

MA3403 Algebraic Topology

Lecturer: Gereon Quick

Lecture 23

23. CLASSIFICATION OF SURFACES

We will first continue the study of bilinear forms, and then use this knowledge to classify all compact connected surfaces, i.e., compact connected two-dimensional manifolds. Then we are going to contemplate a bit more on Poincaré duality. In this lecture, all vector spaces, homology and cohomology groups will be over $\mathbb{F}_2$.

The monoid of nondegenerate symmetric bilinear forms

Last time we showed:

Proposition: Classification of nondegenerate forms

Any finite-dimensional nondegenerate symmetric bilinear form over $\mathbb{F}_2$ splits as an orthogonal sum of forms with matrices

$$I = (1) \text{ and } H = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}.$$

Now let **Bil** be the set of isomorphism classes of nondegenerate symmetric bilinear forms over $\mathbb{F}_2$. This is a **commutative monoid** under the operation of taking orthogonal direct sums (that means it is like a group except that there are no inverses).

Since any such form corresponds to a matrix, we can identify **Bil** also with the **set of invertible symmetric matrices modulo the equivalence relation of similarity**:

- Two matrices A and B are called **similar**, denoted $A \sim B$, if $B = PAP^T$ for some invertible matrix P .
- Every form corresponds to a matrix A determined by $v \cdot w = v^T A w$.
- Assume we have given two vector spaces V_1 and V_2 with nondegenerate symmetric bilinear forms which are represented by matrices A_1 and A_2 , respectively. Then there is an isomorphism $\varphi: V \xrightarrow{\cong} W$ such that

$$\varphi(v \cdot_V w) = \varphi(v) \cdot_W \varphi(w),$$

if and only if A_1 and A_2 are similar.

Hence, in order to understand **Bil** we can aim to understand invertible matrices modulo similarity. Here is a crucial fact:

Lemma

Over $\mathbb{F}_2$ we have the similarity

$$\begin{pmatrix} 0 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix} \sim \begin{pmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{pmatrix}.$$

Proof: The assertion is equivalent to saying there is an invertible matrix P such that

$$\begin{pmatrix} 0 & 1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix} = PP^T.$$

This is the case for $P = \begin{pmatrix} 1 & 1 & 0 \\ 1 & 0 & 1 \\ 1 & 1 & 1 \end{pmatrix}$ where we need to remember that we work over $\mathbb{F}_2$. **QED**

Since neither nI nor mH are similar to other matrices, $I + H = 3I$ is the only relation. As a consequence we get:

Bilinear forms via generators and relations

The commutative monoid **Bil** is generated by I and H modulo the relation

$$I + H = 3I.$$

Now we are going to apply this knowledge to the intersection form.

Intersection form

Let us look at a particular case of Poincaré duality. Let us assume that M is **even-dimensional**, say of dimension $n = 2p$. Then Poincaré duality defines a **symmetric bilinear form** on the $\mathbb{F}_2$ -vector space $H^p(M)$:

$$H^p(M) \otimes_{\mathbb{F}_2} H^p(M) \rightarrow H^p(M).$$

As we observed last time, this can be interpreted as a bilinear form on homology $H_p(M)$. Recall that evaluating this form can be viewed as describing

(modulo 2) the number of points where two p -cycles intersect, after they have put moved in general position, i.e., a position where they intersect transversally.

Intersection form

For a compact manifold of dimension $n = 2p$, the intersection pairing

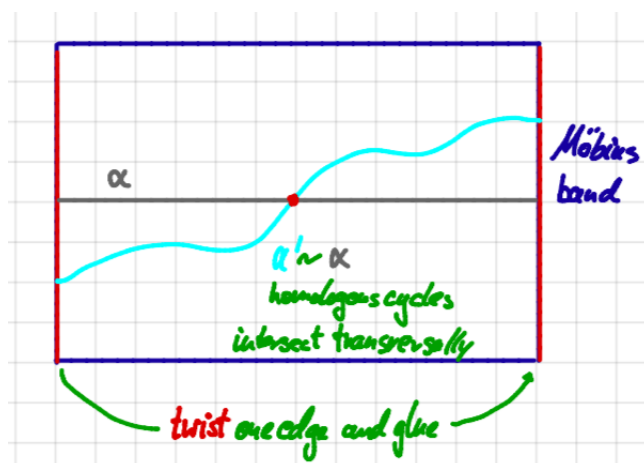
$$H_p(M; \mathbb{F}_2) \otimes H_p(M; \mathbb{F}_2) \rightarrow \mathbb{F}_2, \alpha \cdot \beta := \langle a \cup b, [M] \rangle$$

