

LECTURES ON SHARDS

HUGH THOMAS

Let (W, S) be a Coxeter group, with $S = \{s_1, \dots, s_n\}$. I am only going to discuss the case that W is finite, even though significant parts of what I'm going to say can be adapted to the case that W is infinite.

Let $T = \{ws w^{-1}\}$ be the set of reflections of W . Let V be a real n -dimensional vector space, and consider an action of W on V where $t \in T$ acts by the reflection in the hyperplane H_t (which are all distinct).

Let $\mathcal{H} = \{H_t \in T\}$. The convex geometry of this hyperplane arrangement is intimately connected to W , and to a certain order on W called right weak order: $u \leq v$ if $v = us$, $s \in S$, and $\ell(v) = \ell(u) + 1$.

$W, \leq$ is a lattice, that is to say, given two elements u, v , there is a unique smallest element greater than both u and v , called the join of u and v , written $u \vee v$ and a unique largest element smaller than both u and v , called the meet of u and v , written $u \wedge v$. (We will prove this shortly.) A key structural ingredient of a lattice are the join-irreducible elements (those which can't be written non-trivially as a join.) In a finite lattice, being join-irreducible is equivalent to covering exactly one element. If j is join-irreducible, we write j_* for the unique element it covers.

(There is also a dual notion of meet-irreducible element. There is no reason to prefer join-irreducibles to meet-irreducibles, but for my purposes here it suffices to focus on one or the other.)

1. Shards. Our goal in this minicourse will be to study Nathan Reading's *shards*, which are geometrical objects, coming from the reflection arrangement $\mathcal{H}$, which are in natural bijection with the join-irreducible elements of $(W, \leq)$.

2. Shards and lattice quotients. If L is a lattice, and $\equiv_\Theta$ is an equivalence relation on L which respects the lattice operations, then the equivalence classes L/Θ also have a natural lattice structure. Such a lattice is called a *lattice quotient* of L . One of the reasons to be interested in join-irreducible elements of a lattice is to control lattice quotients.

Specifically, we say that a lattice quotient Θ contracts a join-irreducible j if $j \equiv_\Theta j_*$. A lattice quotient is determined by the set of join-irreducibles it contracts.

Shards help us understand the lattice quotients $(W, \leq)$: Nathan Reading tells us that the quotient amounts to erasing the shards which are contracted and looking at the induced poset on the resulting regions.

3. Shard polytopes (type A_n). We say that the 1-skeleton of a polytope P in a vector space V realizes a lattice L if there is a bijection $p : x \rightarrow p_x$ from L to the vertices of P , such that p_x and p_y are connected by an edge if and only if there is a cover relation between x and y , and there is ρ in V^* such that $x \leq y$ iff $\rho(p_y - p_x) > 0$.

There is a beautiful polytope which realizes the right weak order of W : namely, take any point $p \in V$ which lies on no hyperplane, and take the convex hull of its

W -orbit. Nathan Reading asked whether any lattice quotient of weak order also has some polytope whose 1-skeleton realizes it. This question was answered in the affirmative by Pilaud and Santos [PS], and again by Padrol, Pilaud, and Ritter [PPR].

Padrol, Pilaud, and Ritter associated certain *shard polytopes* to the shards in type A_n and B_n . I will explain this in type A_n (from a slightly different perspective). Shard polytopes give a way to answer this question of Reading's: the necessary polytope is the Minkowski sum of the non-contracted shards.

4. Shards and representation theory. Padrol, Pilaud, and Ritter needed to be clever in order to construct the shard polytopes. I will explain how it would have been possible to construct the shard polytopes with no cleverness, using representation theory. In fact, this approach would also have allowed us to discover the definition of shards! What we need is the notion of semistability of representations of a quiver. I am going to explain this first in type A_n , where we can make our lives slightly simpler, at the expense of choosing the algebra that we are using in a way that does not directly extend to other types. I will also explain how one can define shard polytopes in a uniform way in arbitrary simply-laced type (ADE). It remains a conjecture that they have all the good properties of the type A_n shard polytopes.

