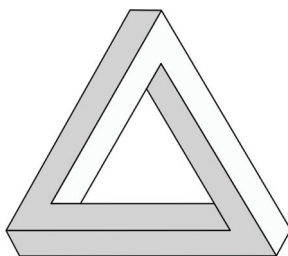


9 Cellular Sheaf Cohomology through Examples

In which we start to look at some more involved and computationally explicit examples, working up toward an extended introduction to cellular sheaf cohomology and giving a brief glimpse into the theory of cosheaves.

Roughly, if sheaves represent local data—or, more precisely, represent how to properly ensure that what is locally the case everywhere is in fact globally the case—sheaf cohomology can be thought of as a tool for systematically measuring and relating *obstructions* to such passages from the local to the global. In particular, sheaf cohomology in low degrees gives an explicit measure of the failure of local data to patch. Moreover, in sheaf theory more generally, one could argue (as does Grothendieck, for instance) that individual sheaves are only of secondary importance—the real power of sheaf theory emerges from the use of constructions involving various sheaves, linked together via sheaf morphisms. The exposition of cellular sheaf cohomology that follows will allow us to begin to appreciate such a perspective. There is a more general theory of sheaf cohomology in terms of homological algebra and derived categories, but the focus here will be on the simpler case of cellular sheaf cohomology, leaving the reader to explore the more general theory on their own.

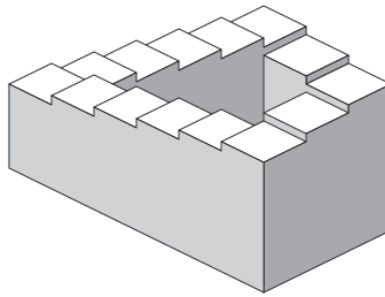
As motivation for the general idea of cohomology, let us briefly consider *impossible objects*. Consider the “Penrose tribar,” first devised by Roger Penrose:



Focusing one’s eye on any sufficiently small pieces of this depicted figure—for instance, regarding the tribar as being assembled from three L-shaped pieces—suggests that the drawing is the projection of a visually viable spatial object. In other words, breaking it up into possible pieces, such a 2-D picture suggests an object that is *locally consistent*. However, while it may initially seem plausible that the pieces of the tribar may assemble together into the closed triangle depicted above, the local depth data cannot in fact

be consistently merged, and there is no plausible interpretation of the entire perspectival drawing as a projection of a viable spatial object that respects standard visual expectations. The 3-D object that the drawing of the Penrose tribar would depict is not a possible object in ordinary Euclidean space—hence, the Penrose tribar is sometimes referred to as an “impossible object.” In other words, *globally* such an object is inconsistent; but *locally* (i.e., as one considers any small enough region of the drawing), there is no such impossibility in what the depiction represents. To adopt another way of seeing the same issue: there is a property—the *impossibility* of such an object—that ceases to be valid locally, that is, is *nonlocal*.

Similarly, consider the “Penrose staircase,” a 2-D depiction of a staircase where the stairs take four 90-degree turns as they descend in a continuous loop, something that is impossible in ordinary 3-D Euclidean space:



Visually, it is to be understood that we might have a number of local snapshots of four individual sides or corners of the staircase, where we would like to patch these together into a viable 3-D globally consistent staircase. Each part of the structure is viable as a representation of a directed flight of stairs. However, any attempt to merge all these partial representations must fail—as the steps are continually descending in a clockwise direction—and so the representation must be of an object that, globally, is inconsistent. Since the implied heights will not all match up, this inconsistency thus presents an obstruction to any such attempted patching. In short, (1) you have a number of local pictures, that is, snapshots of the individual sides of the staircase; which (2) you would like to patch together into a globally consistent object (the full staircase); yet (3) there exists an obstruction to step (2), as what one expects for the heights does not match.

With such impossible figures, we are fundamentally being presented with an impossibility of patching locally consistent data into a consistent whole. Locally, such impossible figures look entirely possible, but globally, they are not. Cohomology can help track such obstructions. A decisive feature of cohomology is that it is fundamentally nonlocal. Such impossible objects can help build some intuition for (sheaf) cohomology, for the “impossibility” of any correct 3-D representation can be captured by an element of cohomology. The obstruction described in the third step above shows up in the form of a nonzero element of some cohomology group.¹³⁴

134. Penrose (1992) uses cohomology to make more precise the fact that such figures are impossible globally, while being viable when viewed locally. Example 231 below covers the main ideas of that paper.

One of the main tools in sheaf theory is sheaf cohomology, which can be seen as a generalization of some classical cohomology theories (de Rham cohomology and Čech cohomology). In a rough sense, sheaf cohomology is fundamentally an invariant that quantifies and tracks obstructions to extensions of local sections of a sheaf to global sections. The so-called first cohomology of a sheaf will accordingly capture the set of things that locally appear just like sections, yet globally may not come from a section.

To begin to tell this story—focusing ultimately on the case of cellular sheaf cohomology—let us first build up some necessary background.

9.1 Simplices and Their Sheaves

Of the many ways to represent a topological space, a particularly computationally friendly way is to perform a triangulation with entities called *simplices*, decomposing the space into simple pieces (thought of as being glued together) whose common intersections or boundaries are lower-dimensional pieces of the same kind. With simplices come certain simplicial maps that, moreover, approximate continuous maps. In this way, simplices play a role in bridging the gap between continuous figures and their discrete representation and approximation via decompositions of spaces into discrete parts. More than that, as we will see, they allow us to develop profound connections between algebra and geometry. Simplices are powerful and easy-to-use devices for understanding qualitative features of data collections, and in general they can be thought to represent n -ary relations between n vertices.

Basically, we use collections of simplices—for now, just think of points, line segments, generalized triangles, or tetrahedra generalized to arbitrary dimensions—to build up what are called *simplicial complexes*. Given a simplex σ , it is common to refer to a (nonempty) simplex τ whose vertices are a subset of the vertices of σ as a *face* of σ . A (geometrical) simplicial complex K is a collection of simplices such that (1) every face of a simplex of K is in K , and (2) the intersection of any two simplices of K is a face of each of them. Roughly, then, one can think of a simplicial complex K as comprising generalized triangles of various dimensions, glued together along common faces. This is really part of a more general story involving *cell complexes* (including cubical complexes, multigraphs, etc.), where one can roughly think of a cell complex as a collection of *closed disks* of various dimensions which are moreover glued together along their boundaries. But we will instead focus on the more computationally tractable combinatorial counterpart to the already simplified notion of a simplex: that of *abstract simplicial complexes* (ASC). Here, we reencode the information of a simplicial complex via the more computationally friendly notion of an ASC, where this is basically just a collection of finite subsets of vertices (elements), closed under the operation of taking subsets. This captures in a purely combinatorial way the geometrical notion of simplicial complex.¹³⁵

Definition 222 An *abstract simplicial complex* $K = (A, K)$ is a set A equipped with a collection of ordered finite nonempty subsets $K \subseteq \mathbb{P}(A)$ that contains all singletons and that is

135. What has been lost is how the simplex is embedded in, say, Euclidean space; however, this specification retains all the data needed to reconstruct the complex up to homeomorphism.

closed under taking subsets (sublists), that is, every nonempty subset of a set in K is also in K . In other words, we must have

- for each $x \in A$, the singleton $\{x\} \in K$; and
- if $\sigma \in K$ and $\tau \subseteq \sigma$, then $\tau \in K$.

Terminologically, each member of K is called a *simplex* (or, looking ahead, *cell*), and given a $\tau \subseteq \sigma$ we say that τ is a *face* of σ . A simplex with $n+1$ elements is called an n -dimensional simplex of K . (But, as one would expect, a 0-face is usually just called a *vertex*, and a 1-face an *edge*.) If all the simplices of an ASC K are of dimension n or less, K is said to be an n -dimensional simplicial complex, that is, the dimension of an ASC is the maximal dimension of its constituent simplices. A *simplicial map* $f: K \rightarrow K'$ from an ASC K (defined on a set A) to an ASC K' (defined on the set B) is a function $f: A \rightarrow B$ with the property that for any σ of K , the image $f(\sigma)$ is an element of K' .¹³⁶

Altogether, this data in fact lets us define the category **SCpx** that has (abstract) simplicial complexes as objects and simplicial maps as morphisms.

To check understanding, notice that a 1-D simplicial complex essentially recovers the notion of a simple graph.

ASCs are particularly easy to describe, but one might worry that certain topologically valuable information gets lost in encoding things in this simplified, set-theoretical fashion. We will ultimately be interested in certain topological information, so it makes sense to construct for a given ASC K a geometrical *realization* $|K|$ of K , allowing K to be realized (basically, “pictured”) as some (generalized) triangles glued together in suitable ways, living in a subspace of $\mathbb{R}^n$. For every simplicial complex K , there in fact exists a unique (up to simplicial isomorphism) geometric realization $|K|$, so we will not bother to distinguish between ASCs and their geometric realizations.¹³⁷ Such a realization better explains why we think of objects (sets) in an ASC K as faces or simplexes, since via the realization, three-element sets correspond to filled-in triangles, two-element sets to edges, singletons to vertices. For instance, given a set $A = \{a, b, c\}$ for which we have an ASC $K = \{\{a, b\}, \{b, c\}, \{a, c\}, \{a\}, \{b\}, \{c\}\}$, then its realization $|K|$ will be the (hollow) triangle (with a natural orientation). Drawing such pictures yields the usual geometrical notion of a simplicial complex, obtained by gluing together the standard simplices along the boundaries; the usual approach then employs these n -simplices to probe a topological space via continuous maps into the space.¹³⁸

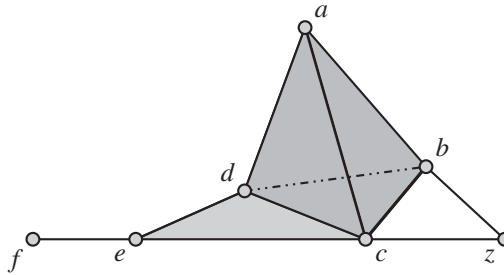
In more detail, observe that an oriented 0-simplex thus corresponds to a signed $(+, -)$ point P , while an oriented 1-simplex is a directed line segment P_1P_2 connecting the points P_1 and P_2 , where we assume that we are traveling in the direction from P_1 to P_2 , that is, $P_1P_2 \neq P_2P_1$ (however, $P_1P_2 = -P_2P_1$). An oriented 2-simplex will be a triangular region $P_1P_2P_3$, with a prescribed order of movement around the triangle. An oriented 3-simplex is given by an ordered sequence $P_1P_2P_3P_4$ of four vertices of a solid tetrahedron. Similar

136. Simplicial maps between simplicial complexes are the natural equivalent of continuous maps between topological spaces.

137. Ensuring the uniqueness of this is precisely the reason for considering *ordered* sets (lists) in the definition of an ASC, instead of just unordered sets.

138. More details on these matters can be found in Ghrist (2014) or Hatcher (2002).

definitions hold for $n > 3$. By gluing together various simplices along their boundaries, we get simplicial complexes such as the following:

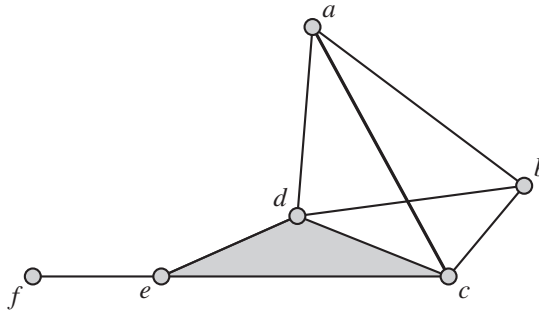


While it is perhaps useful to visualize things in this way, and while this perspective can be important for connections to other concepts, the alert reader might have observed that given the way ASCs were defined, they should already come with a natural topology, letting us bypass the geometric realization, and associate to each ASC a particular topological space. This will be useful to us in the construction of sheaves on such spaces toward which we are building.

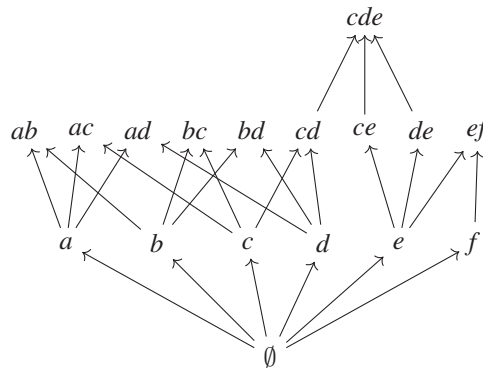
To see this, first observe that ASCs come with a canonical partial order on faces, given by the face subset inclusion (or attachment) relation between vertices, edges, and higher-dimensional faces. We can use this fact to define the *face* (or *cell*) *category*, where this has for objects the elements of K , a cell complex, and (setwise) inclusions of one element/cell of X into another for its morphisms; if a and b are two faces in a complex X with $a \subseteq b$ and $|a| \leq |b|$, we will write $a \rightsquigarrow b$ and say that a is *attached* to b .¹³⁹ We will then identify a complex with its face poset, writing the incidence relation $a \rightsquigarrow b$. Building on the graphical construction of a complex, the attachments between the faces or cells of a complex can then be displayed in attachment diagrams, where the links represent attachments going from lower- to higher-dimensional cells (and where any additional attachments that arise as compositions of attachments are left implicit). Observe that the attachment diagram of a graphical complex is itself just a set partially ordered via the attachment relations, that is, it is a poset. The data of this face relation poset can be displayed with a Hasse diagram. Assume we have the following simplicial complex K , which we imagine has been realized thus:¹⁴⁰

139. Note that technically to make the following construction work we need to use the more general notion of *cell complexes*, the full definition of which can be found in Curry (2014, chap.4); but since ultimately the realization $|K|$ of an ASC K on a finite set, which is the sort of thing we will be dealing with in our example, can be shown to be a cell complex, and simplicial maps $f : K \rightarrow K'$ induce a cellular map $|f| : |K| \rightarrow |K'|$, we will not worry too much about the distinction.

