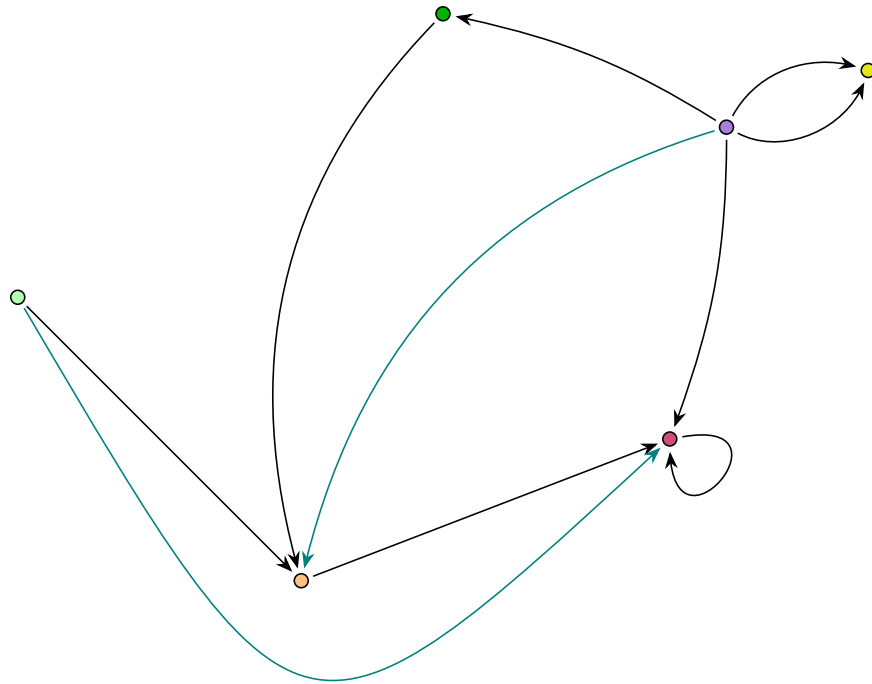
A walk through the land of categories

While the previous sections dealt more with programming, this one leans more to the mathematical side. But be assured that the fruits of all this abstraction will pay well.

A nice thing about types and functions between them is that they form a *category*.⁵ In fact, with all the structure we defined, they have the much more powerful structure of a *Cartesian closed category*. But we aren't going to use this last part today. If you happen to not know what a category is, I'm afraid I can't properly introduce them in here.⁶ But I will go through the basics for the sake of completeness.

A category $\mathcal{C}$ may be imagined as a special kind of graph. There is a collection of vertices, called its *objects* and between each pair of vertices there is a collection of *directed edges*, called *arrows* or *morphisms*. The A and B are objects, the notation for an arrow r from A to B is the same as for functions, $r: A \rightarrow B$. Soon we will see why. What differentiates a category from a simple directed graph is that we have a notion of *composition* between the arrows. That is, if

$f: A \rightarrow B$ and $g: B \rightarrow C$, there always exists another arrow $g \circ f: A \rightarrow C$. You may think of arrows as representing paths and $g \circ f$ as the arrow path going through f and then through g .



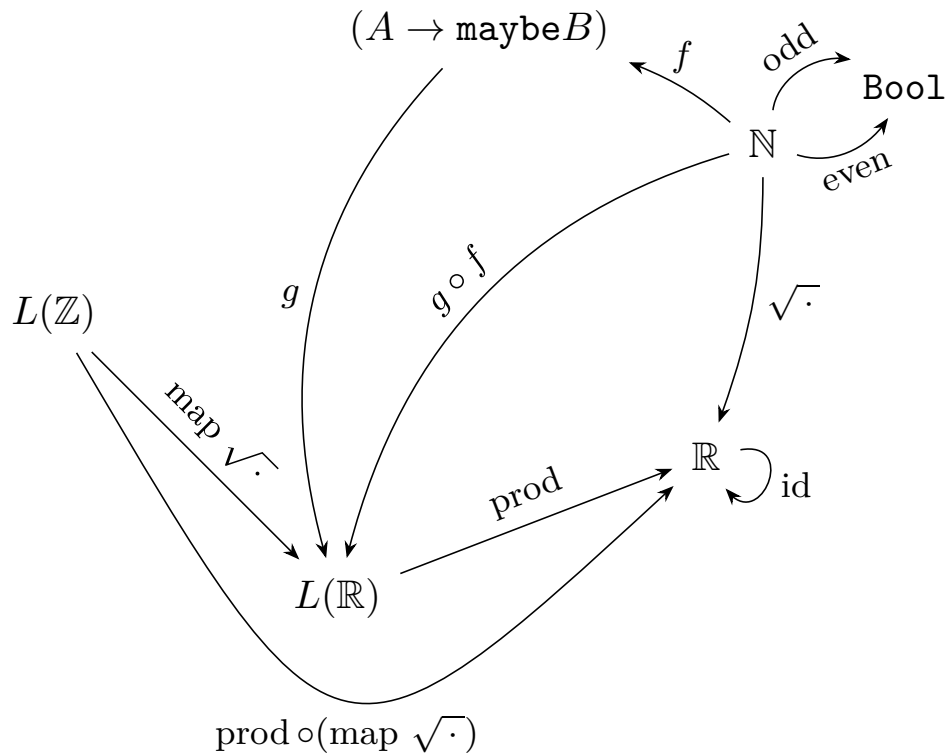
To call $\mathcal{C}$ a category, composition must also satisfy two laws. First, it must be associative. That is, for any composable triple of arrows,