1. SHARDS

1.1. Order on chambers. In setting out to define shards, it costs nothing to work in a slightly more general setting. Let $\mathcal{H}$ be a finite collection of linear hyperplanes in an n -dimensional real vector space V . The *chambers* of $\mathcal{H}$ are the connected components of $V \setminus \bigcup_{H \in \mathcal{H}} H$. It is convenient to assume that the collection is reduced, i.e., that $\bigcap_{H \in \mathcal{H}} H = \{0\}$.

Example 1. *Because we have assumed that our hyperplanes are linear (rather than affine), the pictures we can draw in rank 2 are not very exciting.*

There are two other things we can do to draw higher-dimensional pictures. One is to take an affine slice. This allows us to draw rank three pictures, but it has the disadvantage that we don't see all the chambers.

A third thing we can do is to intersect the hyperplanes with a sphere around the origin, and then stereographically project. The hyperplanes turn into circles.

We choose a chamber B . For any chamber, C , define $s_B(C)$ to be the set of hyperplanes having B and C on opposite sides. We follow [E, BEZ] in ordering the chambers of $\mathcal{H}$ by the inclusion order on $s_B(C)$. We write $L_{\mathcal{H}}$ for the resulting... poset.

Lemma 1. *Cover relations in $L_{\mathcal{H}}$ correspond to pairs of chambers sharing a codimension 1 boundary.*

Proof. It's clear that crossing a hyperplane from one chamber to an adjacent one increases the separation set by one hyperplane, so it is certainly a cover relation. Conversely, if we have two chambers $X > Y$, draw a line from a point in the interior of X to a point in the interior of Y , generically enough that it doesn't pass through any codimension 2 or higher intersections of hyperplanes. Since X and Y are on opposite sides of exactly the hyperplanes in $s_B(X) \setminus s_B(Y)$, the line must cross each of them once, and no others. This gives us an order in which we can pass from X

to Y removing one element from the separating set at each step, and thus crossing from one chamber to another via a facet. $\square$

A chamber of $\mathcal{H}$ is called simplicial if it is cut out by n hyperplanes (rather than more). An arrangement is simplicial if all its chambers are simplicial. Any hyperplane arrangement in two-dimensions is simplicial, but this is not true in dimension 3 and higher.

Proposition 1. *If we assume that all the chambers are simplicial, then $L_{\mathcal{H}}$ is a lattice.*

(In fact, this is a stronger assumption than necessary, but it is good enough for our purposes, since reflection arrangements are simplicial.)

To prove this, we will use a result which goes back to [BEZ], and which Reading refers to as the BEZ Lemma.

Lemma 2. *Let L be a finite poset with a minimum element $\hat{0}$ such that whenever x is covered by y and z , then the join $y \vee z$ exists. The L is a lattice.*

This lemma has a fairly simple proof by induction (which I skipped in the lecture).

Proof. The proof is by induction on the size of L . We will first show that for any pair of elements x, y , the join exists. We can assume that x and y are incomparable, since otherwise the join is the larger of them. Let a be an element below x covering $\hat{0}$, and let b be an element below y covering zero. If $a = b$, then we can reduce to the subset of L above a , and we are done by induction. Otherwise, a, b both cover zero, so $a \vee b$ exists by hypothesis. There is a well-defined join of x and $a \vee b$, since both are contained in the part of L above a , and similarly there is a well-defined join of $a \vee b$ and y . Now the join of $x \vee (a \vee b)$ and $(a \vee b) \vee y$ is well-defined, since this is within the part of L above $a \vee b$. Further, this join must be $x \vee y$.