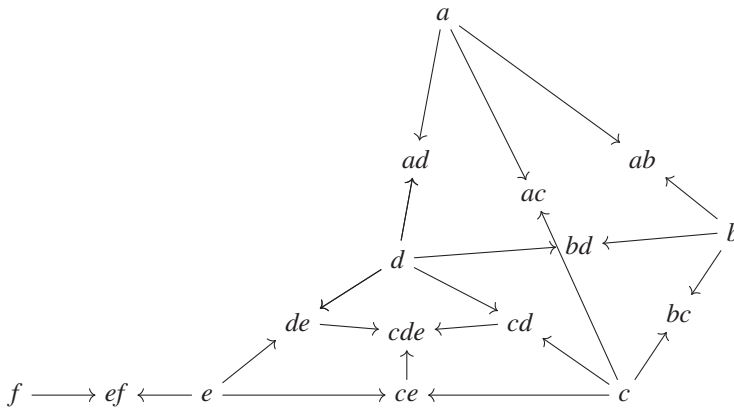
140. Note that the rightmost simplex ($abcd$) is *not* meant to depict a hollow *tetrahedron*; each of the four component triangles are to be thought of as lying in the plane. We have simply spaced it this way to make the sheaf diagrams we build on top of this in a moment a little more readable.



We can form the diagram of the face-subset relations, where the edges and higher faces are given the natural names (and are assumed to be ordered lexicographically):



But it is more revealing to display this in the form of an attachment diagram, as follows:



It is this face/cell incidence poset that we are regarding as a category, $\mathcal{F}_K$. In other words, to a simplicial complex we can associate a category, which is just the face-incidence poset viewed as a category. We can then put the Alexandrov topology on this poset of face-relations. Then, given this Alexandrov topology on a poset, the usual topological notions (such as interiors, boundaries, and closures) can be easily understood in terms of the poset itself.

The basic idea is this: pick a simplex; then look at all the other simplices that include that one as a face (i.e., higher-dimensional simplices adjacent to it); then regard such “upper sets” as the open sets. The sets of the form $\uparrow x = \{y \in \mathcal{P} \mid x \leq y\}$ —that is, the principal upper sets—form a basis for the topology. We can also define the closure of x by $\bar{x} = \{y \in \mathcal{P} \mid y \leq x\}$; and, provided the poset $\mathcal{P}$ is finite, a basis of closed sets is given by these $\bar{x}$.¹⁴¹

In the cellular context, for σ a cell of a cell complex, the analogue of the principal upper set construction is called the *open star* of σ , where this is denoted $\text{st}(\sigma)$ and consists of the set of cells τ such that $\sigma \rightsquigarrow \tau$ —that is, it captures all the higher-dimensional cells containing that cell. Then, in terms of the topology, $\text{st}(\sigma)$ will be the smallest open set of cells containing σ , or $\text{st}(\sigma) = \bigcup_{x \in \sigma} U_x$, with $U_x = \uparrow x$. Taking all the stars and the union of all the stars will give a topology for the complex/simplex—this is the Alexandrov topology. Every intersection of opens in the Alexandrov topology on a poset $\mathcal{P}$ is open. Thus, a star over $A \subseteq X$ is then defined to be the intersection of the collection of all open sets containing A . The resulting collection of stars will be a basis for the Alexandrov topology. While stars need not exist in general, in the Alexandrov topology on a poset, there will exist a star of every subset. There is accordingly a dictionary between cellular complexes and Alexandrov spaces, which can be seen by considering that for a cell complex, every cell Δ_σ has a star, where this is a set consisting of all those cells Δ_τ such that $\Delta_\sigma \leq \Delta_\tau$.

Note that with respect to the inclusions in the face relation poset $\mathcal{F}_K$, the containment relation for the open sets (stars) in the Alexandrov topology is order-reversing. In other words, the Alexandrov construction will yield an order-reversing inclusion functor $\mathcal{P} \rightarrow \mathcal{O}(\mathcal{P})^{op}$, just as we saw in chapter 8. More generally, we have effectively described a contravariant functor $Alxd : \mathbf{Pre} \rightarrow \mathbf{Top}$ —one that, upon applying $Alxd$ to X^{op} , yields a space that has for open sets unions of simplices.

It turns out that a sheaf (of sets) over an ASC K can be defined as a covariant functor from the face category $\mathcal{F}_K$ of its associated face-relation poset to **Set**. More generally, given a functor from a preorder or poset $\mathcal{P}$ to a category **D**, this functor can be used to produce a sheaf on the Alexandrov topology (via a right Kan extension).¹⁴² With this construction, the sheaf gluing axiom for any cover is automatically satisfied! In other words, given a poset endowed with the Alexandrov topology, as anticipated in the last chapter, we do not even need to distinguish between presheaves and sheaves.

The following definition follows Shepard (1985), who defined sheaves for more general *cell complexes*, which are just a collection of closed disks of certain dimensionality that are glued together along the boundaries. Via its realization $|K|$, an ASC K is just a cell complex, so while the definition is more general, it can be applied to a simplicial complex (which is what we will work with in the coming example). Just as with ASCs, since cell complexes are built up from simple pieces (cells), the associated attachment diagram exhibiting the

141. A dual topology thus arises by considering, for a given simplex, all the other simplices that are attached to (included in) it—these are the downsets, which also can serve as the opens of this topology. In the Alexandrov topology, arbitrary intersections of opens are open and arbitrary unions of closed sets are closed; therefore, by exchanging opens with closed sets, we can pass from any Alexandrov space to its dual topology.

The Alexandrov topology construction is really appropriate when the elements of the underlying poset $\mathcal{P}$ represent finite pieces of information (i.e., are compact), something that is typical for many combinatorial and computer science applications; however, if $\mathcal{P}$ includes infinite elements, then the “Scott topology” is called for (see Vickers [1996, chap. 7] for details on this).

142. See Curry (2019) for details.

relations between cells contains the information of the cell complex itself. Attending to the face poset in particular, then, we will define a *cellular sheaf*, following Shepard, as a covariant functor from the face category of a complex K to some other category $\mathbf{D}$.¹⁴³ For concreteness, for the remainder we consider $\mathbf{D} = \mathbf{Vect}$.

Definition 223 A *cellular sheaf* (of vector spaces) F on a cell complex X is

- an assignment of a vector space $F(\sigma)$ to each cell σ of X ,¹⁴⁴ together with
- a linear transformation

$$F_{\sigma \rightsquigarrow \tau} : F(\sigma) \rightarrow F(\tau)$$

for each incident cell pair $\sigma \rightsquigarrow \tau$.

These linear maps have to further satisfy the identity relation $F_{\sigma \rightsquigarrow \sigma} = \text{id}$ and the usual composition condition, namely

$$\text{if } \rho \rightsquigarrow \sigma \rightsquigarrow \tau, \text{ then } F_{\rho \rightsquigarrow \tau} = F_{\sigma \rightsquigarrow \tau} \circ F_{\rho \rightsquigarrow \sigma}.$$

The reader may be wondering if there is a typo in the direction of the maps described in this definition of a cellular *sheaf*. A sheaf, after all, is a particular presheaf, so one would have expected (order-reversing) *restriction* maps. But this is not a typo, and in fact goes to the heart of the underlying construction. The reason for the seemingly “wrong” direction of the arrows—typically, sheaf restriction maps reverse the direction of the arrows, while cosheaves preserve them—is explained (as correct) by a result we have already encountered.

Recall that, in general, when dealing with a poset $\mathcal{P}$, we can regard a sheaf on $\mathcal{P}$ —once this has been equipped with the upper Alexandrov topology—as a plain old pre-cosheaf (covariant functor) on $\mathcal{P}$ (which could, in turn, be regarded as a presheaf on $\mathcal{P}^{op}$).¹⁴⁵ But the face category $\mathcal{F}_X$ with which we identify a complex X is a poset. So, using our general result¹⁴⁶

$$\mathbf{Sh}(\mathcal{U}(\mathcal{P})) \simeq \mathbf{D}^{\mathcal{P}}$$

letting us move freely between sheaves (valued in $\mathbf{D}$) on the upper Alexandrov topology placed on $\mathcal{P}$ and covariant $\mathbf{D}$ -valued functors (pre-cosheaves) on $\mathcal{P}$, we know that we can freely regard a plain old *covariant functor* $\mathcal{P} \rightarrow \mathbf{D}$ as a *sheaf* on $\mathcal{U}(\mathcal{P})$, that is, on the upper Alexandrov topology on $\mathcal{P}$. The inclusion taking a poset into its upper sets is order-reversing, and the underlying functor of a sheaf (now on the upper sets) is itself order-reversing—their composition, equivalent to the original covariant functor, is accordingly covariant. This accounts for why the definition of a cellular *sheaf* seems to just contain the data of a definition of a covariant functor—that is indeed all it says! Theorem 220,

143. While this, and the subsequent cellular sheaf cohomology, is original to Shepard, his thesis was never published and fell into oblivion for some time. Part of the intent of Curry (2014) was to revive Shepard’s contribution, while providing a more modern account, supplemented with missing details, demonstrating why cellular sheaves are actually sheaves. The account here thus owes a great debt to Curry (2014), which the reader is encouraged to consult for more extensive treatment of these matters.

144. This vector space $F(\sigma)$ is then the *stalk* of F at (or over) σ .

145. Also recall that, taking $\mathcal{P}^{op}$ instead, sheaves on $\mathcal{P}$ —where this is equipped with the *lower* Alexandrov topology (using that $\mathcal{D}(\mathcal{P}) \simeq \mathcal{U}(\mathcal{P}^{op})$)—are just presheaves on $\mathcal{P}$.

146. Earlier, in chapter 8, we just described it in terms of set-valued functors; but in fact, the equivalences hold for (pre)sheaves valued in a category $\mathbf{D}$, provided that category is complete and cocomplete.

discussed at the end of the last chapter, lets us conflate these two perspectives. In short, we have the slogan:

cellular sheaves are covariant functors from the face category into some other category $\mathbf{D}$.

In particular, a covariant functor from $\mathcal{F}_X$ to $\mathbf{Vec}$ is already just a sheaf of vector spaces.

Dually, we could define:

Definition 224 A *cellular cosheaf* (of vector spaces) $\hat{F}$ on a cell complex X is

- an assignment of a vector space $\hat{F}(\sigma)$ to each of the cells σ of X —this vector space $\hat{F}(\sigma)$ is called the *costalk* of $\hat{F}$ at (or over) σ —together with
- a linear transformation $\hat{F}_{\sigma \rightsquigarrow \tau} : \hat{F}(\tau) \rightarrow \hat{F}(\sigma)$ for each incident cell pair $\sigma \rightsquigarrow \tau$.

These maps—called the *corestriction maps*—have to further satisfy the identity relation $\hat{F}_{\sigma \rightsquigarrow \sigma} = \text{id}$ and the usual composition condition, namely

$$\text{if } \rho \rightsquigarrow \sigma \rightsquigarrow \tau, \text{ then } \hat{F}_{\rho \rightsquigarrow \tau} = \hat{F}_{\rho \rightsquigarrow \sigma} \circ \hat{F}_{\sigma \rightsquigarrow \tau}.$$

In the extended example to follow, we focus on cellular sheaves; at the end of the chapter, an example with cellular cosheaves is presented. Before turning to the example, let us also highlight a few more explicit definitions of the corresponding cellular notions that are more or less as you would expect for a sheaf.

Definition 225 For F a cellular sheaf on X , we define a *global section* x of F to be a choice $x_\sigma \in F(\sigma)$ for each cell σ of X , where this satisfies

$$x_\tau = F_{\sigma \rightsquigarrow \tau} x_\sigma$$

for all $\sigma \rightsquigarrow \tau$.

Fundamentally, the data of a cellular sheaf on a complex X amounts to a specification of spaces of local sections on a cover of X (namely, the one given by open stars of cells). Ultimately, we will be able to form the category of *all sheaves* $\mathbf{Sh}(X)$ over a fixed complex X , adopting the only notion of morphisms that there could be, namely as the natural transformations between the corresponding functors.¹⁴⁷ Explicitly,

Definition 226 A *morphism* $f : F \rightarrow G$ of *sheaves* (or *sheaf morphism*) on a cell complex X is an assignment of a linear map

$$f_\sigma : F(\sigma) \rightarrow G(\sigma)$$

to each cell σ of X , where for each attachment $\sigma \rightsquigarrow \tau$, the usual (natural transformation) compatibility condition holds, making the diagram

$$\begin{array}{ccc} F(\sigma) & \xrightarrow{f_\sigma} & G(\sigma) \\ \downarrow F(\sigma \rightsquigarrow \tau) & & \downarrow G(\sigma \rightsquigarrow \tau) \\ F(\tau) & \xrightarrow{f_\tau} & G(\tau) \end{array}$$

commute.