defines a nondegenerate symmetric bilinear form on $H_p(M; \mathbb{F}_2)$, called the **intersection form**. Here a and b are Poincaré dual to α and β , respectively.

- We have seen the **example of the torus** for which the intersection form is **hyperbolic**, i.e., can be described in terms of the basis α and β by the matrix

$$H = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}.$$

- For another **example**, take $M = \mathbb{RP}^2$. We know $H_1(\mathbb{RP}^2) = \mathbb{F}_2$. Moreover, $\mathbb{RP}^2$ can be viewed as a **Möbius band with a disk glued along the boundary**. On the Möbius band, there is a nontrivial intersection. Hence the **intersection form is nontrivial** and therefore given by I according to our classification, since on a one-dimensional space there only options. As a consequence we see that in whatever way try to move the boundary of the Möbius band in $\mathbb{RP}^2$, it will always **intersect itself in an odd number of points**.



Note that the **open Möbius band** itself is a two-dimensional manifold, but it is **not compact**. While the **closed** Möbius is compact, it is not a manifold

according to the definition we stated last time. Though it is a **manifold with boundary**. The story is different if we allow boundaries.

Connected sums

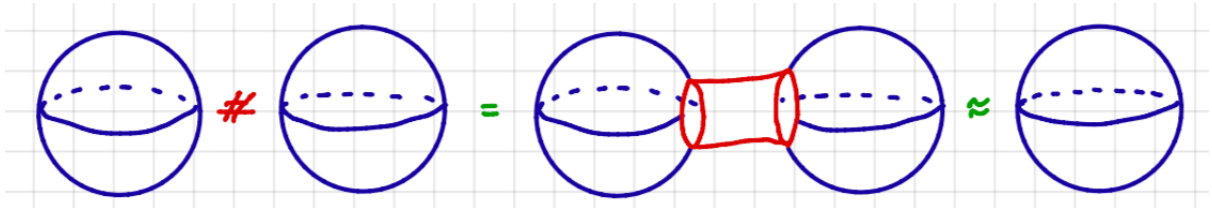
There is an interesting **geometric operation** on manifolds which **produces new ones out of old**:

Given two compact connected manifolds M_1 and M_2 both of dimension n . Then we can

- cut out a small open n -dimensional disk D^n of each one, and
- sew them together along the resulting boundary spheres S^{n-1} , i.e., identify the boundaries via a homeomorphism.
- The resulting space is called the **connected sum** of M_1 and M_2 and is denoted by $M_1 \# M_2$. Note $M_1 \# M_2$ is a **connected compact n -dimensional manifold**.

Let us see two **examples**:

- There is not much happening if we take $S^2 \# S^2$ as it is homeomorphic to S^2 :



- But we get a new surface for $T^2 \# T^2$:



Lemma: Homology of connected sums

There is an isomorphism

$$H_i(M_1 \# M_2) \cong H_i(M_1) \oplus H_i(M_2) \text{ for all } 0 < i < n.$$

Proof: We start with the pair $(M_1 \# M_2, S^{n-1})$. Since M_1 and M_2 are manifolds of dimension n , there is an open neighborhood around S^{n-1} in $M_1 \# M_2$ which retracts onto S^{n-1} . Thus, by a result we showed some time ago when we discussed cell complexes and wedge sums, we know

$$H_*(M_1 \# M_2, S^{n-1}) \cong H_*((M_1 \# M_2)/S^{n-1}, \text{pt}) \cong \tilde{H}_*(M_1 \vee M_2).$$