This guarantees that we have joins. To show that $x \wedge y$ exists, take the join of all the elements below x and below y . This set is non-empty, since $\hat{0}$ is in it. The join of all these elements must itself be below x and below y , so it is the meet of x and y . $\square$

Proof of Proposition 1. To apply the BEZ Lemma, we must consider what happens when a chamber x is covered by two other chambers y, z . Necessarily, the picture looks like the rank 2 picture: the two codimension 1 faces separating x from y and x from z meet in codimension 2 (by the assumption of simplicialness). Clearly the top element in the rank 2 subposet is the join. $\square$

1.2. Join-irreducible elements of $L_{\mathcal{H}}$. Recall that a non-minimal element is join-irreducible if it cannot be written as the join of two smaller elements. It is easy to see that this is equivalent to covering exactly one element.

There is a map from join-irreducible chambers to hyperplanes, sending a join-irreducible to the unique hyperplane which one can cross and go down.

Looking again at the rank 2 situation, we see that we can't simply assign a join-irreducible to each hyperplane (as we might have guessed): sometimes two join-irreducibles correspond to the same hyperplane. So, we do the simplest thing imaginable: we split in two the ones that correspond to two join-irreducibles.

In general, what we are going to do is to assign to each join-irreducible which corresponds to a hyperplane H a “shard” of H . These shards will cover H and will be disjoint except on boundaries.

The way this works is as follows. Given several hyperplanes which all intersect in codimension 2, we know what the picture looks like: it is exactly like the two-dimensional case, with a further $n - 2$ dimensions in which nothing is happening. Given such a codimension 2 intersection, we temporarily forget all the hyperplanes that do not contain the intersection. The base region B lies in some region of this arrangement with fewer hyperplanes, so we can identify two *basic* hyperplanes, which are the two between which B lies.

Given a hyperplane $H \in \mathcal{H}$, we remove from it those codimension 2 intersections for which H is not basic. Since these intersections are codimension 2 in V , they are codimension 2 in H , so they disconnect H into multiple components.

The closures of these components are the shards.

Example 2. *The A_3 example.*

1.3. Shards for reflection arrangements. If we identify the chambers with elements of W by $w \rightarrow wB$, then the order on chambers turns onto the weak order.

For each hyperplane, we choose a normal vector in V^* which pairs positively with the points in the base region B . The collection of all these normal vectors, we call Φ^+ .

Let $e_1, \dots, e_n$ be the normal vectors corresponding to the hyperplanes adjacent to B . (In the language of root systems, these are the simple roots.)

Hyperplanes that intersection in codimension 2 correspond to normal vectors lying in a 2-plane. The hyperplane H_γ will be split by H_α and H_β if γ is a non-negative linear combination of α and β .

1.3.1. Shards in type A_n . Let $\Phi^+ = \{e_i + \dots + e_j \mid i \leq j\}$. This defines the type A_n arrangement. (This convention makes it slightly less obvious that it is really type A_n , giving us the symmetric group $\mathfrak{S}_{n+1}$. On the other hand, this is the uniform perspective.)

Consider the hyperplane $H_{e_i + \dots + e_j}$ for $i \leq j$. A pair of basic hyperplanes cutting this hyperplane are necessarily of the form $H_{e_i + \dots + e_k}, H_{e_{k+1} + \dots + e_j}$. So the hyperplane $e_i + \dots + e_j$ is cut $j - i$ times. Recall that the cuts are codimension 1 inside $H_{e_i + \dots + e_j}$, and it is easy to see that their normals are linearly independent. Thus, $H_{e_i + \dots + e_j}$ is divided into 2^{j-i} shards.

We see that this agrees with what we saw in A_3 .

1.3.2. Shards in type D_4 . Number the vertices of the D_4 diagram 0,1,2,3, so that 0 is the central vertex. For all the roots except the highest one, the situation is as in type A_n : there are at most three cuts, and their normals are linearly independent.