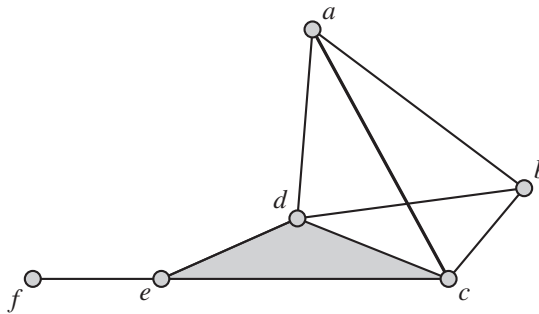
¹⁴⁷ In the next definition, we focus on the corresponding notion of morphism for sheaves of vector spaces; but if we were to work with sheaves valued in some other category $\mathbf{D}$, then we would require that the maps defined below be the appropriate structure-preserving map (e.g., for sheaves of groups, just homomorphisms).

A *sheaf isomorphism*, then, is also defined in the inevitable way, as a morphism where each of the f_σ is an isomorphism.¹⁴⁸

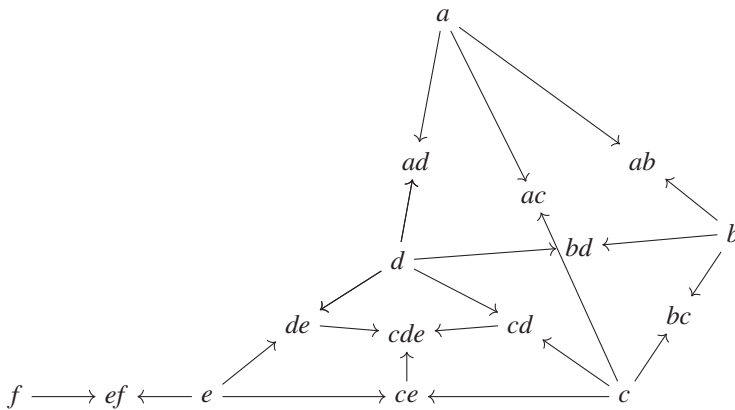
In brief: if X is a cell (or simplicial) complex with the associated face poset category $\mathcal{F}_X$, we identify the complex with its face category, and then a cellular sheaf is just a vector space-valued functor $F : \mathcal{F}_X \rightarrow \mathbf{Vect}$, while a cellular cosheaf is a vector space-valued functor $F : \mathcal{F}_X^{op} \rightarrow \mathbf{Vect}$. Thus, to avoid confusion, realize that in the above definitions of cellular (co)sheaves, X was really short for $\mathcal{F}_X$, to which we associate the cell complex, so that a cellular sheaf on X (per the definition) is really a sheaf on $\mathcal{F}_X$ where this has been equipped with the Alexandrov topology, which in turn is uniquely determined by a functor $\mathcal{F}_X \rightarrow \mathbf{Vect}$ (and *this* last functor is what the definition is describing).

In the example that follows, we illustrate these ideas in a concrete fashion by constructing a cellular sheaf on a particular simplicial complex.

Example 227 Recall our simplicial complex X from before, which we imagine has been realized thus:

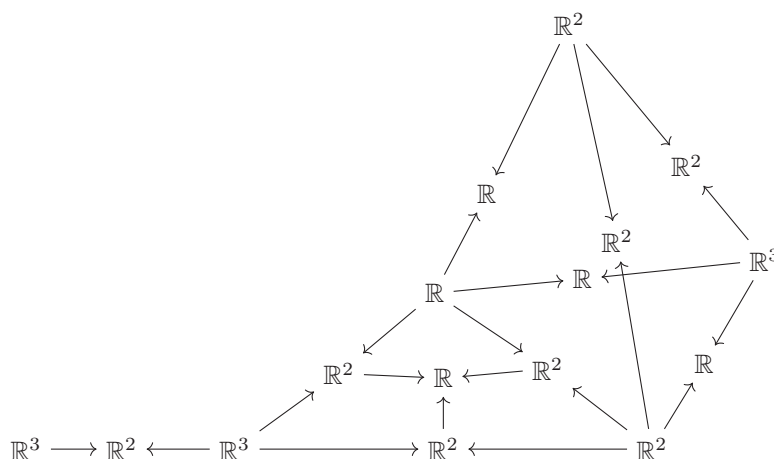


together with its associated attachment diagram displaying the poset of face relations (i.e., its face category $\mathcal{F}_X$):

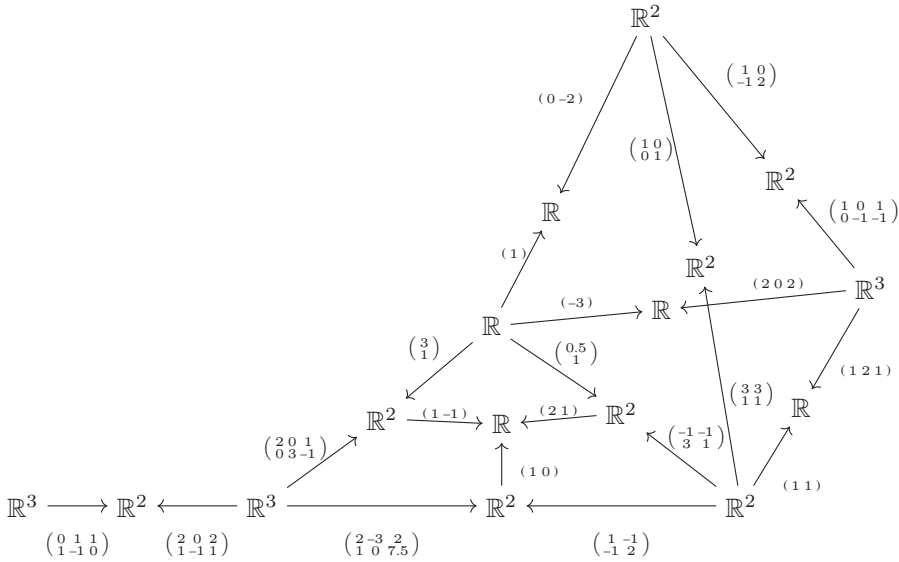


148. It is also easy to show that a morphism between sheaves of vector spaces, such as that given above, induces a linear map between the spaces of global sections of the sheaves; moreover, isomorphic sheaves will have isomorphic spaces of global sections.

Thus, following the definition of a cellular sheaf, for the values of the sheaf over cells, this will just amount to the specification or assignment of values (spaces) to each of the cells of the simplices, data that comes in the form of vectors, for example,

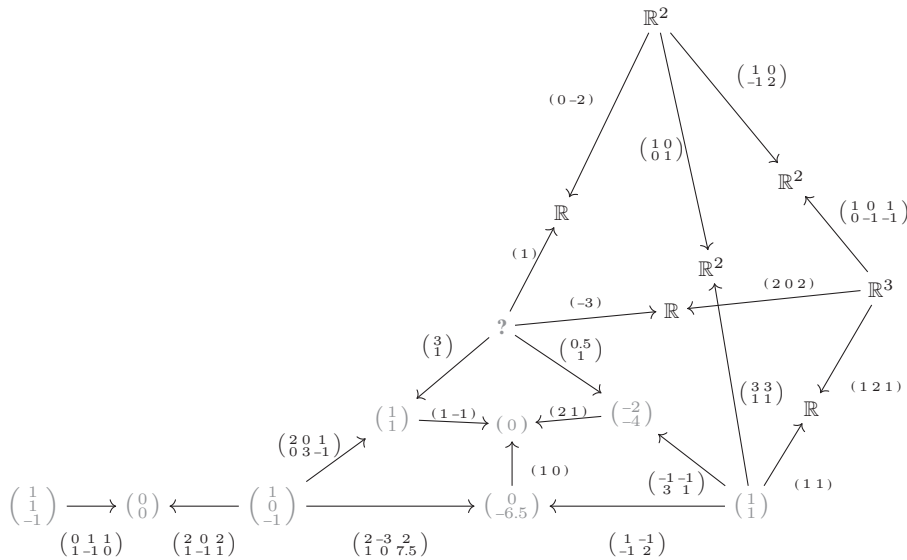


But we need maps as well. The maps, for their part, may be thought of as representing some sort of local constraints or as enforcing certain relations between the data. In general, when the stalks $F(\sigma)$ have structure—for instance, here they are vector spaces—then a sheaf of that type (i.e., valued in the relevant category) is obtained when the restriction maps preserve this structure. In other words, we should have a function (in our particular case of vector space assignments, these will be given by a linear map) assigned to each inclusion of faces in such a way that the diagram commutes, that is, the composition of functions throughout the diagram is path independent. The following sheaf diagram nicely displays all these ideas:



A sheaf is generated by its values specified on individual simplices of X , that is, by local sections specified on the vertices. But a sheaf is not just this data. The restriction maps of a sheaf are an essential part of the construction, as they encode how any local sections can be extended into sections over a larger part of the diagram (ultimately throughout the entire diagram), and so it is precisely via the restriction maps that it is made explicit how the local sections can be glued together. The sheaf assignment given over *all* of X will be specified by a collection of local sections that can be extended along all the restriction maps to higher-dimensional faces. There may be some flexibility or freedom in the actual data assignments over a vertex, but they are not entirely arbitrary, for the restriction maps encode how local assignments—values specified on certain parts of the diagram—can be extended to other parts of the diagram, and so the maps will constrain the assignments in various ways.

If the reader would like to get a good working understanding of the important distinction between a local section and a global section, it would be useful to closely consider what happens in this concrete case when we assign, for instance, the value $\begin{pmatrix} 1 \\ 0 \\ -1 \end{pmatrix}$ to the stalk over the vertex e versus what happens when we assign, for instance, the value $\begin{pmatrix} -1 \\ 2 \\ 2 \end{pmatrix}$ to the same vertex. In the first case, one finds that we can “extend” or propagate this particular selection along *some* of the arrows to those stalks highlighted in gray, but then there is a problem, an obstruction to our continuing this process any further upward along the edges of the diagram:



Observe that there is simply *no value* that might be placed at the stalk over vertex d (hence the “?”) that would allow us to continue with this extension process. If we assigned (-4) to the stalk over d , this would indeed be consistent with the map $\begin{pmatrix} 0.5 \\ 1 \end{pmatrix}$ proceeding down and to the right and landing in the stalk over cd , which would in turn land us, perfectly consistently, in the stalk over cde with the value (0) , as required by the other restriction maps. However, by following what happens to that same assignment -4 under the action of the map down and to the left via $\begin{pmatrix} 3 \\ 1 \end{pmatrix}$, we see that we would get $\begin{pmatrix} -12 \\ 4 \end{pmatrix}$, which, when further mapped under $\begin{pmatrix} 1 & -1 \end{pmatrix}$ would yield (-8) . We thus cannot assign -4 —or *anything* for that matter—to the stalk over d , given our original assignment over e .

The original assignment at the stalk over e , then, is said to describe a strictly *local section*, one that importantly cannot be extended globally, that is, over the entire complex. By beginning with other values at the same (or, if one desires, another) vertex, the reader can explore various other solutions that are merely local versus those that manage to be global. In this way, the reader will not only discover that some local solutions or sections are “more local” than others, but will discover other types of obstructions to the extension of local sections. For instance, whereas with our test assignment above it turned out that there was simply *no value at all* that could be assigned to the stalk over vertex d , while maintaining consistency with the other stalk assignments required by the restriction maps, another (less serious) issue one frequently encounters is that a specific assignment on one stalk ends up requiring two *different* assignments at some stalk.¹⁴⁹

149. We might here take the opportunity to mention that in a variety of applications of sheaves, including beyond sheaves of vector spaces on complexes, it is possible that *exact equality* of assignments will be unattainable or the least valuable thing to consider. There are a few ways to develop machinery to accommodate this, but it is beyond the scope of this book to cover them in detail. Instead, for one particularly friendly approach, the reader is referred to the work of Michael Robinson, who has proposed to deal with this situation via formalizing a “consistency structure” with a corresponding notion of *distance* between assignments (say, with the structure of a pseudo-metric space). Moreover, the notion of *pseudosections* can be developed, and Robinson has shown that in fact pseudosections are already sections, just with respect to a different sheaf; for instance, pseudosections of a sheaf over an ASC X are veritable sections of another sheaf over the barycentric subdivision of X . See Robinson

cellular sheaves on (possibly) different spaces. We would like to extend the notion of a morphism of cellular sheaves to provide maps between sheaves on different spaces. Of course, we could consider a morphism between sheaves on a fixed space X as a special case of this approach, but we have really been heading towards the more general case of a sheaf morphism involving different spaces.