Now we consider the long exact sequence of the pair $(M_1 \# M_2, S^{n-1})$:

$$\cdots \rightarrow \tilde{H}_i(S^{n-1}) \rightarrow H_i(M_1 \# M_2) \rightarrow H_i(M_1 \# M_2, S^{n-1}) \rightarrow \tilde{H}_{i-1}(S^{n-1}) \rightarrow \cdots$$

Since only $\tilde{H}_{n-1}(S^{n-1})$ is nonzero, we deduce

$$\tilde{H}_i(M_1 \# M_2) \cong H_i(M_1 \# M_2, S^{n-1}) \cong \tilde{H}_*(M_1 \vee M_2) \text{ for all } i < n - 1.$$

Hence, for $0 < i < n - 1$, we have

$$H_i(M_1 \# M_2) \cong H_i(M_1) \oplus H_i(M_2).$$

The remaining part of the long exact sequence is then

$$0 \rightarrow H_n(M_1 \# M_2) \rightarrow H_n(M_1 \vee M_2) \rightarrow H_{n-1}(S^{n-1}) \rightarrow H_{n-1}(M_1 \# M_2) \rightarrow H_{n-1}(M_1 \vee M_2) \rightarrow 0$$

where the zeros on both ends are explained by the vanishing of the corresponding homologies of S^{n-1} .

Since fundamental classes are natural, the map

$$(1) \quad H_n(M_1) \oplus H_n(M_2) \xrightarrow{\cong} H_n(M_1 \vee M_2) \rightarrow H_{n-1}(S^{n-1})$$

sends the fundamental classes of both M_1 and M_2 to the fundamental class of S^{n-1} . Thus, this map is surjective and we deduce from the exactness of the sequence that

$$H_{n-1}(M_1 \# M_2) \cong H_{n-1}(M_1) \oplus H_{n-1}(M_2).$$

We also see that $H_n(M_1 \# M_2)$ is the kernel of the map in (1).

QED

Lemma: Connected sums and intersection forms

Assume both M_1 and M_2 are of dimension $n = 2p$. Then the isomorphism

$$H_p(M_1 \# M_2) \xrightarrow{\cong} H_p(M_1) \oplus H_p(M_2)$$

is **compatible with the intersection form**.

Proof: Fundamental classes are natural in the sense that the homomorphism

$$H_n(M_1 \# M_2) \xrightarrow{\cong} H_n(M_1 \vee M_2) \xrightarrow{\cong} H_n(M_1) \oplus H_n(M_2), [M_1 \# M_2] \mapsto [M_1] + [M_2]$$

sends the fundamental class of $[M_1 \# M_2]$ to the sum of the fundamental classes of M_1 and M_2 .

Moreover, the cup product is natural so that we get a commutative diagram

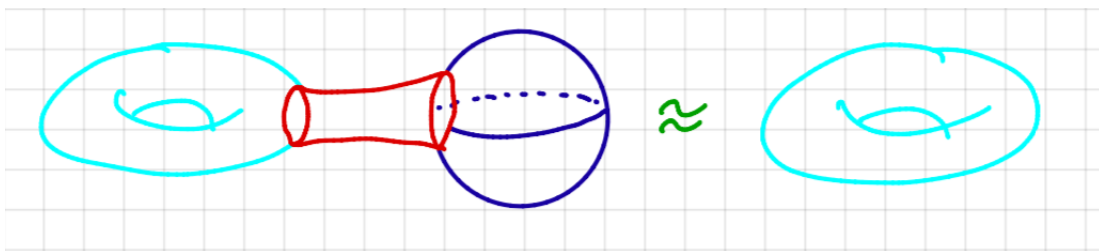
$$\begin{array}{ccccc} H^p(M_1 \# M_2) \otimes H^p(M_1 \# M_2) & \xrightarrow{\cup} & H^{2p}(M_1 \# M_2) & \xrightarrow{\langle -, [M_1 \# M_2] \rangle} & \mathbb{F}_2 \\ \downarrow & & \downarrow & & \downarrow \\ H^p(M_1) \otimes H^p(M_1) \oplus H^p(M_2) \otimes H^p(M_2) & \xrightarrow{\cup} & H^{2p}(M_1) \oplus H^{2p}(M_2) & \xrightarrow{\langle -, [M_1] \rangle + \langle -, [M_2] \rangle} & \mathbb{F}_2. \end{array}$$

Now it remains to translate this into the intersection pairing in homology which proves the claim. **QED**

Classification of surfaces

Motivated by the examples of the torus and real projective plane we are going to focus now on **the case $n = 2$** , i.e., two-dimensional manifolds which we are going to call **surfaces**. In fact, we are going to study **compact surfaces**. In this case we have an intersection form on $H_1(M)$.