The highest positive root is $2e_0 + e_1 + e_2 + e_3$. The hyperplane $H_{2e_0 + e_1 + e_2 + e_3}$ is cut by the hyperplanes $H_{e_0 + e_1}, H_{e_0 + e_2 + e_3}$, and cyclically — and also the hyperplanes $H_{e_0}, H_{e_0 + e_1 + e_2 + e_3}$. As in type A , the first three pairs of hyperplanes are linearly independent and cut the hyperplane into 8 regions — but then there is an additional hyperplane. It cuts six of the remaining 8 regions, making a total of 14 shards.

1.4. Shards are in bijection with join-irreducibles. Let S be a shard of some hyperplane H , and consider the chambers within S . We put an order on them by considering the corresponding $\mathcal{H}$ -chambers. The goal will be to show that a

minimal chamber in S corresponds to a join-irreducible $\mathcal{H}$ -chamber, and that there is exactly one minimal element with respect to this order, thus giving us what we want.

The approach I am taking here is closer to that of [M] than the approach of Reading. Reading's approach uses some more powerful lattice theory, which is lovely, but I would either have had to not prove a bunch of things, or else develop a lot of machinery.

1.5. Shards are in bijection with join-irreducibles. Let S be a shard of some hyperplane H , and consider the chambers within S . We put an order on them by considering the corresponding $\mathcal{H}$ -chambers. The goal will be to show that a minimal chamber in S corresponds to a join-irreducible $\mathcal{H}$ -chamber, and that there is exactly one minimal element with respect to this order, thus giving us what we want.

The approach I am taking here is closer to that of [M] than the approach of Reading. Reading's approach uses some more powerful lattice theory, which is lovely, but I would either have had to not prove a bunch of things, or else develop a lot of machinery.

Lemma 3. *If a chamber x of S is minimal, then the corresponding $\mathcal{H}$ -chamber, X , is join-irreducible.*

Proof. Assume that the chamber X is not join-irreducible. Therefore, there are at least two ways to go down from it, crossing H and crossing some other hyperplane J . Draw the rank 2 picture for H and J . We can see that there is another chamber lower than X with H lying below it, and such that it lies in the same shard, because, H being basic in this arrangement, it is not cut. $\square$

What are we really doing when we make this “rank 2” argument? We are saying: take a generic point in the codimension 2 intersection in question, and take a little two-dimensional slice transverse to it.

Lemma 4. *Adjacent chambers in S are comparable.*

Proof. We draw the rank 2 picture for the subspace where the chambers intersect. The hyperplane H must be basic, since otherwise the two chambers we are considering would belong to different shards. Now it is visible in the rank 2 picture. $\square$

The following lemma is slightly confusing to state, but useful.

Lemma 5. *The order defined on chambers of S by transitive closure of the induced order on adjacent pairs is the same as the induced order.*

The point of the lemma is that we could define another order on the chambers of S , just by taking adjacent pairs of chambers and ordering them according to the order on the corresponding $\mathcal{H}$ -chambers, and then taking the transitive closure. This is somehow more in keeping with the idea of localizing our attention to S .

Proof. Suppose that we have x, y chambers in S , such that the corresponding $\mathcal{H}$ -chambers are X, Y , with $X < Y$. There is therefore some hyperplane J defining a wall of X by which we can go up while staying below Y . This wall cuts S along a facet of x . Let z be the chamber in S on the opposite side of J from x , and Z the corresponding $\mathcal{H}$ -chamber. Then

$$s(Z) = s(X) \cup \{\text{the hyperplanes cutting } H \text{ along } H \cap J\}.$$

But $s(Y)$ certainly contains $s(Z)$, so $Y > Z$, and we are done by induction. $\square$

Lemma 6. *If we have three chambers x, y, z in S with x covering y and z , then there is also a chamber $w \in S$ less than y, z .*