Definition 228 A *sheaf morphism* $s : F \rightarrow G$ from a sheaf F over a space Y to a sheaf G over the space X consists of the following data:

- a cellular map $f : X \rightarrow Y$,
- a collection of (linear) maps $l_\sigma : F(f(\sigma)) \rightarrow G(\sigma)$ such that for each attachment map $\sigma \rightsquigarrow \tau$ in X , the following diagram commutes:

$$\begin{array}{ccc} F(f(\sigma)) & \xrightarrow{l_\sigma} & G(\sigma) \\ \downarrow F(f(\sigma) \rightsquigarrow f(\tau)) & & \downarrow G(\sigma \rightsquigarrow \tau) \\ F(f(\tau)) & \xrightarrow{l_\tau} & G(\tau). \end{array}$$

In other words, a sheaf morphism takes data in the stalks over two sheaves and relates them through linear maps in such a way that the resulting diagram commutes.

The reader may observe how this in fact makes use of the *pullback sheaf* notion, first mentioned in chapter 8, where for a map $f : X \rightarrow Y$, and a sheaf F on Y , the pullback f^*F will be a sheaf on X defined by $(f^*F)(\gamma) = F(f(\gamma))$ and $(f^*F)(\sigma \rightsquigarrow \tau) = F(f(\sigma) \rightsquigarrow f(\tau))$.

Such notions (and others, such as the pushforward sheaf) are clearly useful for switching base spaces. Sheaf morphisms can also be composed, under certain conditions, leading to *sequences of sheaves*, linked together by sheaf morphisms. Certain sequences will even exhibit special algebraic properties, like *exactness*, which will have significance in a number of applications. We take up these matters in the next section, after a brief discussion of a few further operations on sheaves.

We have seen a few constructions—such as that of subsheaves, pullback sheaves, and pushforward sheaves—where old sheaves are used to generate new ones. Here is a very brief look at just a few other important things one can do *to* or *with* sheaves, to generate new sheaves; specifically, we focus on indicating a few of the *algebraic* operations one can perform on sheaves. These notions can be defined in greater generality, but we will stick to the cellular context.

Definition 229 For F and G , two sheaves of vector spaces on a cell complex X , we can define $F \oplus G$, their *direct sum*, in the natural way:

$$(F \oplus G)(\gamma) = F(\gamma) \oplus G(\gamma)$$

and

$$(F \oplus G)(\sigma \rightsquigarrow \tau)(v, w) = (F(\sigma \rightsquigarrow \tau)v, G(\sigma \rightsquigarrow \tau)w)$$

for $v \in F(\gamma)$ and $w \in G(\gamma)$.

In a similar fashion, we could define $F \otimes G$, the *tensor product* of two sheaves F, G , in the expected way; but there is a subtlety here when we try to think of this in terms of sheaves, and we will not be making use of this, so instead we will just indicate an example of a direct sum of sheaves.

Example 230 Consider a network, that is, a 1-D cell complex with oriented edges. Ghrist (2014) provides some nice applications of cellular sheaves over such networks, called *flow sheaves*, where these represent the flow of a commodity (as in supply chains or various information or transportation of goods moving through networks). The underlying graph supports certain viable flow values, and one of the purposes of the sheaf is to encode these feasibility conditions or constraints. Basically a sheaf on such a network will supply some algebraic structure encoding a particular collection of logical, numerical, stochastic, or other constraints on the “flows” or transport of commodities through a network. One of the advantages of using sheaves, in this setting, is that we can easily generalize beyond numerical constraints on a network to other (perhaps noisy or logical) constraints.

Here is a very rough sketch of how this works.¹⁵⁰ A *flow* (or *flow sheaf*) on a network X is an assignment of coefficients (e.g., in $\mathbb{Z}$ or $\mathbb{N}$) to each edge of X , in such a way that a particular “conservation” condition is met (namely the sum of the incoming edge flow values equals the sum of the outgoing edge values, except at the special “external” vertex, where they are not conserved). Each such value can be imagined as representing an amount of a commodity or resource in transit at a location of the underlying graph. Restrictions $F(v \rightsquigarrow e)$ are then projections onto components. The *direct sum* of two flow sheaves $F \oplus G$ could then be used to represent the transportation of two *different* resources being transported along the same given network. In other words, the sum $(F \oplus G)$ would represent the number of both sorts of resources, so $(F \oplus G)(e) = \mathbb{N}^2$ would represent the number of F -items and G -items being carried along the edge e .

9.2 Sheaf Cohomology

The impatient reader might well be wondering at this point: “Okay, I understand what a sheaf is already! But what good is all this?” One glib answer, following Hubbard, might be that “without cohomology theory, they aren’t good for much”!¹⁵¹ While this seems too pessimistic—after all, even if all sheaves on their own did was organize a wealth of particular and highly disparate constructions involving local data into a powerfully general framework, this would be immensely valuable—it is a perspective that gets at something important. If sheaves represent local data—or, more precisely, represent how to properly ensure that what is locally the case everywhere is in fact, more than that, globally the case—sheaf cohomology is a device that lets us extract global information from local data and systematically explore, represent, and relate obstructions to the extension of the local to the global. In this way, sheaf cohomology can cope with situations where the local-to-global passage breaks down; and this is of immense value, since we would like to be able to handle and talk about structures that somehow fall short of assembling into sheaves. Moreover, as we mentioned at the outset of the chapter, we would like to appreciate a fact that Grothendieck insisted on, namely that individual sheaves are only of secondary importance—the real power of sheaf theory emerging from the use of constructions involving various sheaves, linked together via sheaf morphisms. Sheaf cohomology will allow us to glimpse this.

150. The reader who desires a less rough sketch of these notions is invited to look at Ghrist (2014).

151. See Hubbard (2006, 383).

In line with our categorical approach thus far, in our presentation sheaf cohomology will emerge as a *functor*, specifically as one with domain the category of sheaves (together with their sheaf morphisms) and codomain the category of vector spaces. The cellular sheaves that we will continue to work with, together with their cohomology, have the nice property that it is just as easy to compute with them as to compute the usual cellular cohomology of a cell complex, something one can learn about in more elementary contexts. Sheaf cohomology is particularly important to understand, though, so we do not assume that the reader already knows or recalls all they need to know about the basic notions of (co)homology. Over the next few pages, we accordingly build up to sheaf cohomology by first reviewing the basic notions of (co)homology with respect to ordinary simplices.¹⁵²

9.2.1 Primer on (Co)Homology

Given oriented simplices (or cell complexes), as described above, a very natural thing is to look at the *boundary* of a given n -simplex. As one might expect, the boundary of a 1-simplex P_1P_2 is simply the vertices of the edge; however, we must now carefully attend to the issue of orientation. Taking the boundary, an operation denoted by ∂ , is more precisely defined as taking the formal “difference” between the endpoint and the initial point, that is, $\partial_1(P_1P_2) = P_2 - P_1$. Similarly, the boundary of a 2-simplex $P_1P_2P_3$ is given by

$$\partial_2(P_1P_2P_3) = P_2P_3 - P_1P_3 + P_1P_2.$$

This in fact corresponds to traveling around what we intuitively think of as the boundary of a triangle in the direction indicated by the orientation arrow. The boundary of a 3-simplex is then defined as

$$\partial_3(P_1P_2P_3P_4) = P_2P_3P_4 - P_1P_2P_3 + P_1P_2P_4 - P_1P_2P_3.$$

The pattern should be clear, allowing us to define the boundary operator ∂_n more generally for $n > 3$:

$$\partial_k(\sigma) = \sum_{i=0}^k (-1)^i (v_0, \dots, \widehat{v_i}, \dots, v_k),$$

where the oriented simplex $(v_0, \dots, \widehat{v_i}, \dots, v_k)$ is the i -th face of σ obtained by deleting its i -th vertex. Notice that each individual summand (i.e., the positive terms) of the boundary of a simplex is just a *face* of the simplex.

We can associate some *groups* to a given complex X . The group $C_n(X)$ of oriented n -chains of X is defined to be the free abelian group generated by the oriented n -simplices of X . Every element of $C_n(X)$ is a finite sum $\sum_i m_i \sigma_i$, where the σ_i are n -simplices of X and $m_i \in \mathbb{Z}$. Then the addition of chains is carried out by algebraically combining the coefficients of each occurrence in the chains of a given simplex. For instance, considering the surface of a tetrahedron S (oriented in an obvious way), the elements of $C_2(S)$ will look like $m_1P_2P_3P_4 + m_2P_1P_3P_4 + m_3P_1P_2P_4 + m_4P_1P_2P_3$, while an element of $C_1(S)$ will look like $m_1P_1P_2 + m_2P_1P_3 + m_3P_1P_4 + m_4P_2P_3 + m_5P_2P_4 + m_6P_3P_4$.

152. Though, again, technically we ought to be working with the more general *cell complexes* and their cellular maps.

Now observe that if σ is an n -simplex, then applying the boundary operator to σ will land us inside the group of $(n-1)$ -chains, that is, $\partial_n(\sigma) \in C_{n-1}(X)$.¹⁵³ Moreover, since $C_n(X)$ is a free abelian group—thus enabling us to describe a *homomorphism* of such a group by specifying its values on generators—it is clear that ∂_n describes a boundary homomorphism mapping $C_n(X)$ into $C_{n-1}(X)$. In other words,

$$\partial_n \left(\sum_i m_i \sigma_i \right) = \sum_i m_i \partial_n(\sigma_i).$$

For instance,

$$\begin{aligned} \partial_1(3P_1P_2 - 4P_1P_3 + 5P_2P_4) &= 3\partial_1(P_1P_2) - 4\partial_1(P_1P_3) + 5\partial_1(P_2P_4) \\ &= 3(P_2 - P_1) - 4(P_3 - P_1) + 5(P_4 - P_2) \\ &= P_1 - 2P_2 - 4P_3 + 5P_4. \end{aligned}$$

But since we have a homomorphism, we are naturally drawn to look at two things: the *kernel* and the *image*. The kernel of ∂_n will consist of those n -chains with boundary zero, and so the elements of the kernel are just n -cycles. We sometimes denote the kernel of ∂_n , the group of n -cycles, by $Z_n(X)$. So for instance, if $q = P_1P_2 + P_2P_3 + P_3P_1$, then $\partial_1(q) = (P_2 - P_1) + (P_3 - P_2) + (P_1 - P_3) = 0$. Note that this q corresponds to a cycle around the triangle with vertices P_1, P_2, P_3 (oriented in the obvious way). Furthermore, we can consider the image under ∂_n , namely the group of $(n-1)$ -boundaries, which consists of those $(n-1)$ -chains that are boundaries of n -chains. We sometimes denote this group $B_{n-1}(X)$.

Homomorphisms compose and it is a well-known fact that the composite homomorphism $\partial_{n-1}\partial_n$ taking $C_n(X)$ into $C_{n-2}(X)$ takes everything into zero, that is, for each $c \in C_n(X)$, we have that $\partial_{n-1}(\partial_n(c)) = 0$, or $\partial^2 = 0$. A corollary of this is that $B_n(X)$ is a *subgroup* of $Z_n(X)$, allowing us to form the quotient or factor group $Z_n(X)/B_n(X)$ which we denote $H_n(X)$ and call the n -dimensional *homology group* of X . While perhaps obvious, it is important to realize that this quotient simply puts an equivalence relation on Z_n with respect to B_n , that is, $\omega \sim \sigma \iff \omega - \sigma \in B_n$, and so is technically represented as some coset.

The important thing to realize now is that we can form the *sequence of chain groups* linked together by such boundary homomorphisms, a sequence we call the *chain complex*:

$$\cdots \xrightarrow{\partial_{n+2}} C_{n+1} \xrightarrow{\partial_{n+1}} C_n \xrightarrow{\partial_n} C_{n-1} \xrightarrow{\partial_{n-1}} \cdots$$

Moreover, if $C = \langle C, \partial \rangle$ is the previous (in principle doubly infinite) sequence of abelian groups together with the collection of homomorphisms satisfying the condition that each map descends by one dimension and that $\partial^2 = 0$, then we can extend all the above reasoning to the sequences themselves and immediately see that under these conditions the image under ∂_k will be a subgroup of the kernel of ∂_{k-1} . In brief, we can define the kernel $Z_k(C)$ of ∂_k as the group of k -cycles, the image $B_k(C) = \partial_{k+1}[C_{k+1}]$ as the group of k -boundaries, and then the factor group $H_k(C) = Z_k(C)/B_k(C) = \ker \partial_k / \text{image } \partial_{k+1}$ as the k -th homology group of C . In other words, H_k gives all the vectors that are annihilated in stage k that were not already present in stage $k+1$. If for all k in a sequence we have that the image under ∂_k

153. For the record, if we define $C_{-1}(X) = \{0\}$, the trivial group of one element, then $\partial_0(\sigma) \in C_{-1}(X)$.

is equal to the kernel of ∂_{k-1} , then we have what is called an *exact sequence*. While exact sequences are chain complexes, the converse is not true, since a chain complex need only satisfy that the image (of the prior map) is contained in the kernel (of the subsequent map). The important thing to realize here is that *homology just measures the difference* between the image and the kernel maps.

For simplicial complexes X and Y , a map f from X to Y induces a mapping (i.e., homomorphism) of homology groups $H_k(X)$ into $H_k(Y)$. This arises if we consider that for certain triangulations of X and Y , the map f will give rise to a homomorphism f_k of $C_k(X)$ into $C_k(Y)$, which moreover commutes with ∂_k , that is, $\partial_k f_k = f_{k-1} \partial_k$.