$$h \circ (g \circ f) = (h \circ g) \circ f.$$

Second, for each object X , a special arrow $\text{id}_X: X \rightarrow X$ must exist, called the *identity* for X , with the special property that for any $f: A \rightarrow B$,

$$f \circ \text{id}_A = \text{id}_B \circ f = f.$$

Now, back to types. Probably the most striking feature of functional programming is that it is stateless. That is, we don't think of functions as a series of procedures mutating some registers, we think of them as mathematical functions and as such, we must have that $f(x)$ returns the exact same value no matter when we execute it. For us, the main consequence of this is that function composition becomes associative. Thus, there is a category **Types** whose objects are types and arrows are functions between types. Composition is the usual for functions and the identities are the identity functions $\text{id}_A(x) = x$.



Category theory is not only concerned with categories but also with transformations between them. As in most of mathematics, we study the transformations that preserve the structure of composition and identity, called *functors*. More precisely, a functor from $\mathcal{C}$ to $\mathcal{D}$ is comprised of a function taking objects of $\mathcal{C}$ to objects of $\mathcal{D}$ and a function taking arrows of $\mathcal{C}$ to arrows of $\mathcal{D}$ (normally both denoted by F) such that composition and identities are preserved,

$$F(g \circ f) = F(g) \circ F(f)$$

$$F(\text{id}) = \text{id}.$$

As an example, the type operator **maybe** is a functor from **Types** to itself when coupled with the operation on arrows

$$\text{maybe}(f)(\text{nothing}) = \text{nothing}$$

$$\text{maybe}(f)(\text{just}(x)) = \text{just}(f(x)).$$

More generally, if a type operator G is defined using only products, coproducts, other functors, and function types with the variable appearing only on the right-hand side there is a canonical way to turn it into a functor from **Types** to itself. Some compilers, such as Haskell's GHC [can even do that automatically for us](#). Instead of going through the constructions gory details, I think doing an example is better to elucidate how the method. Say we have a type operator

$$\begin{aligned}
GX &= \text{pig } A \\
&\quad | \text{yak } A \times X \times (\text{maybe } X) \\
&\quad | \text{cow } B \times (C \rightarrow X).
\end{aligned}$$

To turn it into a functor, we define an action on functions as

$$\begin{aligned}
(Gf)(\text{pig}(a)) &= \text{pig}(a) \\
(Gf)(\text{yak}(a, x, m)) &= \text{yak}(a, f(x), \text{maybe}(f)(m)) \\
(Gf)(\text{cow}(b, g)) &= \text{cow}(b, f \circ g)
\end{aligned}$$

Essentially, what Gf does is to unwrap each constructor and pass f to each of its arguments according to one of the following rules:

1. If the term is of type X , apply f to it.
2. If the term is of type $(A \rightarrow X)$, compose f with it.
3. If the term is of type FX for some functor F , apply Ff to it.
4. Otherwise, just return the term itself.

Functors and their fixed points

Now we know that both the natural numbers and lists aren't just the fixed points of ordinary type operators but of functors from **Types** to itself. But is there anything special in this? As you may expect, the answer is *yes*. If a type is the fixed point of a functor, then we can define recursion schemes for it that work very much like structural induction. Since this construction works on any category, it won't bite to do that in full generality. We begin by fixing a category $\mathcal{C}$ and a functor F from $\mathcal{C}$ to itself. From F , we will construct an auxiliary category, called the category of *F -algebras*.

A F -algebra is an object X of $\mathcal{C}$ together with an arrow $f: FX \rightarrow X$. Morphisms between F -algebras $f: FX \rightarrow X$ and $g: FY \rightarrow Y$ are given by an arrow $h: X \rightarrow Y$ such that the following diagram commutes

$$\begin{array}{ccc}
FX & \xrightarrow{Fh} & FY \\
f \downarrow & & \downarrow g \\
X & \xrightarrow{h} & Y
\end{array}$$

You can check that this definition turns F -algebras into a category where the identity for $f: FX \rightarrow X$ is given by id_X itself. On this category, we are interested in a special kind of object called an *initial F -algebra*. That is a F -algebra in with the special property that for any other F -algebra f there is a *unique* morphism going from in to f . Ok, time for some action. We haven't proved any theorem in this post yet.

Theorem.

Any initial F -algebra $\text{in}: FI \rightarrow I$ is an isomorphism in $\mathcal{C}$.

Before we prove this theorem, notice that it implies that I is a fixed point of F ! It's not every fixed point that defines an initial algebra, but the just the ones that are "smallest" in a certain sense.⁷ Nevertheless, finite length inductive data structures are generally initial.