Proof. We draw the ...rank three picture! We see that X , the $\mathcal{H}$ -chamber corresponding to x is at the very top (since there are three ways to go down from it). So the hyperplanes incident to it are all basic and don't ever get cut. The consequence is that H (restricted to this picture) is a full two-dimensional hyperplane. Its bottom region w , corresponds to the $\mathcal{H}$ -chamber W , which is only separated from B by H . Thus, W is below Y and Z , so the same holds for w, y, z . $\square$

Lemma 7. *If a chamber m in S is minimal, then it is actually the minimum.*

Proof. Assume otherwise, so m is minimal but not the minimum. Since the poset of chambers of S is connected, there must be an element x which lies over m and also lies over some element z which does not lie over m . Choose x to be minimal among such elements, x must cover z . By Lemma 5, x also covers some y which lies over m . Now applying Lemma 6, we get a chamber w below z and y . Now w cannot lie over x , but y lies over w and over m . So our initial choice of x was not minimal. $\square$

Theorem 1. *The natural map from join-irreducibles to shards, sending a join-irreducible chamber J to the shard including the unique panel below J , is a bijection.*

Proof. For each shard S , we know that the order induced on its chambers has a unique minimum m (by Lemma 7). The corresponding $\mathcal{H}$ -chamber, M is join-irreducible by Lemma 3. Any other chamber x in the shard is not minimal, so it has at least one lower neighbour in the poset of chambers in the shard, and therefore there is at least one way to descend from the corresponding $\mathcal{H}$ -chamber X , other than crossing S . So X has at least two lower covers, and it is therefore not join-irreducible. $\square$

We will later need a lemma whose proof fits naturally here:

Lemma 8. *For a shard S , there is a unique minimal element among all the elements having an upper facet on S , and it is J_* , where J is the join-irreducible corresponding to S .*

Proof. Let U be a chamber with a lower facet on S . We have already seen that $J \leq U$, and in fact that there is a sequence of adjacent chambers $x_0, \dots, x_r$ in H where, if we write X_i for the $\mathcal{H}$ -chamber having x_i as its lower facet, then $X_i < X_{i+1}$. But looking at the corresponding rank two picture, we also see that if we write Y_i for the $\mathcal{H}$ chamber having x_i as its upper facet, the same analysis shows that $Y_i < Y_{i+1}$. The desired conclusion follows. $\square$

2. SHARDS AND LATTICE QUOTIENTS

An equivalence relation $\equiv_\Theta$ on a lattice L respects the lattice operations if $x \equiv_\Theta x'$ and $y \equiv_\Theta y'$ imply that $x \vee y \equiv_\Theta x' \vee y'$ and $x \wedge y \equiv_\Theta x' \wedge y'$. Obviously, this means that the lattice operations descend to well-defined operations $\vee$ and $\wedge$ on the equivalence classes.

Note that in a lattice L , we have $x \geq y$ iff $x \vee y = x$ iff $x \wedge y = y$. This makes clear how to define the lattice operations on $L / \equiv_\Theta$.

Proposition 2. *There is a lattice structure on the equivalence classes, where $[x]_{\Theta} \geq [y]_{\Theta}$ iff $[x]_{\Theta} \vee [y]_{\Theta} = [x]_{\Theta}$ iff $[x]_{\Theta} \wedge [y]_{\Theta} = [y]_{\Theta}$*

Proof. Exercise. $\square$

Lemma 9. *The equivalence classes of a lattice congruence are intervals.*

Proof. The join of all of the elements of an equivalence class must also lie in the same equivalence class, thus the equivalence class has a maximum element. The dual argument shows that each equivalence class has a minimum element. Further, by antisymmetry of the poset structure on the equivalence classes, any element in the interval must be in the equivalence class as well. $\square$

Given a lattice congruence $\equiv_{\Theta}$, we say that it *contracts* a cover relation $b \succ a$ iff $b \equiv_{\Theta} a$. We say it contracts a join-irreducible j if and only if $j \equiv_{\Theta} j_*$.