We can dualize this entire account to get an account of *cohomology*, and we do so briefly now since it will be important for what follows. Consider a simplicial complex X . For an oriented n -simplex σ of X , we can define the *coboundary* $\delta^n(\sigma)$ of σ as the $(n+1)$ -chain summing up all the $(n+1)$ -simplices τ that have σ as a face. In other words, we are summing those τ that have σ as a summand of $\partial_{n+1}(\tau)$. For instance, if we let X be the simplicial complex consisting of the solid tetrahedron, then $\delta^2(P_3 P_2 P_4) = P_1 P_3 P_2 P_4$, while $\delta^1(P_3 P_2) = P_1 P_3 P_2 + P_4 P_3 P_2$.

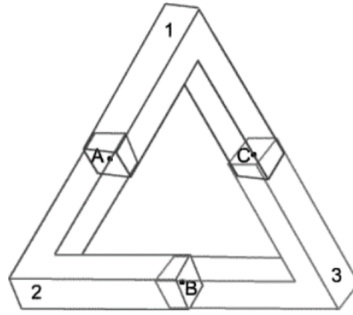
We can also define the group $C^n(X)$ of n -cochains as the same as the group $C_n(X)$. However, the coboundary maps δ^n go the other way from the boundary maps, that is, we have $\delta^n : C^n \rightarrow C^{n+1}$, defined by

$$\delta^n \left(\sum_i m_i \sigma_i \right) = \sum_i m_i \delta^n(\sigma_i).$$

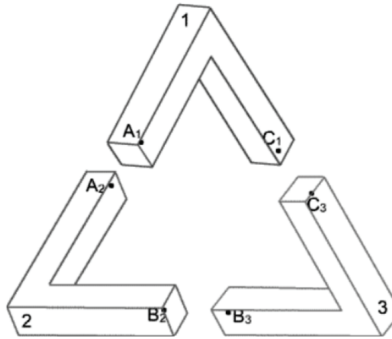
We can then build up sequences of cochain groups into cochain complexes, just as one would expect. Cochain complexes in general can be thought of as looking at how objects are related to larger superstructures instead of to smaller substructures (as was the case for chain complexes). Just as before, we could show that $\delta\delta=0$, that is, that $\delta^{n+1}(\delta^n(c))=0$ for each $c \in C^n(X)$. Moreover, we can define the group $Z^n(X)$ of n -cocycles of X as the kernel of the coboundary homomorphism δ^n , the group B^n of n -coboundaries of $C^n(X)$ as the image of δ^{n-1} , that is, $\delta^{n-1}[C^{n-1}(X)]$; and since we have that $\delta\delta=0$, again $B^n(X)$ will be a subgroup of $Z^n(X)$. This last fact allows us to define the n -dimensional *cohomology group* $H^n(X)$ of X as $Z^n(X)/B^n(X)$, that is, the kernel of the map “going out” mod the image of the map “coming in.” A cocycle that is not a coboundary is topologically informative—and this difference gives rise to a “nontrivial” cohomology group (or nontrivial cocycle).

Example 231 Let us look a little closer at these ideas through the lens of the Penrose tribar, introduced at the beginning of the chapter. Let Q be the region of the plane on which the tribar is drawn. We can consider Q as being pasted together from three pieces, as in:¹⁵⁴

154. This image, and the next, are modified from Phillips (2021). The exposition that follows is derived from Penrose (1992) and Phillips (2021).



The three pieces overlap and we can imagine having chosen points in each of the three overlapping regions—for example, A in the overlap between piece 1 and piece 2; B in the overlap between piece 2 and piece 3; and C in the overlap between piece 1 and piece 3. Then imagine “pulling apart” the tribar into these components, and relabeling the points appropriately, so that A_1 in piece 1 and A_2 in piece 2 both correspond to A , points B_2, B_3 correspond to B , and C_1, C_3 correspond to C , as in:



We take the tribar to be patched together from the overlapping smaller drawings, where each of these three components is taken to be a perspective drawing of a corresponding object in space. Observe how such a perspective drawing—representing 3-D space in the plane—implies that certain of the points are meant to be further from or closer to the eye E of a viewer than other points. For instance, letting $d(E, A_1)$ denote the implied distance from a viewer E to point A_1 in space, this may differ from $d(E, A_2)$ —in accordance with conventions of perspective where, say, the corner of piece 2 is implied as closer than the corner of piece 1, forcing A_1 to be further from the eye than A_2 . Thus, we can consider the ratio

$$d_{12} = \frac{d(E, A_1)}{d(E, A_2)},$$

which describes piece 1’s implied distance in relation to that of piece 2. Similarly, we can let $d_{13} = \frac{d(E, C_1)}{d(E, C_3)}$ and $d_{23} = \frac{d(E, B_2)}{d(E, B_3)}$. Of course, the other ratios d_{21}, d_{31} , and d_{32} would then be defined by just taking the reciprocal, for example, $d_{21} = \frac{1}{d_{12}}$. Observe that these values d_{ij} do not depend on the chosen points—since any two points in the relevant pieces would yield the same value for the ratio.

While the drawing of each piece is a consistent rendering of a 3-D structure, there is some “ambiguity” in any interpretation of the perspective drawing of the spatial objects:

the distance of the object depicted, in relation to the viewer's eye, is not known. The implied object could be three times as big and three times further away and appear the same as the original. This perceived distance can accordingly be tracked with positive real numbers in $\mathbb{R}^+$, and we can use the factors λ to track how the perceived distances change, for example, if the perceived distance of piece 1 is changed by a factor of λ then d_{12} and d_{13} will have to be multiplied by λ , and d_{21} and d_{31} divided by λ . Clearly, if the hypothetical tribar (or any other drawing of a figure) could indeed be consistently realized in 3-D space—that is, if the pieces 1, 2, and 3 could indeed be fused into a viable 3-D object—then we would expect that we could effectively rescale each of the three pieces by factors $\lambda_1, \lambda_2, \lambda_3$ until they came together into one consistent structure. And their “coming together into one consistent structure” would be captured by rescalings that make each of the $d_{ij} = 1$. In other words, we would be able to find three positive real numbers such that $(\lambda_i/\lambda_j)d_{ij} = 1$ for each different i, j , or (what is the same)

$$d_{12} = \frac{\lambda_2}{\lambda_1}, d_{13} = \frac{\lambda_3}{\lambda_1}, d_{23} = \frac{\lambda_3}{\lambda_2}.$$

On the other hand, if no such factors $\lambda_1, \lambda_2, \lambda_3$ exist, then the hypothetical object would be impossible to realize in 3-D space.

This situation, and the decisive impossibility, can be codified in the language of cohomology, specifically attending to the first cohomology group. First, let us consider some terminology. Here, a collection $\{d_{ij}\}$ will be a *cocycle*. If it respects the equations shown above, then the cocycle is a *coboundary*. Two cocycles are regarded as equivalent if they can be converted to one another using a rescaling factor λ , that is, when we shift the distance from which object Q_i is being viewed by replacing the pair (d_{ij}, d_{ik}) with $(\lambda d_{ij}, \lambda d_{ik})$ for some positive real number λ . Under such an equivalence, we are left with the elements of the *cohomology group* $H^1(Q, \mathbb{R}^+)$, the quotient group of 1-cocycles modulo 1-coboundaries. The trivial “unit” element of $H^1(Q, \mathbb{R}^+)$ is thus given by the coboundaries—so, checking whether the figure depicted in Q is an impossible figure amounts to checking whether the element of $H^1(Q, \mathbb{R}^+)$ is the unit.

More explicitly, our object Q is the union of $n=3$ subsets—given by the three parts Q_1, Q_2, Q_3 —where, technically, each of the Q_i is topologically a solid ball, as are the intersections $Q_{ij} = Q_i \cap Q_j$. With our setup, a 0-dimensional cochain (taking values in $\mathbb{R}^+$) assigns a number $\lambda_i \in \mathbb{R}^+$ to each of the Q_i . A 1-dimensional cochain assigns a number (say, a_{ij}) to each of the Q_{ij} , where $a_{ji} = 1/a_{ij}$, allowing us to take account of orientation and treat Q_{ij} and Q_{ji} as different (even though they correspond to the same set). A two-dimensional cochain would assign a number to each threefold intersection—yet, in this example, there are no such intersections. Supposing some collection of λ_i is our 0-cochain, its coboundary $\delta\lambda$ just assigns to each Q_{ij} the number λ_j/λ_i . If Q had threefold intersections, then the coboundary of the 1-cochain $\{a_{ij}\}$ would be the 2-cochain δa that assigns the number $\delta a_{ijk} = a_{ij}a_{jk}/a_{ik}$ to Q_{ijk} . Observe that we would have $\delta\delta = 1$, which 1 is the unit in the group $\mathbb{R}^+$ of coefficients.

As before, a cochain with coboundary equal to 1 (the unit) is called a *cocycle*. Of course, if there happen to be no cochains in the next higher dimension—as is the case with 1-cochains of the tribar—then every cochain is automatically a cocycle. A cocycle that is *not* a coboundary is said to be “nontrivial”—and it is these cocycles that are of interest.

In particular, if a drawing is of an impossible object, then decomposing it into possible pieces, selecting points in the intersections of those pieces, and defining the numbers d_{ij} will yield a nontrivial cocycle, that is, one that is not a coboundary. In looking at the first cohomology group of Q , $H^1(Q, \mathbb{R}^+)$, we would ordinarily really be working with the cohomology of the cover $\mathcal{U} = \{Q_1, Q_2, Q_3, Q_{12}, Q_{13}, Q_{23}\}$ —however, since each of these elements of the cover is topologically equivalent to a solid ball, we are able to work directly with Q .

In this construction, there is an implied evaluation map ϵ that takes a 1-cochain d to the product $d_{12}d_{23}d_{31}$, which will be an element of $\mathbb{R}^+$. A 1-cochain d will then be a coboundary iff $\epsilon(d) = 1$. To see this, observe that if there were a 0-chain $\lambda = \{\lambda_1, \lambda_2, \lambda_3\}$ with $d = \delta\lambda$ (i.e., such that $d_{ij} = \frac{\lambda_j}{\lambda_i}$), then all the factors of λ in the product $d_{12}d_{23}d_{31}$ would cancel; conversely, if $d_{12}d_{23}d_{31} = 1$, then one could just define $\lambda_1 = 1, \lambda_2 = d_{12}, \lambda_3 = d_{13}$ and then get $d = \delta\lambda$. It would remain to check that $d_{23} = \frac{\lambda_3}{\lambda_2}$, but since $\epsilon(d) = 1$, this would be $d_{13}/d_{12} = 1/(d_{31}d_{12}) = d_{23}$.

Considering the tribar, then, we would just calculate

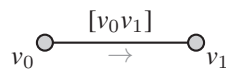
$$\epsilon(d) = \frac{d(E, A_1)}{d(E, A_2)} \frac{d(E, B_2)}{d(E, B_3)} \frac{d(E, C_3)}{d(E, C_1)}.$$

Considering pairs of the three pieces, and adopting a counterclockwise orientation (starting at A), observe that ordinary perspective implies A_1 is further from the eye than A_2 —making the first factor $\frac{d(E, A_1)}{d(E, A_2)}$ greater than 1. B_2 and B_3 , for their part, are viewed as being the same distance away, making the factor $\frac{d(E, B_2)}{d(E, B_3)} = 1$. Finally, C_3 is regarded as further from the eye than C_1 , making the last factor $\frac{d(E, C_3)}{d(E, C_1)}$ greater than 1. Altogether, we have one factor equal to 1 and the other two factors greater than 1—thus, whatever their particular values, their product $\epsilon(d) \neq 1$. In other words, there is no way of rescaling the distances of the pieces to make them all fit together to assemble into the tribar. This impossibility is codified by the nontriviality of the above cocycle d , that is, by the nontriviality of the first cohomology group H^1 .

(Co)chain complexes can be thought of as representing a cell complex within the context of linear algebra, expressing the action of taking the boundary of a cell in terms of a linear transformation. If we put this latter approach together with the development of ASCs from before, we get what is usually called simplicial (co)homology. We can thus start with some arbitrary ASC X or its realization and turn it into a chain complex (for instance). This then allows us to compute its homology. Note that what we are doing here is moving *via functors* from **Top** (after having placed the appropriate topology on X) to the category of chain complexes **Chn**, finally landing in the category **Vect**. The idea here is that we can use the algebraic properties exhibited by the composite of functors $H_\bullet$ to view the topological properties of the original ASC. Initially, the C_k and their maps may be vector spaces over some field like $\mathbb{F}_2$ (the field of two elements, 0 and 1), and so the simplicial homology of X , denoted $H_\bullet(X; \mathbb{F}_2)$ will take coefficients in $\mathbb{F}_2$ (“on” or “off”). But we could also let (co)homology take coefficients elsewhere—for instance, as we saw a moment ago, in $\mathbb{R}$ or $\mathbb{R}^+$ (thereby describing simplices’ intensities, say, as opposed to the simple “on-off” of $\mathbb{F}_2$). Eventually, the idea here is to abstract further and let (co)homology take *sheaves as coefficients*.