Proof. First, notice that if $\text{in}: FI \rightarrow I$ is a F -algebra, so is $F\text{in}: F(FI) \rightarrow FI$. Since in is initial, there is a unique arrow $g: I \rightarrow FI$ making the following diagram commute

$$\begin{array}{ccc} FI & \xrightarrow{Fg} & F(FI) \\ f \downarrow & & \downarrow F\text{in} \\ I & \xrightarrow{g} & FI \end{array}$$

Since in itself may be viewed as a morphism between the F -algebras $F\text{in}$ and in , we get that their composition is a morphism from in to itself represented by the following diagram, where all paths still commute

$$\begin{array}{ccccc} FI & \xrightarrow{Fg} & F(FI) & \xrightarrow{F\text{in}} & FI \\ \text{in} \downarrow & & \downarrow F\text{in} & & \downarrow \text{in} \\ I & \xrightarrow{g} & FI & \xrightarrow{\text{in}} & I \end{array}$$

Since in is initial, the unique arrows going from it to itself is the identity. Thus, we must have $g \circ \text{in} = \text{id}$. Moreover, from the definition of functor and the previous diagram,

$$\text{in} \circ g = (Fg) \circ (F\text{in}) = F(g \circ \text{in}) = F(\text{id}) = \text{id}.$$

Therefore, g is an inverse to in , concluding the proof. □

Using catastrophes in your favor

We finally have the necessary tools to define our first recursion scheme in all its glory and generality. Consider, just as in the previous section, a functor F from **Types** to itself and call its initial algebra $\text{in}: FA \rightarrow A$. Given a F -algebra $f: FX \rightarrow X$, its *catamorphism*, which we will denote by $\text{cata } f$, is the unique arrow from A to X given by the initially of in .

Before proceeding, I must allow myself a little rant: sometimes mathematicians are just terrible name-givers. Like, what in hell is a catamorphism? What kind of intuition should it elicit? Well... as a matter of fact the name makes a lot of sense if you happen to be speaking in ancient Greek. Since, despite the arcane name, the catamorphism is a super cool concept, I think it deserves that we do a little etymological break to explain where its name comes from.

The word catamorphism has two parts, *cata-* + *morphism*. The latter comes from the Greek ‘μορφή’ and means “*form*” or “*shape*”. Its use is common in category theory and stems from the fact that usually, the arrows in a category are structure (or shape) preserving functions. For example, in the category of F -algebras, the morphisms are functions that commute with the algebras, thus “preserving” its application. The prefix “cata-” comes from the Greek ‘κατά’ and means something like a “*downward motion*”. It is the same prefix as in “cataclysm” or “catastrophe” and this is the intuition we shall keep in mind. If you think of an inductive data structure as a great tower piercing the sky, the catamorphism is a divine smite that, in a catastrophic manner, collapses the tower into a single value.⁸

Ok, back to mathematics.

As you may have imagined, the catamorphism is a generalization of the **reduce** operator. The difference is that while **reduce** only works for lists, **cata** can collapse the initial algebra of any functor. But, in all this generality, how do we actually compute a catamorphism? Let’s start by taking a look at its type signature:

$$\text{cata}: (FX \rightarrow X) \rightarrow (A \rightarrow X).$$

It is a higher-order operator that takes a function from FX to X and turns it into another function, now from the fixed point A to X . Intuitively, what it encapsulates is that if we know how to collapse one step of a data structure, then we can use structural induction to collapse the entire structure.

To properly calculate `cata`, let's take a look at the commutative diagram defining it,

$$\begin{array}{ccc} & FA & \xrightarrow{F(\text{cata } f)} FX \\ \text{in}^{-1} \nearrow & \downarrow \text{in} & \downarrow f \\ & A, & \xrightarrow{\text{cata } f} X \end{array}$$

In this diagram, there are two ways to go from A to X . From its commutativity, we know that they are equal. Hence, `cata` must satisfy

$$\text{cata } f = f \circ F(\text{cata } f) \circ \text{in}^{-1}.$$

So, that's a pretty concise formula. Let's unpack it with some examples before trying to give a full explanation.

The first inductive data type we encountered where the natural numbers, defined as the least fixed point of

$$FX = \text{zero} \mid \text{succ } X.$$

For it, the initial algebra and its inverse are defined as

$$\begin{aligned} \text{in}(\text{zero}) &= 0, & \text{in}^{-1}(0) &= \text{zero}, \\ \text{in}(\text{succ } n) &= n + 1, & \text{in}^{-1}(n) &= \text{succ}(n - 1). \end{aligned}$$

If these definitions seem too obvious, remember that the $\mathbb{N}$ are the natural numbers as we think of them everyday while `zero` and `succ` are the formal constructors of the functor F . Keeping track of this boilerplate is essential.