Lemma 10. *The Hasse diagram of $L / \equiv_{\Theta}$ is obtained from that of L by contracting the cover relations contracted by $\equiv_{\Theta}$, and replacing multiple edges joining a given pair of (new) vertices by a single edge.*

Proof. Contracting the cover relations replaces equivalence classes (which are intervals by the previous lemma) by a single vertex, so the given procedure produces the right collection of vertices. Uncontracted cover relations certainly remain order relations. Suppose that a cover relation $X \succ Y$ became a relation $[X] \succ [Y]$ which was not a cover relation, i.e., we have $[X] \succ [Z] \succ [Y]$ for some Z . But in this case, $Z \vee Y \succ Y$ and $X \succ Y$, so $(Z \vee Y) \wedge X \geq Y$. But $[(Z \vee Y) \wedge X] = Z$, so we have $X \succ (Z \vee Y) \wedge X \succ Y$, which is a contradiction. $\square$

So, let us analyze which cover relations from the lattice of chambers survive in the quotient.

Lemma 11. *If $X \succ Y$ is separated by a shard that is contracted, then $X \succ Y$ is contracted.*

Proof. Let J be the join-irreducible corresponding to the facet in question. We have $J \leq X$ and $J_* \leq Y$. Further, $J \vee Y = X$, but $J_* \vee Y = Y$. Since $J \equiv_{\Theta} J_*$, we must have $X \equiv_{\Theta} Y$. $\square$

Lemma 12. *If $X \succ Y$ is separated by a shard that is not contracted, then $X \succ Y$ is not contracted.*

Proof. Let J be the join-irreducible corresponding to the separating shard. We prove the contrapositive. Suppose $X \equiv_{\Theta} Y$. Then $X \wedge J = J$, but $Y \wedge J = J_*$. So $J \equiv_{\Theta} J_*$. $\square$

The previous discussion proves the following theorem.

Theorem 2. *The Hasse diagram for a lattice quotient of the lattice of $\mathcal{H}$ -chambers is obtained by erasing the contracted shards, putting a vertex in each resulting region, and putting covering relations through all the remaining walls.*

2.1. Which sets of shards can be contracted? We have seen that a lattice quotient of $L_{\mathcal{H}}$ contracts some set of shards, which determine the quotient. It is natural to ask which sets of shards can be contracted.

We say that a shard S forces a shard T if any quotient in which S is contracted also has T contracted.

Looking at the rank two situation, we see that either of the shards corresponding to a basic hyperplane forces all the others, except the other basic hyperplane, while none of the other shards force any other shard.

By looking at the rank two configurations in $\mathcal{H}$, we obtain the following:

Lemma 13. *Near a generic point p in a codimension 2 intersection in $\mathcal{H}$, we have shards S_1 and S_2 containing a neighbourhood around p of lying of the basic hyperplanes H_1 and H_2 for this intersection. We have an additional two shards, with p on their boundary, for each of the other hyperplanes containing $H_1 \cap H_2$. Then each of S_1 and S_2 force all the additional shards.*

In the above situation, we say the shards S_1 and S_2 cut the additional shards.

This gives us a necessary condition for a set of shards to be contracted:

Proposition 3. *In order for a set of shards to be contracted by some lattice equivalence, it must be closed under the cutting relation defined above.*

Reading also proves the converse of this proposition, but we will not see the proof.

3. SHARD POLYTOPES (TYPE A_n)

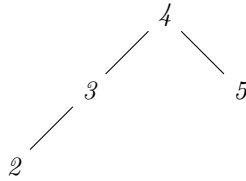
Let S be a shard of type A_n . Padrol, Pilaud, and Ritter define a polytope associated to this shard.

The shard is in bijection with a join-irreducible element of $\mathfrak{S}_{n+1}$. Let's suppose the shard is in the hyperplane $H_{e_i + \dots + e_{j-1}}$. Reflection in this hyperplane corresponds to the transposition (ij) , so the join-irreducibles corresponding to regions just above the hyperplane are those permutations consisting of two increasing words, the first ending at j and the second starting at i . It follows that the permutation is of the form $1 \dots i-1, L, j, i, R, j+1, \dots, n$, where each number from $i+1$ to j belongs to one of L or R .