When we start computing simplicial homology, we must choose and fix an ordering on the list of vertices in each simplex. We use coefficients other than $\mathbb{F}_2$, like $\mathbb{R}$, and the boundary maps will be used to track orientation.¹⁵⁵ In brief, for an arbitrary simplicial complex X , $C_k(X)$ will be a vector space whose dimension is the number of k -simplices of X , that is, the vector space whose basis (each row as one of the k -simplices, with some coefficients) is the list (given some fixed ordering) of k -simplices of X . The boundary maps $\partial_k : C_k(X) \rightarrow C_{k-1}(X)$ go down in dimension, from the chain space built of the k -dimensional simplices to the chain space built of the $k-1$ simplices. In other words, the map acts on the ordered set $[v_0, \dots, v_k]$ and should be a linear map. Indeed it takes some linear combination of k -simplices and returns some linear combination of $(k-1)$ -simplices.

As a very simple example illustrating these ideas, consider the following ASC $X = \{[v_0], [v_1], [v_0v_1]\}$, realized geometrically as



Then the chain complex associated to this is given by

$$C_2 \xrightarrow{\partial_2} C_1 \xrightarrow{\partial_1} C_0 \xrightarrow{\partial_0} C_{-1}$$

$$0 \xrightarrow{\partial_2} \mathbb{R} \xrightarrow{\partial_1 = \begin{pmatrix} -1 \\ 1 \end{pmatrix}} \mathbb{R}^2 \xrightarrow{\partial_0} 0.$$

By inspection, one can see that $H_2(X)$ (and all higher homology groups) must be zero (the trivial group). Moreover, the kernel of the ∂_1 map is trivial, and the image of ∂_2 (a 1×0 matrix map) is zero, so $H_1(X)$ must be zero. As for H_0 , on the other hand, by inspection we observe that the kernel of ∂_0 is everything, that is, $\mathbb{R}^2$. The image of ∂_1 is spanned by the $\partial_1 = \begin{pmatrix} -1 \\ 1 \end{pmatrix}$ matrix. Taking the quotient, then, we see that the dimension of H_0 must be 1. Strictly speaking, H_0 is a coset, and has a vector space that is isomorphic to a single copy of $\mathbb{R}$. Thinking of it in terms of its coset representation, parameterized by one free parameter, one can think of this as *identifying* the two points on account of the fact that they happen to be connected via an edge. As H_0 effectively represents the connected components, this should make good sense. Since the dimension of $H_1(X)$ can be thought of as picking out the number of *cycles* in the graph (i.e, among the vertices and edges) that are not filled in by two-dimensional simplices, it should also be intuitively clear that H_1 ought to be trivial in this example. If we had found an H_1 not equal to zero, say for another simplex, this might be indicating that the simplices all fit together in some fashion but that they cannot be glued together into one big construct, on account of some kind of obstruction or “hole.” The dimension of $H_k(X)$ is usually referred to as k -cycles, and for nontrivial values this can be thought of as picking out $(k+1)$ -dimensional “voids” or holes.¹⁵⁶

155. Really, though, we would like to generalize beyond *field* coefficients, say to $\mathbb{Z}$ coefficients, making each C_k a $\mathbb{Z}$ -module. In this case, the chains will record finite collections of simplices with some orientation and multiplicity. We then move to a chain complex over an R -module, where R is a ring, the boundary maps being module homomorphisms. Then we have a sufficiently general definition: a chain complex $\mathcal{C} = (C_\bullet, \partial)$ is any sequence of R -modules C_k with homomorphisms $\partial_k : C_k \rightarrow C_{k-1}$ satisfying $\partial_k \circ \partial_{k+1} = 0$.

156. We could further develop this story in a number of directions, for instance going on to define the Betti numbers, which indicate various levels of obstructions, stringing them together as k varies; but we leave the curious reader to pursue these matters on their own.

We now return to simplicial or cellular maps.¹⁵⁷ Just as X can be unfolded into a chain complex $C_\bullet(X)$, cellular maps $f: X \rightarrow Y$ between cell complexes X and Y can be unfolded to yield a sequence $f_\bullet$ of homomorphisms $C_k(X) \rightarrow C_k(Y)$. Since f is continuous, it induces a *chain map* $f_\bullet$ that “plays nicely” with the boundary maps of $C_\bullet(X)$ and $C_\bullet(Y)$, meaning that the following diagram is made to commute:

$$\begin{array}{ccccccc} \cdots & \longrightarrow & C_{n+1}(X) & \xrightarrow{\partial} & C_n(X) & \xrightarrow{\partial} & C_{n-1}(X) \xrightarrow{\partial} \cdots \\ & & \downarrow f_\bullet & & \downarrow f_\bullet & & \downarrow f_\bullet \\ \cdots & \longrightarrow & C_{n+1}(Y) & \xrightarrow{\partial'} & C_n(Y) & \xrightarrow{\partial'} & C_{n-1}(Y) \xrightarrow{\partial'} \cdots \end{array}$$

Since the squares in this diagram commute, meaning the chain map respects the boundary operator, we have that f will act not just on chains but on cycles and boundaries as well, which entails that it induces the homomorphism $H(f): H_\bullet(X) \rightarrow H_\bullet(Y)$ on homology. We are dealing with functors! In particular, the functoriality of homology means that the induced homomorphisms will reflect, algebraically, the underlying properties of continuous maps between spaces. Chain maps, in respecting the boundary operators, send neighbors to neighbors, and thus capture an essential feature of the underlying continuous maps. Via such functors, we can build up big “quiver” diagrams with composable maps between chain complexes, that is, between one complex representation and the next.

Overall, the idea is that homology should algebraically capture changes that happen to the underlying complex structure. Homology proceeds by first replacing topological spaces with complexes of algebraic objects. It then takes a hierarchically ordered sequence of these parts “chained together,” that is, a chain complex, as input and returns the global features. Then other topological concepts—like continuous functions and homeomorphisms—have analogues at the level of chain complexes. *Cohomology*, for its part, is just the homology of the cochain complex.

In terms of the big picture, then, homology can be seen as a way of translating topological problems into algebraic ones. Specifically, it will (in the general approach) translate a topological problem into a problem about modules over commutative rings; but when we can take coefficients in a *field*, this actually amounts to a translation of the topological problem into one of linear algebra. And this is one of the principal motivations! Linear algebra is generally much simpler than topology, in part because of how dimension can classify finite-dimensional vector spaces (thus the centrality of “dimension formula” that relates the kernel of a linear transformation to its image). The “long exact sequences” we started to look at are basically fancy versions of the dimension formula. In the context of sheaf theory, we can develop powerful ways of building long exact sequences of cohomology spaces from short exact sequences of sheaves, where such sequences can already tell us a lot on their own.

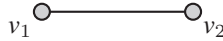
9.2.2 Cohomology with Sheaves

We now proceed, at last, to the construction of the cochain complex for a sheaf F , where each C^k will comprise the stalks, stacked together, over the k -simplices, and the coboundary maps (denoted with δ^k) are built by gluing together a bunch of restriction maps. Once we

157. We ignore more sophisticated issues here, like maps between different dimensions.

have a sheaf, the tactic for computing all sections at once is to build a chain complex (where we go up in dimension, so that really we have a *cochain* complex) and examine its zero-th homology. The difference between the (co)homology of spaces and sheaf (co)homology mostly just has to do with the fact that with sheaves we are again looking at functions on a space, but the range of these functions is allowed to vary, that is, we will have a collection of possible outputs, where the output space of the functions can change as we move around the domain space.

The basic idea here can be nicely explained as follows.¹⁵⁸ Suppose we have a simplicial complex, say, for simplicity



with the attachment diagram (displayed on top) and the corresponding data of the sheaf (below):

$$\begin{array}{ccccc} v_1 & \longrightarrow & e & \longleftarrow & v_2 \\ F(v_1) & \xrightarrow{F(v_1 \rightsquigarrow e)} & F(e) & \xleftarrow{F(v_2 \rightsquigarrow e)} & F(v_2) \end{array}.$$

Now suppose that s is a global section of this sheaf. Then obviously we must have

$$F(v_1 \rightsquigarrow e)s(v_1) = s(e) = F(v_2 \rightsquigarrow e)s(v_2),$$

where $F(v_i \rightsquigarrow e)$ is the restriction map, $s(v_1)$ is a section belonging to $F(v_1)$, and $s(v_2)$ is a section belonging to $F(v_2)$. But now we need only observe that this equation in fact holds in a vector space, which means that we can rewrite it $F(v_1 \rightsquigarrow e)s(v_1) - F(v_2 \rightsquigarrow e)s(v_2) = 0$, or in matrix form:

$$\left(\begin{array}{c|c} +F(v_1 \rightsquigarrow e) & -F(v_2 \rightsquigarrow e) \end{array} \right) \begin{pmatrix} s(v_1) \\ s(v_2) \end{pmatrix} = 0.$$

Note that these $F(v_i \rightsquigarrow e)$ entries are just the restriction maps, so in the context of sheaves of vector spaces, they will in general be (potentially large) matrices the entries of which will be given by the linear (restriction) maps.¹⁵⁹

Extending this reasoning to arbitrary simplicial complexes (which we assume comes with a listing of vertices in a particular order, say lexicographic for concreteness), we can observe that computing the space of global sections of a sheaf is equivalent to computing the kernel of a particular matrix. Moreover, the matrix

$$\left(\begin{array}{c|c} +F(v_1 \rightsquigarrow e) & -F(v_2 \rightsquigarrow e) \end{array} \right)$$

generalizes into the coboundary map

$$\delta^k : C^k(X; F) \rightarrow C^{k+1}(X; F),$$

158. The next two paragraphs closely follow Robinson (2014).

159. Note that, in the context of the (co)homology groups defined earlier in terms of equivalence relations and cosets, saying that the difference of the two restriction maps is equal to their value along the edge (i.e., $s(e)$), is effectively the same as saying that their difference can be regarded as *zero*.

which takes an assignment s on the k -faces to another assignment $\delta^k(s)$ whose value at a $(k+1)$ -face b is

$$(\delta^k(s))(b) = \sum_{\text{all } k\text{-faces } a \text{ of } X} [b : a] F(a \rightsquigarrow b) s(a),$$

where $[b : a]$ is defined to be 0 if a (a k -simplex) is not a face of b (a $(k+1)$ -simplex) and $(-1)^n$ if you have to delete the n -th vertex of b to get back a .¹⁶⁰ This makes sense since a and b must differ by one dimension, so either a is not a face of b , or a is a face of b , in which case they will differ by exactly one vertex (i.e., one need only delete one of the vertices of b to get a). The sign mechanism tied to the deleted vertex allows us to track and respect the chosen orientation given to the complex.

The matrix kernel construction enables us to extend this approach to higher dimensions and over simplices that are arbitrarily larger. Then we can show that we have a cellular cochain complex for a sheaf F on some simplicial complex X . To form the cochain spaces $C^k(X; F)$, we just collect the stalks over vertices (the domain) and edges (the codomain) together via direct sum, so that an element of $C^k(X; F)$ comes from the stalk at each k -simplex:

$$C^k(X; F) = \bigoplus F(a),$$

where a is a k -simplex. Moreover, the coboundary map $\delta^k : C^k(X; F) \rightarrow C^{k+1}(X; F)$ is defined as the block matrix where for row i and column j , the (i, j) -th entry is given by $[b_i : a_j] F(a_j \rightsquigarrow b_i)$, where the $[b_i : a_j]$ term is either 0, +1, or -1, depending on the relative orientation of a_j and b_i , assuming one is a face of the other. Carrying on in this way, we have defined the cellular cochain complex:

$$\dots \xrightarrow{\delta^{k-2}} C^{k-1}(X; F) \xrightarrow{\delta^{k-1}} C^k(X; F) \xrightarrow{\delta^k} C^{k+1}(X; F) \xrightarrow{\delta^{k+1}} \dots$$

Using the cellular sheaf cohomology group definition,

$$H^k(X; F) = \ker \delta^k / \text{image } \delta^{k-1},$$

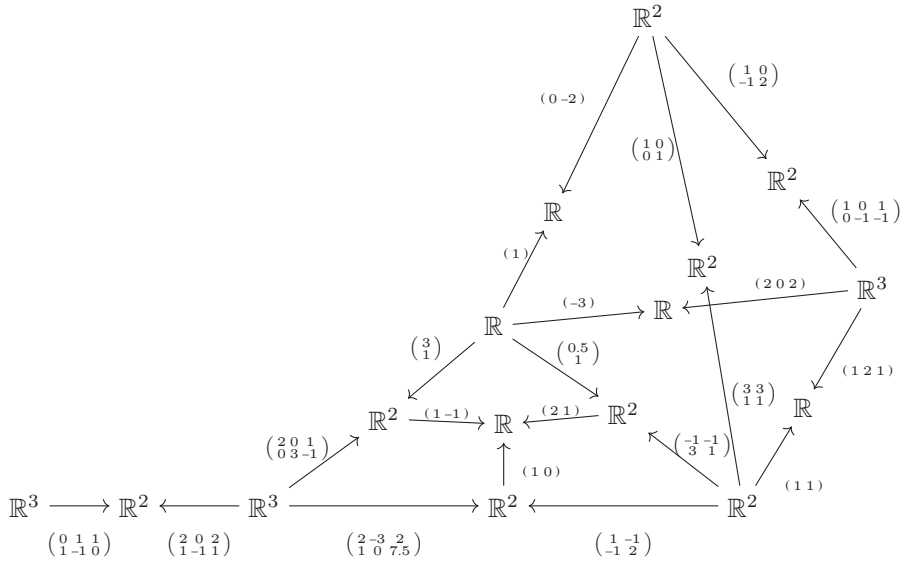
we will recover all the cochains that are *consistent* in dimension k (kernel part) but that did not yet show up in dimension $k-1$ (image part). Notice that $H^0(X; F) = \ker \delta^0$, that is, the zero-th cohomology classes, will be those assignments s that are global sections of the sheaf F . For $k > 0$, the nontrivial elements of H^k will represent some calculable obstruction to globally consistent fusions. For instance, the appearance of nontrivial cohomology in H^1 supplies an explicit measure of the failure of local data to patch.¹⁶¹

All of this is perhaps better illustrated with our running example.