As our first example, let's construct the function `exp`: $\mathbb{N} \rightarrow \mathbb{R}$, which takes a natural n to the real number e^n , as a catamorphism. Exponentiation is recursively defined by the equations

$$\begin{aligned} e^0 &= 1 \\ e^{n+1} &= e \cdot e^n. \end{aligned}$$

In terms of an algebra $f: F\mathbb{R} \rightarrow \mathbb{R}$, this means that on the constructor **zero** we must return **1** while in each **succ** step, we must return e times the already accumulated value. Then, $\text{exp} = \text{cata } f$, where

$$\begin{aligned} f(\text{zero}) &= 1 \\ f(\text{succ } x) &= e \cdot x. \end{aligned}$$

Let's use the catamorphism formula to unwrap the calculations of exp . First, in^{-1} writes a number as the **succ** of its predecessor. Then, the way we defined the functor action of a data type means that $F(\text{cata } f)$ unwraps **succ** n to apply $\text{cata } f$ to the natural number n stored inside it and then reapplies the constructor, leaving us with **succ**(($\text{cata } f$)(n)). This process recursively continues until we encounter a **zero**. But this time $F(\text{cata } f)$ *does not* reapplies $\text{cata } f$; instead, it returns **zero** as is. At this point, our call stack consists of the function f applied exactly n times to the constructor **succ** applied exactly n times to **zero**. For example, let's take a look at the traceback of calling $\text{exp}(2)$:

$$\begin{aligned} \text{exp}(2) &= (f \circ F(\text{cata } f) \circ \text{in}^{-1})(2) \\ &= f(F(\text{cata } f)(\text{in}^{-1}(2))) \\ &= f(F(\text{cata } f)(\text{succ } 1)) \\ &= f(\text{succ}(\text{cata } f(1))) \\ &= f(\text{succ}(f(F(\text{cata } f)(\text{in}^{-1}(1))))) \\ &= f(\text{succ}(f(F(\text{cata } f)(\text{succ } 0)))) \\ &= f(\text{succ}(f(\text{succ}(\text{cata } f(0)))) \\ &= f(\text{succ}(f(\text{succ}(f(F(\text{cata } f)(\text{in}^{-1}(0))))) \\ &= f(\text{succ}(f(\text{succ}(f(F(\text{cata } f)(\text{zero})))))) \\ &= f(\text{succ}(f(\text{succ}(f(\text{zero})))))) \\ &= f(\text{succ}(f(\text{succ } 1))) \\ &= f(\text{succ}(e \cdot 1)) \\ &= e \cdot (e \cdot 1). \end{aligned}$$

Ok, these were a lot of parentheses. However, if you actually followed that mess above, the catamorphism's pattern should be clear by now.

As a final example of catamorphism, let's write a little calculator that supports addition, multiplication and exponentiation by a natural. A calculator should take an arithmetic expression and return a real number. We define an expres-

sion recursively. It may be a real number, the sum of two other expressions, the multiplication of two expressions, or an expression raised to a natural exponent. This is represented by a type **Expr** which is the least fixed point of the functor

$$\begin{aligned} F X = & \text{const } \mathbb{R} \\ & | \text{add } X \times X \\ & | \text{mult } X \times X \\ & | \text{pow } X \times \mathbb{N}. \end{aligned}$$

To evaluate an expression, we need an appropriate F -algebra $f: F\mathbb{R} \rightarrow \mathbb{R}$. As with natural numbers, the idea in here is to treat the constructor **const** where X don't appear as a base case and to think of the others constructors as storing an already solved problem on the X . With this in mind, the evaluator F -algebra is

$$\begin{aligned} f(\text{const}(a)) &= a \\ f(\text{add}(x, y)) &= x + y \\ f(\text{mult}(x, y)) &= x \cdot y \\ f(\text{pow}(x, n)) &= x^n. \end{aligned}$$

And the evaluator **eval**: **Expr** $\rightarrow \mathbb{R}$ is just **cata** f .

Another application of a catamorphism is if instead of evaluating the expression, we want to print it as a string. Let's say we have a method **str** that converts numbers to strings and an operator $\diamond$ that concatenates two strings. Erring on the side of too many parentheses, a candidate F -algebra is

$$\begin{aligned} g(\text{const}(a)) &= \text{str}(a) \\ g(\text{add}(x, y)) &= x \diamond "+" \diamond y \\ g(\text{mult}(x, y)) &= "(" \diamond x \diamond ")" * "(" \diamond y \diamond ")" \\ g(\text{pow}(x, n)) &= "(" \diamond x \diamond ")" ^ " \diamond \text{str}(n). \end{aligned}$$

As you probably already expect, the function **show**: **Expr** $\rightarrow$ **String** that converts an expression to a string is just **cata** g .

Least fixed points are recursive data structures. The catamorphism abstracts the process of transforming these structures into another type via structural induction. All recursion occurs inside the formula for **cata**. If you just use it as a black box that turns functions $FX \rightarrow X$ into functions $A \rightarrow X$ (where A is the

fixed point), you will never actually see the recursion's autoreference. The operator `cata` abstracts recursion just like a `for` loop abstracts the `goto` away. It is still there but following a fixed structure.

Since `cata` is a more rigid structure than general recursion, it is easier to reason about it. For example, one of the strongest properties of catamorphisms is that if the recursive data structure is finite and f is total function, then `cata f` is also guaranteed to eventually stop; As a lot of applications use finite data, this facilitates a lot the act of debugging a program.