We write **Surf** for the set of **homeomorphism classes of compact connected surfaces**. The connected sum operation provides it with the structure of a **commutative monoid**. The **neutral element** being S^2 , since $S^2 \# \Sigma \approx \Sigma$ for any surface Σ .



There is the following important result:

Theorem: Classification of surfaces

Associating the intersection form to a surface defines an **isomorphism** of commutative monoids

$$\text{Surf} \xrightarrow{\cong} \text{Bil.}$$

This theorem is **great** because it gives us a **complete algebraic classification of a class of geometric objects**. This is one reason why algebraic topology is so useful.

Actually, we are not finished with proving the theorem yet. Our examples show us that T^2 corresponds to H and $\mathbb{RP}^2$ corresponds to I . And S^2 is sent to the neutral element.

It remains to **show the relation** (and that this is the only relation)

$$T^2 \# \mathbb{RP}^2 \cong \mathbb{RP}^2 \# \mathbb{RP}^2 \# \mathbb{RP}^2.$$

One way to do this is to **triangulate the surfaces** involved. This requires too much geometric thinking for us today.

Instead, we make the following **observation**. We have not defined an orientation, but assuming we know what that means it is a surprising fact that even though we never assumed anything about orientations and worked with F_2 -coefficients, the theorem tells us what the **orientable surfaces** look like.

For, the **orientable surfaces** correspond to the forms gH where g is the **genus** of the surface. This follows from the facts that T^2 is **orientable** whereas $\mathbb{RP}^2$ is **not**.

In other words, any compact connected **orientable** surface Σ of genus g is homeomorphic to the connected sum of g tori

$$\Sigma \cong T^2 \# \cdots \# T^2.$$

Since S^2 is also an orientable surface, we allow $g = 0$ for this case too.

The real projective plane is not orientable. Therefore, any surface which is homeomorphic to a connected sum of at least one copy of $\mathbb{RP}^2$ is **not orientable**.

Quadratic refinement and the Kervaire invariant

Recall that away from characteristic 2 there is a bijection between quadratic forms and symmetric bilinear forms. However, since we are working over F_2 , we can ask whether there is a quadratic refinement q of the intersection form such that

$$q(x + y) = q(x) + q(y) + x \cdot y.$$

For such a refinement to exist requires $x \cdot x = 0$ for all $x \in H_1(M; \mathbb{F}_2)$, since

$$0 = q(2x) = q(x) + q(x) + x \cdot x = x \cdot x.$$

Hence we can only expect such a refinement on a sum of tori, i.e., on an orientable surface.

The existence of a quadratic refinement is an additional structure associated with the intersection form. Geometrically, it corresponds to a trivialization of the normal bundle of an embedding into an $\mathbb{R}^N$ for some N sufficiently large. Such a trivialization is called a **framing**. There is an invariant for quadratic forms in characteristic two, called the **Arf invariant**. In the case of a surface, or more generally a manifold of dimension $4k + 2$ (the only dimension where interesting things happen for this invariant), this invariant is called the **Kervaire invariant**. This invariant is a measure for if we can do certain **surgery** maneuvers on a manifold or not. Kervaire and Milnor used this invariant to study the differentiable structures on spheres.

But there were certain dimensions they could not completely explain. To settle the missing dimensions remained an open problem for about 60 years until Mike **Hill**, Mike **Hopkins**, and Douglas **Ravenel** finally solved the mystery (almost completely as there is one dimension left, it is 126) in a groundbreaking work in 2009 (published in 2016) using highly sophisticated methods in **equivariant stable homotopy theory**.