Example 232 We reproduce our running example below for convenience:

¹⁶⁰ We start counting at 0.

¹⁶¹ In this vein, we saw how Penrose (1992) displayed the global impossibility of the locally realistic figure of his tribar in terms of a nontrivial element of first cohomology—that is, the first cohomology group $H^1(Q, \mathbb{R}^+)$ of some region of the plane on which the drawing of the impossible figure is made, with coefficients in $\mathbb{R}^+$. But that example could have also been presented, in fundamentally the same way, by taking coefficient values in the sheaf of positive functions (under pointwise multiplication). For a number of interesting examples and discussion of further aspects of sheaf cohomology, see Curry (2014) and Robinson (2014).



For such a sheaf F over our given complex X , we have the following:

$$\begin{aligned}
 H^0(X; F) &\xrightarrow{\delta^0} H^1(X; F) \xrightarrow{\delta^1} H^2(X; F) \xrightarrow{\delta^2} H^3(X; F) \\
 \mathbb{R}^2 \oplus \mathbb{R}^3 \oplus \mathbb{R}^2 \oplus &\quad \mathbb{R}^2 \oplus \mathbb{R}^2 \oplus \mathbb{R} \oplus \mathbb{R} \oplus \\
 \mathbb{R} \oplus \mathbb{R}^3 \oplus \mathbb{R}^3 &\xrightarrow{\delta^0} \mathbb{R} \oplus \mathbb{R}^2 \oplus \mathbb{R}^2 \oplus \mathbb{R}^2 \oplus \mathbb{R}^2 \xrightarrow{\delta^1} \mathbb{R} \xrightarrow{\delta^2} 0 \\
 \begin{pmatrix} s(a) \\ s(b) \\ s(c) \\ s(d) \\ s(e) \\ s(f) \end{pmatrix} &\xrightarrow{\delta^0} \begin{pmatrix} s(ab) \\ s(ac) \\ s(ad) \\ s(bc) \\ s(bd) \\ s(cd) \\ s(ce) \\ s(de) \\ s(ef) \end{pmatrix} \xrightarrow{\delta^1} (s(cde)) \xrightarrow{\delta^2} 0
 \end{aligned}$$

and where δ^0 is given by

$$\begin{array}{lcl}
 & \begin{array}{cccccc} a & b & c & d & e & f \\ \downarrow & \downarrow & \downarrow & \downarrow & \downarrow & \downarrow \end{array} & \\
 \begin{array}{l} [ab] \rightarrow \\ [ac] \rightarrow \\ [ad] \rightarrow \\ [bc] \rightarrow \\ [bd] \rightarrow \\ [cd] \rightarrow \\ [ce] \rightarrow \\ [de] \rightarrow \\ [ef] \rightarrow \end{array} & \rightarrow & \left[\begin{array}{cccccc} -F(a \rightsquigarrow ab) & F(b \rightsquigarrow ab) & 0 & 0 & 0 & 0 \\ -F(a \rightsquigarrow ac) & 0 & F(c \rightsquigarrow ac) & 0 & 0 & 0 \\ -F(a \rightsquigarrow ad) & 0 & 0 & F(d \rightsquigarrow ad) & 0 & 0 \\ 0 & -F(b \rightsquigarrow bc) & F(c \rightsquigarrow bc) & 0 & 0 & 0 \\ 0 & -F(b \rightsquigarrow bd) & 0 & F(d \rightsquigarrow bc) & 0 & 0 \\ 0 & 0 & -F(c \rightsquigarrow cd) & F(d \rightsquigarrow cd) & 0 & 0 \\ 0 & 0 & -F(c \rightsquigarrow ce) & 0 & F(e \rightsquigarrow ce) & 0 \\ 0 & 0 & 0 & -F(d \rightsquigarrow de) & F(e \rightsquigarrow de) & 0 \\ 0 & 0 & 0 & 0 & -F(e \rightsquigarrow ef) & F(f \rightsquigarrow ef) \end{array} \right]
 \end{array}$$

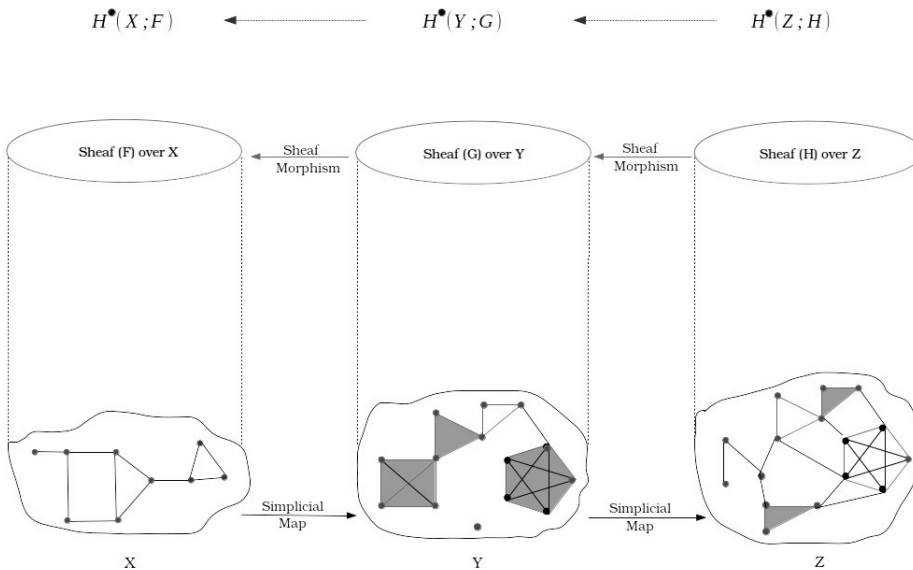
and δ^1 by

$$\begin{array}{lcl}
 & \begin{array}{ccccccccc} [ab] & [ac] & [ad] & [bc] & [bd] & [cd] & [ce] & [de] & [ef] \\ \downarrow & \downarrow & \downarrow & \downarrow & \downarrow & \downarrow & \downarrow & \downarrow & \downarrow \end{array} & \\
 [cde] \rightarrow & \left[\begin{array}{ccccccccc} 0 & 0 & 0 & 0 & 0 & F(cd \rightsquigarrow cde) & -F(ce \rightsquigarrow cde) & F(de \rightsquigarrow cde) & 0 \end{array} \right]
 \end{array}$$

and δ^2 is trivial. Computing the cohomologies here will require doing some linear algebra and row reduction, which is left to the reader. Observe that in general by adding in more higher-dimensional consistency checks—that is, filling in data corresponding to the “missing” higher-dimensional simplices of the underlying simplex—we would be able to reduce the kernel of δ^1 , and thus get closer to reducing the nontrivial H^1 .

It is worth emphasizing a more general truth, namely that the space of global sections of a sheaf F on a cell complex X will be isomorphic to $H^0(X; F)$. Moreover, this H^k (i.e., cohomology with sheaves as coefficients) is a functor, which means that when we have sheaf morphisms between sheaves, they will induce linear maps between their cohomology spaces, allowing us to extend things still further. Ultimately, it is also worth noting that while we have been focused on cellular sheaves, *general sheaf cohomology* can be shown to be isomorphic to cellular sheaf cohomology, so this sort of example is really part of a much more general story. And the vector space version of homological algebra we have used in order to provide a concrete example—while useful for many applications and for building intuition—does not display the full power of these notions. Ultimately one would like to (and could) retell the story using rings, modules, and other categories.

Via cohomology, global (in)compatibilities between pieces of local data can display global qualitative features of the data structure. At a very high level, the general idea of all this is suggested by the following figure, illustrating how we get a comparatively simple algebraic representation of features of spaces:



9.3 Philosophical Pass: Sheaf Cohomology

Box 9.1

Philosophical Pass: Sheaf Cohomology

If the sheaf compatibility conditions require controlled transitions from one local description to another, enabling progressive patching of information over overlapping regions until a unique value assignment emerges over the entire region—something that is captured by the vanishing of the group H^0 (yielding our global sections)—higher (nonvanishing) cohomology groups basically detect and measure (in an algebraic fashion) obstructions to such local patching and consistency relations among various dimensional subsystems. In other words, it can be thought of as measuring (for some cover) how many incompatible (purely local) systems we would have to “discard” in order to be left with only the compatible systems. In this way, sheaf cohomology moreover allows us to examine the relationship between information valid globally and the underlying topology of the space.

A proper discussion of the possible invariants that emerge in cohomology, especially as we ascend in dimension, would require a much longer and more detailed discussion. Instead, here are some very general reflections on the idea of sheaf cohomology. Referencing the high-level figure on the previous page, in forming sheaves (middle level) over the discrete approximations of spaces via their triangulations (bottom level), a more “continuous” perspective is recaptured. However, the nonvanishing cohomology groups (top level) give an algebraic (more “discrete”) representation of something like the *resistance* of certain information (assigned to a part of a space) to integration into a more global system. In short, if the collation condition in the sheaf construction aligns them with continuity in the sense that it ensures smooth passage from the local to the global, cohomology with sheaves is something like its discrete counterpart providing us an algebraic measure of when such local-global passages might be blocked. In this respect, sheaf cohomology could intuitively be thought of as capturing—in a dialectic between the continuous and discrete—the nonglobalizability or nonextendibility of a given information structure in relation to other overlapping structures. Both in its algebraic representation and in this general interpretation, then, the nonvanishing cohomology groups, like H^1 , might be thought of as giving us a picture of just how nonintegrated a system of information over a space may be. On the other hand, vanishing cohomology groups indicate the mutual compatibility or globalizability of local information systems (since they tell us about the global sections). In this way, sheaf cohomology emerges as a tool for representing (algebraically) what might be thought of as the degree of *generality* (or lack thereof) of a given system of locally specified data or interlocking ways of assigning information to a space.

Some years before the invention of sheaf theory, Charles Peirce argued that “continuity is shown by the logic of relations to be nothing but a higher type of that which we know as generality. It is relational generality” (Peirce 1997, 6.190). Such a suggestive, if somewhat cryptic, remark provokes us to take a closer look at the connections between generality and continuity that emerge in the context of sheaves. We know that a sheaf enables a collection of local sections to be patched together uniquely given that they agree (or that there exists a translation system for making them agree) on the intersections. Consider the satellite image mosaic sheaf introduced in example 127 (chapter 5). Recall the way in which the sheaf (collation/gluing) condition ensures a systematic passage from local sections (images of parts of the glacier) to a unique global section (the image of the entire glacier). Where the localizing step of the sheaf construction might be thought of as analytic, decomposing an object into a multitude of individual parts (local), the gluing steps are synthetic in restoring systematic relations between those parts and thereby securing a unique assignment over the entire space

(global). The global sections of such a sheaf should not be thought of as a single (topmost) image but rather as the entire network of component parts welded together via certain compatibility relations or constraints. In this connection, generality can be understood in terms of the systematic passage from the local to the global, a passage that is strictly *relational* in that the action of the component restriction maps is precisely an enforcing of certain relations or mutual constraints between the local sections (that are then built up, along the lines of these relations, into a global section), and these are an ineliminable part of the construction.

9.4 A Glimpse into Cosheaves

In the cellular case, we can easily talk about *sheaf homology* by just reversing the direction of the arrows, technically producing a *cosheaf* (with its “corestriction” maps). Simply by observing that the vector space dual of every corestriction map in a cosheaf produces a sheaf over the poset, we arrive at homology for a cosheaf. More explicitly, this reversal gives us a *cosheaf* $\hat{F}$ of vector spaces on a complex, as in definition 224. As with a sheaf, a cosheaf can come to serve as a system of coefficients for homology that varies as the space varies. However, the globality of the cosheaf sections will be found in the top dimension (unlike how the global sections were built up from the *vertices* in the case of cellular sheaves).

More generally, following Curry (2014), we can define a pre-cosheaf as one might expect:

Definition 233 A *pre-cosheaf* is a functor $\hat{F}: \mathcal{O}(X) \rightarrow \mathbf{D}$.¹⁶² Whenever $V \subseteq U$, the co-restriction (or extension) map for the pre-cosheaf is written as $r_U^V: \hat{F}(V) \rightarrow \hat{F}(U)$.