Taking your functions to the gym

Collapsing a data structure into a value has many use cases but is not the end of the story. Sometimes we want to do the exact opposite: take a value as a seed to construct a list or another structure from it. This notion is dual to the catamorphism and, of course, there is also a recursion scheme to encapsulate it: the *anamorphism*.⁹

Again we have a pretty arcane name in our hands. Let's take a look at its etymology in order to clarify things. The prefix *ana-* comes from the Greek 'ἀνα' and generally means "upward motion" but is also used to mean "strengthening" or "spreading all over". It is the same prefix of *analogy* and of *anabolic steroids*. In both these words, the prefix means that we are building something up.

In order to describe the catamorphism in all its generality and yet with a concise formula, we had to dive into some categorical abstract nonsense. As you might expect, the same applies to the anamorphism. Fortunately, there is a tool in category theory called *duality*. If we simply reverse the arrows in a theorem, its conclusion still holds but also with some arrows reversed.

Given a category $\mathcal{C}$ and a functor $F: \mathcal{C} \rightarrow \mathcal{C}$, we define a F -coalgebra to an object X of $\mathcal{C}$ together with an arrow $f: X \rightarrow FX$. The F -coalgebras form a category where the morphisms between $f: X \rightarrow FX$ and $g: Y \rightarrow FY$ are arrows $h: X \rightarrow Y$ such that the diagram below commutes.

$$\begin{array}{ccc}
 FX & \xrightarrow{Fh} & FY \\
 f \uparrow & & \uparrow g \\
 X & \xrightarrow{h} & Y
 \end{array}$$

The dual notion to an initial object F -algebra is a *terminal* F -coalgebra. Namely, an algebra $\text{out}: S \rightarrow FS$ such that there is a unique morphism from any other F -coalgebra to out . As you probably already expect, terminal coalgebras are always isomorphisms. The proof is essentially the same as the one for initial F -algebras. The only difference being some arrows reversed.

Theorem.

Any terminal F -coalgebra $\text{out}: C \rightarrow FS$ is an isomorphism.

The most direct consequence of this theorem is that S must be a fixed point of F . Since there is an arrow from every other coalgebra to S , it is called the *greatest fixed point* of F . For this presentation, there is no need to discriminate between least and greatest fixed points. Specially because there are languages such as Haskell where they are the same. Just keep in mind that when working with coalgebras, there is no guarantee that functions must eventually stop. We are no longer on the land of finite data structures but in the much greater land of possibly infinite data structures.

Given a F -coalgebra $f: X \rightarrow FX$, its anamorphism $\text{ana } f$ is defined as the unique arrow from f to the terminal object

$$\begin{array}{ccc}
 FS & \xleftarrow{F(\text{ana } f)} & FX \\
 \text{out}^{-1} \left(\begin{array}{ccc} \uparrow & & \uparrow \\ \text{out} & & f \end{array} \right. & & \\
 S & \xleftarrow{\text{ana } f} & X
 \end{array}$$

Thus, the type signature of ana viewed as a higher-order operator is

$$\text{ana}: (X \rightarrow FX) \rightarrow (X \rightarrow S).$$

It turns a coalgebra, which we may view as one construction step, into a function that constructs an entire inductive data structure. If we pay attention to the commutative diagram above, we can see that there is a concise recursive formula defining ana ,

$$\mathbf{ana} \, f = \mathbf{out}^{-1} \circ F(\mathbf{ana} \, f) \circ f.$$

This is the same formula we had for `cata`, but the order of composition reversed! This reversion, however, means that we are building structures instead of collapsing them.

Let's do some examples with lists. Recall that the type $L(A)$ of lists with elements of type A is a fixed point of the functor

$$PX = \mathbf{nil} \mid \mathbf{cons} \, A \times X.$$

Using the usual notation $[a, b, c, \dots]$ to represent lists, the F -coalgebra `out` is defined as

$$\begin{aligned} \mathbf{out}([\] &= \mathbf{nil} \\ \mathbf{out}([a, b, c, \dots]) &= \mathbf{cons}(a, [b, c, \dots]). \end{aligned}$$

One of the simplest list anamorphisms is a function that receives a natural number n and returns a decreasing list from n to 1. It is defined as `ana g`, where g is the coalgebra

$$g(n) = \begin{cases} \mathbf{nil}, & n < 1 \\ \mathbf{cons}(n, n - 1), & \text{otherwise.} \end{cases}$$

In practice, it is most common to use an increasing list, however. The induction on this one is a little trickier but nothing an anamorphism can't handle. To make it more fun, let's construct a function `range`: $\mathbb{Z} \times \mathbb{Z} \rightarrow L(\mathbb{Z})$ taking a pair of integers to the closed integer interval between them. We want `range(-2, 1)` = $[-2, -1, 0, 1]$ and `range(7, 5)` = $[\]$, for example. To achieve this, we will need a family of coalgebras

$$g_b(n) = \begin{cases} \mathbf{nil}, & n > b \\ \mathbf{cons}(n, n + 1), & \text{otherwise.} \end{cases}$$

Then, our range function is the anamorphism

$$\mathbf{range}(a, b) = (\mathbf{ana} \, g_b)(a).$$

If you ever dealt with plotting libraries, you are probably familiar with a `linspace` method. It receives two real arguments a , b , and a natural argument n and returns a list of n points uniformly distributed between a and b . The construction of such a function

$$\text{linspace}: \mathbb{R} \times \mathbb{R} \times \mathbb{N} \rightarrow L(\mathbb{R})$$

is a simple variation on the **range** we defined earlier. We start with a family of coalgebras

$$g(b, n)(x) = \begin{cases} \text{nil}, & x > b \\ \text{cons}(x, x + \frac{b-a}{n}), & \text{otherwise,} \end{cases}$$

and define it to be the anamorphism

$$\text{linspace}(a, b, n) = (\text{ana } g(b, n))(a).$$

To end this section, let's turn to an example from number theory. We will construct the **sieve of Eratosthenes**.¹⁰ This algorithm receives a natural number n and returns a list of all primes below n . The idea is to start with a list of numbers between **2** and n , and to recursively refine it eliminating all multiples of a given prime. Let $\text{test}(p)$ be the predicate that tests if a number is not divisible by p . You can implement it by testing if the remainder of the input by p equals zero, for example. This refinement process is encapsulated by the function $\text{sieve} = \text{ana } g$ where g is the coalgebra

$$\begin{aligned} g([]) &= \text{nil} \\ g([p, x, \dots]) &= \text{cons}(p, (\text{filter}(\text{test}(p))([x, \dots]))). \end{aligned}$$

And our prime-listing function is the composition

$$\text{era}(n) = \text{sieve}(\text{range}(2, n)).$$

This process works because after each filtering, the list's first element cannot be divisible by any number below it. Thus it must be prime. Notice that this function eventually stop because as soon as we reach an empty list, the algorithm returns a **nil**.

Finally, to show the power of anamorphisms coupled with lazy evaluation, let's use **sieve** to compute a list of all primes. We begin by writing a list l of all natural numbers starting at **2**,

$$\begin{aligned} h(x) &= \text{cons}(x, x + 1) \\ l &= (\text{ana } h)(2). \end{aligned}$$

This list must be infinite because, since h never returns `nil`, the anamorphism recurses forever. Unsurprisingly, the list of all primes is defined as `primes = sieve(l)`. At first, this may not seem very useful since any computer would take infinite time to calculate `primes`. But, after some thought, you will notice that the way `sieve` calculates its output is rather special. After it produces an element of the list, it never touches that element again! Thus, although it's impossible to properly calculate `primes`, we have that for any natural N , the first N elements of `primes` are calculated in a finite amount of time. So you can think of `primes` as an iterator that generates `era(n)` in finite time.

This last property is common to all anamorphisms. Although their output can be a possible infinite data structure (even if the input is finite), it is produced in an extremely structured manner. When a value is computed, we go to the next and it is never touched again. Therefore, we can still use an anamorphism to calculate finite data so long as we properly say when we wish to stop. In fact, if we take a look at programs such as games or operational systems, this is exactly the behaviour we want from them. Imagine if any OS terminated after a finite amount of steps... No, what we want is for the OS to run indefinitely producing some well-defined actions in finite time.

Building up to a collapse

By now, we've seen two kinds of recursion schemes. Anamorphisms start with a seed value and end with a data structure constructed from it, while catamorphisms start with a data structure and collapse it to end with result. Both of these are powerful tools but their type signatures are too constrained. They must explicitly return or receive a data structure. What if we want to write a recursive function from a primitive value to another primitive value? Our next (and last for today) recursion scheme addresses exactly this.

Our strategy will be to use an inductive type as our middle man. The *hylomorphism* takes a value, applies an anamorphism to turn it into a data structure and then applies a catamorphism to collapse it into another value.¹¹ This means that we don't need to go through yet another category to define it. No, given a functor F , the type signature for `hylo` is

$$\text{hylo}: (FB \rightarrow B) \times (A \rightarrow FA) \rightarrow (A \rightarrow B)$$

and the definition is simply¹²

$$\mathbf{hylo} \, f \, g = \mathbf{cata} \, f \circ \mathbf{ana} \, g.$$

The etymology for hylomorphism is a little harder to motivate but let's try anyway. The prefix *hylo-* comes from the Greek 'ύλη' and means "wood" or "matter". And as we've previously seen, the term morphism comes from the Greek word for form. The name hylomorphism is a pun on an Aristotelian theory of the same name which says that the being is composed from matter and form. Since I never really understood Aristotle's writings and was unable to find another word starting with *hylo-* outside the context of philosophy, I will just try to improvise an intuition. Let's think of algebras and coalgebras as ways to give form or create matter and a **hylo** combines them into a single being. It's kinda lame but was the best I could think of.