To construct a cosheaf, we again need the notion of open cover, and a cosheaf will merge the notion of covers with that of data (given by the pre-cosheaf). We can think of an open cover here in terms of a function from the “nerve” construction, which basically acts on covers to produce the ASC consisting only of those finite subsets of the cover whose intersection is nonempty. More formally, if we suppose $\mathcal{U} = \{U_i\}$ is an open cover of U , then we can take the *nerve* of the cover to yield an ASC $N(\mathcal{U})$, which will have for elements the subsets $I = \{i_0, \dots, i_n\}$ for which it holds that $U_I = U_{i_0} \cap \dots \cap U_{i_n} \neq \emptyset$. $N(\mathcal{U})$ is then the category with objects the finite subsets I where $U_I \neq \emptyset$, and unique arrows from I to J whenever $J \subseteq I$. Finite intersections of opens are open, so we get the functors $\iota_{\mathcal{U}}: N(\mathcal{U}) \rightarrow \mathcal{O}(X)$ and $\iota_{\mathcal{U}}^{\text{op}}: N(\mathcal{U})^{\text{op}} \rightarrow \mathcal{O}(X)^{\text{op}}$. In general, as the colimit of a cover $N(\mathcal{U}) \rightarrow \mathcal{O}(X)$ is the union $U = \bigcup_i U_i$, the data we associate to U ought to be expressible as the colimit of data assigned to the nerve. With this expectation in mind, we arrive at the following definition:

Definition 234 $\hat{F}$ is a *cosheaf* on $\mathcal{U}$ if the unique map from the colimit of $\hat{F} \circ \iota_{\mathcal{U}}$ to $\hat{F}(U)$, supplied by

$$\hat{F}[\mathcal{U}] := \lim_{I \in N(\mathcal{U})} \hat{F}(U_I) \rightarrow \hat{F}(U),$$

is in fact an isomorphism. Then $\hat{F}$ is a *cosheaf*, period, if for every open set U and every open cover $\mathcal{U}$ of U , the map $\hat{F}[\mathcal{U}] \rightarrow \hat{F}(U)$ is an isomorphism.

162. Looking ahead to the cosheaf conditions, technically we ought to insist that $\mathbf{D}$ be not just any category, but rather a category with enough (co)limits.

For simplicity, suppose $\mathbf{D} = \mathbf{Set}$ and take a cover $\mathcal{U} = \{U_1, U_2\}$ of U by just two open sets. The sheaf condition would of course stipulate that for two functions or sections $s_1 \in F(U_1), s_2 \in F(U_2)$ to give an element in $U = U_1 \cup U_2$, the sections must agree on the overlap $U_1 \cap U_2$. This constraint serves to pick out the consistent choices of elements over the local sections that can then be glued together into a section over the larger set. With a similar setup—that is, $\mathbf{D} = \mathbf{Set}$ and $U = U_1 \cup U_2$ —the cosheaf condition requires not that we find consistent choices, but rather that we use *quotient objects*. We do indeed still form the union of the two sections, but in the process we identify those elements that would be double-counted on account of coming from the intersection. Formally,

$$\hat{F}(U) \cong \left(\prod_{i=1,2} \hat{F}(U_i) \right) / \sim,$$

where $s_1 \sim s_2$ iff there exists an s_{12} (a section over the intersection) such that $s_1 = r_{U_1}^{U_{12}}(s_{12})$ and $s_2 = r_{U_2}^{U_{12}}(s_{12})$. This makes sense, since in accordance with duality, we would expect that the equalizer definition of the sheaf condition would be converted, in passing to cosheaves, into an underlying coequalizer diagram.

A sheaf is constructed in such a way that the values of its sections on larger sets in the Alexandrov topology will determine values on smaller sets. A cosheaf basically reverses this dependence. While so far we have seen many examples of sheaves in contexts where it makes sense to perform *restrictions* of assignments of data from larger spaces to data over smaller spaces, building up global assignments from the “bottom up” via local pieces, cosheaves may roughly be thought of as instead proceeding “top down,” involving *extensions* of data given over smaller spaces to larger spaces.¹⁶³ While in some sense the paradigmatic example of a sheaf was given by the restriction of continuous functions, the paradigmatic example of a cosheaf might be given by the cosheaf of compactly supported continuous functions where, instead of restricting along inclusions, we extend by zero (in the other direction). Ghrist (2014, chap. 9) gives one nice way of appreciating the duality between sheaves and cosheaves, in terms of an application to sensing problems: in this setting, a sheaf fundamentally amounts to *sensing* (the global sections of the sheaf yielding the sensorium), while a cosheaf arises in terms of what the sensors allow one to *infer* (the global sections of the cosheaf being supplied by constraint satisfactions that are consistent with sensing).

Importantly, while in the cellular context the difference between sheaf and cosheaf is somewhat immaterial, simply a matter of which direction makes the most sense for the framing of the problem, things can be far more subtle in the context of sheaves and cosheaves over opens sets for a continuous domain. In insisting on the more general functorial perspective, allowing us to make use of duality, one might suspect that the differences between sheaves and cosheaves are merely formal and not worth discussing. For instance, a sheaf is a particular functor that commutes with limits in open covers. As one might expect, a cosheaf is a functor that preserves colimits in open covers. However, in more

163. In the context of simplices, because of the reversal of direction involved in the topology given on the face relation construction, such extensions will go from higher-dimensional simplices to lower.

general contexts than the cellular one, especially with open sets coming from a continuous domain, the differences can reflect much more than a preference for direction of arrows.¹⁶⁴

The next example explores a particularly fascinating connection between sheaves and cosheaves in the context of probabilities and Bayes nets.

Example 235 Imagine we are given a set of random variables $X_0, X_1, \dots, X_n$.¹⁶⁵ We can consider the set $P(X_0, X_1, \dots, X_n)$ of all joint probability distributions over these random variables, that is, the nonnegative measures or generalized functions with unit integral. Now, there is a very natural map, one that will be familiar to anyone with some exposure to probability theory:

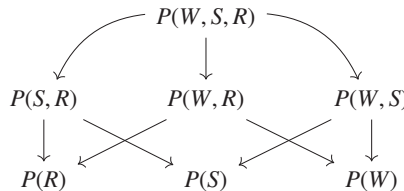
$$P(X_0, X_1, \dots, X_n) \rightarrow P(X_0, X_1, \dots, X_{n-1}).$$

This map is accomplished via *marginalization*, that is, we have

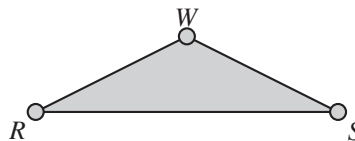
$$f(X_0, X_1, \dots, X_{n-1}) = \int f(X_0, X_1, \dots, X_n) dX_n,$$

and where there are similar maps for marginalizing out the other random variables. The important point is that this in fact yields a cosheaf on the complete n -simplex. We elaborate on this with an example.

Assume given the random variables $X_0 = W$, $X_1 = S$, $X_2 = R$ (the reason for renaming these random variables thus will be clear in a moment). Then the space of probability measures $P(W, S, R)$ is a function from the Cartesian product of the random variables to the (nonnegative) reals such that the integral is zero. Now, we can perform the marginalization operation, for instance, marginalizing out the R by integrating along R , yielding a map $P(W, S, R) \rightarrow P(W, S)$. We can do this for each of the random variables, and continue down in dimension until we reach the measurable functions over a single random variable. In other words, we have:



which in fact represents the attachment diagram of a complete 2-simplex,



164. For instance, when working with the Alexandrov topology on a poset (or when working with locally finite topological spaces), we can ignore the distinction between pre(co)sheaves and (co)sheaves; however, while for general topological spaces, there is a sheafification functor that allows us to pass from a presheaf to the unique smallest sheaf consistent with the given presheaf, there is no analogous cosheafification functor for general topological spaces. For much more on cosheaves, and a number of interesting connections and differences with sheaves, we again refer the reader to Curry (2014). Together with Robinson (2016a), Curry contains more details on some of the dualities in the sheaf-cosheaf perspective, as well as instances of asymmetry (when certain constructions are natural for sheaves but not for cosheaves).

165. The idea for this example was inspired by Robinson (2016a).

The commutative diagram, together with the appropriate marginalization maps, gives a cosheaf on this ASC.

Now, the reader may have already wondered about maps going the other way, namely:

$$P(X_0, X_1, \dots, X_{n-1}) \rightarrow P(X_0, X_1, \dots, X_n),$$

an operation that is parameterized by functions C

$$P(X_0, X_1, \dots, X_n) = C(X_0, X_1, \dots, X_n)f(X_0, X_1, \dots, X_{n-1}),$$

where usually one writes the arguments to C as follows

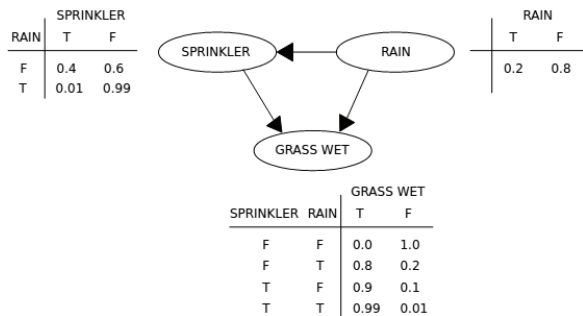
$$C(X_n|X_0, X_1, \dots, X_{n-1}).$$

The reader may recognize that we are just describing *conditional probabilities* and Bayes’s rule. The key observation is that such conditional probability maps yield a sheaf over a portion of the n -simplex.

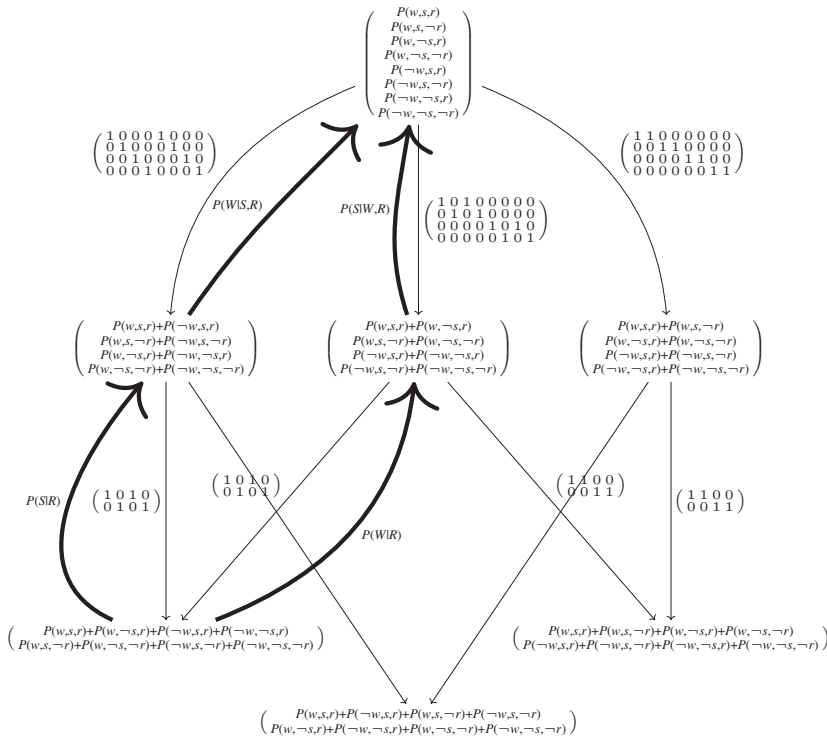
The reader may be familiar with the construction of a *Bayes net*, given by a directed acyclic graph (with an induced topology) together with local conditional probabilities. A Bayes net encodes joint distributions and does so as a product of local conditional distributions, that is,

$$P(X_1, X_2, \dots, X_n) = \prod_{i=1}^n P(X_i | \text{parents}(X_i)).$$

For the sake of concreteness, we consider the following very simple example of a Bayes net, the standard one given in most introductions to the device, involving probabilities of “grass being wet” given that it rained, for instance, or that the sprinkler was running, illustrated with some sample probability assignments:



The perhaps surprising result is that the content of this Bayes net is in fact entirely captured by the paired sheaf-cosheaf construction given below, where the marginalization cosheaf is given by the entire diagram (arrows going down the page), while the paired conditional probability sheaf is given in bold (going up the page) over a part of the underlying attachment diagram.



The paired sheaf-cosheaf construction contains all the data of a Bayes net; and, in fact, a solution to the Bayes net is just a global section that is a section for *both the sheaf and the cosheaf*!

Before ending this chapter and moving into discussion of toposes, we take the opportunity to make an important but frequently overlooked observation. Consider the (co)sheaf construction above. Now, consider that relatively small Bayes nets, say, one with only eight or nine nodes, are already rather simple compared to those that will be of use in practice. One might thus be suspicious of just how complicated the corresponding (co)sheaf might look for even only slightly more involved examples, not to mention the issue of storing the relevant sections for such sheaves. This indeed seems to be a real issue. One might also suspect that computing the sheaf cohomology (and global sections) on “monster” (extremely large) sheaves would be extremely difficult. As the discussion of this Bayes net sheaf-cosheaf construction suggests, most real-life (co)sheaves one would meet in practice may very well turn out to be monsters in the sense of being so large as to cause difficulties in storage, representation, or computation—difficulties we have simply avoided by confining our attention to more modest constructions. This is an issue that deserves to be recognized and pondered.¹⁶⁶

166. For some ways to reduce the difficulty, in one setting, see Smith (2014, 527–533); see Curry (2014, 64–66) for some ways to think about preprocessing the input data so as to deal with the “too many sections” problem. The reader may also find Curry, Ghrist, and Nanda (2015) highly relevant, as this shows how you can “collapse” the data structure if your restriction maps are nice enough. Thanks to Michael Robinson (personal correspondence) for pointing me in the direction of this paper and observing the connection.