Although the first recursive function introduced in this post was the factorial (way back in the [motivation](#)) we haven't rewritten it as a recursion scheme yet. It's time to fix that. If you open the definition of **fat**, you will see that **fat**(*n*) stand for the product of the first *n* natural numbers,

$$\mathbf{fat}(n) = \prod_{k=1}^n k.$$

Thus, an algorithm to compute the factorial is to first construct a decreasing list of integers from *n* to 1 and then collapse it by multiplying all its elements. We've already constructed both these functions but let's rewrite their definitions for the sake of completeness. We start with the coalgebra and algebra,

$$\begin{aligned} g(n) &= \begin{cases} \mathbf{nil}, & n < 1 \\ \mathbf{cons}(n, n-1), & \text{otherwise,} \end{cases} \\ f(\mathbf{nil}) &= 1, \\ f(\mathbf{cons}(x, y)) &= x \cdot y, \end{aligned}$$

and finally define the factorial as **fat** = **hylo** *f* *g*.

Time for a more complex example, let's see how to use hylomorphisms to better shape the call procedure tree and get faster algorithms. Recall how we defined **exp** as a catamorphism. For what we knew at the time, it was fine but, at its $O(n)$ complexity, it's just too slow. With a hylomorphism, we can do better than that. [The trick](#) is noticing that for even values, the exponential satisfies $e^{2n} =$

$(e^{n/2})^2$, which gives us a much better recursive relation. Thus, instead of multiplying e n times, we can construct a much smaller call tree and collapse it to get the desired result. We define a type **CallTree** as the fixed point of the functor

$$TX = \text{leaf } \mathbb{R} \mid \text{square } X \mid \text{mult } \mathbb{R} \times X.$$

This type encapsulates the call tree. As base cases, n equal to zero or one means that we store e^n on a **leaf** node. If n is even, we construct a **square** node and pass $n/2$ to it, meaning that the value of this node will be squared on the algebra. Finally, if n is odd, we store e and $n - 1$ on this node, which will be multiplied when folding with the algebra. The coalgebra that construct this structure is

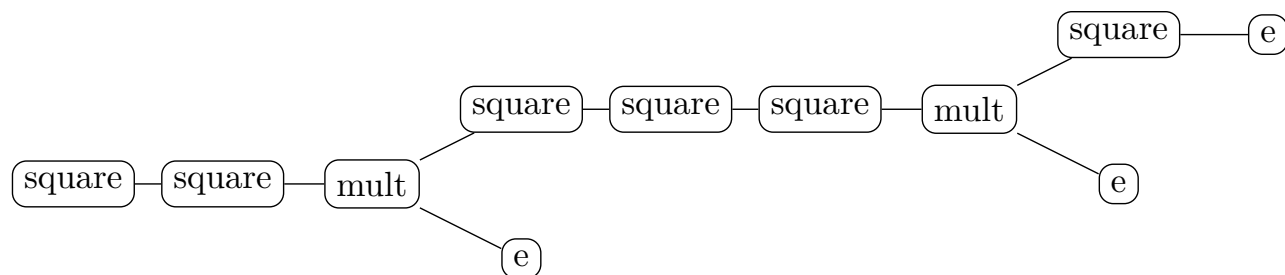
$$g(n) = \begin{cases} \text{leaf}(1), & n = 0 \\ \text{leaf}(e), & n = 1 \\ \text{square}(\frac{n}{2}), & n \text{ is even} \\ \text{mult}(e, n - 1), & n \text{ is odd.} \end{cases}$$

To collapse the tree, we will use the algebra

$$\begin{aligned} f(\text{leaf } x) &= x \\ f(\text{square}(x)) &= x^2 \\ f(\text{mult}(x, y)) &= x \cdot y. \end{aligned}$$

Finally, we define the exponential as the hylomorphism $\text{exp} = \text{hylo } f g$.

When we analyzed the traceback for the exponential as a catamorphism, we noticed that it did a multiplication for each **succ** that appeared until we reached **zero**. Since a natural number n is the n -th successor of zero, this amounts to n multiplications. Viewing a natural number in such a unary representation has its theoretical advantages but is too slow and cumbersome in practice. With our new definition using a hylomorphism, we let the binary representation of n aid us in designing a more efficient call tree. Since the number of multiplications necessary is proportional to $\log_2 n$, we get an algorithm with complexity $O(\log n)$. This is much better! If you're not impressed, take a look at the call tree generated for $n = 100$ and see how it requires much less than 100 computations! Better yet, to compute e^{200} , this tree would have to be augmented by only one node.



Hylomorphisms really shine when an algorithm implements an intermediate data structure that is used but not returned. A special class of such are the divide-and-conquer algorithms, whose structure lend them to have elegant implementations as hylomorphisms. There are many examples such as merge sort, quick sort, convex hull and many others that can be written as an anamorphism that divides the input in a structured manner followed by a catamorphism that conquers the output from it. Please try to write any of these as a **hylo**! If you pick your favorite algorithm, there is a great chance it can be written in a clear manner as a hylomorphism.

To end this section (and this post), we will do a last derivation showing that hylomorphisms do not need to be defined in the way they were. No, they accept a recursive definition in terms of themselves as any good recursion scheme should.

Theorem.

Given a F -algebra f and a F -coalgebra g , their hylomorphism satisfies

$$\mathbf{hylo} \, f \, g = f \circ F(\mathbf{hylo} \, f \, g) \circ g.$$

Proof. Since the least and greatest fixed points of F are equal, we can assume that `in` and `out` are inverses to each other. By pasting together the diagrams defining the anamorphism and the catamorphism we get

$$\begin{array}{ccccc}
FX & \xrightarrow{F(\text{ana } g)} & FS & \xrightarrow{F(\text{cata } f)} & FY \\
g \uparrow & & \text{in} \left(\begin{array}{c} \uparrow \\ \downarrow \end{array} \right) \text{out} & & \downarrow f \\
X & \xrightarrow{\text{ana } g} & S & \xrightarrow{\text{cata } f} & Y
\end{array}$$

Going from X to Y through the bottom path amounts for our definition of hy-lomorphism. Since the diagram is commutative, it must be equivalent to trace the top path. This, together with the functoriality of F , implies

$$\begin{aligned}
\text{hylo } f g &= f \circ F(\text{cata } f) \circ F(\text{ana } g) \circ g \\
&= f \circ F(\text{cata } f \circ \text{ana } g) \circ g \\
&= f \circ F(\text{hylo } f g) \circ g.
\end{aligned}$$

Thus we conclude the proof. □

Although our definition of a **hylo** is good to reason about, this formula has two advantages over it. The first one is practical: because there is less boilerplate happening, we only have half the function calls happening. As you can see in the proof's second line, our definition requires **hylo** to first recurse through **ana**, stacking a pile of calls to **in**, to immediately thereafter recurse through **cata**, canceling each **in** with a **out**. This does not change the algorithm's complexity but still hinders its efficiency. In very large programs that run for weeks, even the complexity constant hurts. There is a practical difference between a program that runs in a week and a program that runs in two.

The second reason is more theoretical: This last formula has no reference to fixed points. There is no explicit mention of any data structure whatsoever. All we have to do is give **hylo** one rule saying how to construct each step and one rule saying how to destruct each step and the call stack becomes our intermediate data structure! This is much cleaner and spare us from thinking about least and greatest fixed points. For short, when thinking about hylomorphisms, use the definition as an anamorphism followed by a catamorphism but when actually implementing one, use this formula instead.

Summary

This was a long journey but we finally got to the end. Throughout it we met with our three heroes: catamorphism, anamorphism, and hylomorphism, and I truly hope that the ideas they encapsulate will appear on your future programs.

Although these three are powerful tools, they're not the only recursion schemes available. Oh no, there is an entire **zoo of them** and, as you may expect, every single one has an arcane ancient Greek prefix. Each recursion scheme encapsulates a common pattern of recursion and comes in two flavours: construction and destruction of an inductive type. The ones we've seen today are the schemes for structural recursion but if instead we wanted primitive recursion, we would use a para- or apomorphism. It is even possible to combine different

schemes to end up with monsters such as the infamous [zygohistomorphic pre-promorphisms](#). Moreover, as one can expect of anything related to category theory, all of them are special cases of [a certain categorical construction](#).

Another topic I haven't given the deserved attention are the fusion laws for a recursion scheme. The recursive formulas we derived for each scheme may be seen as certain algebraic equations that they must obey. Besides these, there are some other laws related to composition of schemes or other contexts. These give us some of the guarantees of structured programming we talked earlier. If a compiler knows that the code only use certain recursion schemes, it can use the fusion laws to substitute your code for something semantically equivalent but much more efficient. You get the type safe, organized code and the machine gets the faster, optimized one. Win-win.

Well, that's all for today. Good recursion for y'all!

1. This sum example is simple to code with a loop or linear recursion. Recursion schemes really shine when manipulating more complex data structures, such as trees with varying arity.↩
2. This is also called `Option` in some languages and is specially useful to safely model partial functions and computations that may fail.↩
3. Choosing the letter N was no accident.↩
4. Also known as `accumulate` or `foldr` in some languages. The 'r' means that this function folds a list from the right.↩
5. This is not entirely true. Questions such as non-termination or laziness/strictness may break this structure in some languages, see [this link](#) for a discussion concerning Haskell. Nevertheless, types are similar enough to a category to be worth it to think of them as forming one.↩
6. Take a look on chapter 1 of [Emily Riehl's book](#) if you want to learn this (and much more) properly. Trust me, it's a great book for the mathematically inclined.↩
7. The technical term is *least fixed point*.↩
8. Fortunately, as functional languages have no side-effects, applying a catamorphism doesn't make everyone start talking in a different language.↩
9. Also called `unfold` in the context of lists.↩

10. This example is adapted from an anamorphism in the lecture notes [Programming with Categories](#) by Brendan Fong, Bartosz Milewski and David I. Spivak.↵
11. As it stands, hylomorphisms are the theoretical analogous of me playing any construction game as a kid.↵
12. To write that, we must assume that the least and the greatest fixed point of F coincide. Although this assumption may sound strange, it always holds for languages such as Haskell.↵