In order to define the shard polytope, we need to draw a poset P_S on $i, \dots, j-1$ where $k < k+1$ if $k+1 \in L$ and $k > k+1$ if $k+1 \in R$.

Then the shard polytope is the convex hull of the indicator vectors of the order ideals of P_S . (I have changed basis with respect to the definition given by [PPR]. For the equivalence, see [AB+].)

Example 3. Consider $(ij) = (26)$, $L = \{3, 4\}$, $R = \{5\}$.



The shard polytope is the convex hull of the indicator vectors:

00000 01000 01100
 00001 01001 01101 01111

The Minkowski sum of polytopes is defined as follows:

$$P + Q = \{p + q \mid p \in P, q \in Q\}.$$

The goal of Padrol, Pilaud, and Ritter is the following result:

Theorem 3. *Let Θ be a lattice congruence on $\mathfrak{S}_{n+1}$. Then S_{n+1}/Θ is realized as the one-skeleton of the polytope obtained by taking the sum of all the shard polytopes corresponding to join-irreducible elements not contracted by Θ .*

Example 4. $\mathfrak{S}_3$.

Example 5. *The Tamari lattice T_n is a well-known quotient of $\mathfrak{S}_n$. A polytope whose 1-skeleton realized the Tamari lattice is called an associahedron. One way to construct an associahedron is as the Minkowski sum of simplices $0, e_i, e_i + e_{i+1}, \dots, e_i + \dots + e_j$ for $1 \leq i \leq j \leq n-1$. These are the shard polytopes that come from exactly the join irreducibles with $R = \emptyset$.*

4. SHARDS AND REPRESENTATION THEORY

In this section, we will show how shards correspond exactly to the region of semistability of an indecomposable representation of a certain quiver with relations.

I'm going to focus on type A_n , where the combinatorics is straightforward and can be done by hand, but I also want to explain (without proof) how the same approach applies to arbitrary ADE types.

Representations of quivers. Let $Q = (Q_0, Q_1)$ be a quiver (a directed graph), with vertex set Q_0 and arrow set Q_1 . A representation X of Q over a field K is an assignment to each $i \in Q_0$ a vector space X_i and to each arrow $i \xrightarrow{\alpha} j$ a linear map $X_\alpha : X_i \rightarrow X_j$.

For two explicitly given representations X, Y to be isomorphic, there is a change of basis which takes one to the other.

It will be important for us to consider the subrepresentations of a representation. A subrepresentation of a representation X is a choice of a subspace $Y_i \leq X_i$ for each $i \in Q_0$, such that for each $i \xrightarrow{\alpha} j$, we have $X_\alpha(Y_i) \subseteq Y_j$. That is to say, we choose subspaces Y_i so that the maps X_α restrict to the subspaces Y_i to define a representation.

If we have two representations X, Y of Q , we define $X \oplus Y$ by $(X \oplus Y)_i = X_i \oplus Y_i$, and for $i \xrightarrow{\alpha} j$ we define $(X \oplus Y)_\alpha = X_\alpha \oplus Y_\alpha$. A concept of fundamental importance in the representation theory of quivers is that of *indecomposable representation*. A non-zero representation X is indecomposable if $X \simeq X_1 \oplus X_2$ implies that X_1 or X_2 is 0.

A quiver is called of type A_n if it is an orientation of a type A_n Dynkin diagram, i.e., it is some orientation of a path. It is relatively easy to convince oneself that if take such a Q , and we take a sub-path, then, if we put copies of K at those vertices, with identity maps between them, and zero vector spaces elsewhere, these representations will be indecomposable.

Example 6. *Some indecomposable representations.*

In fact every indecomposable representation of this quiver is isomorphic to one of this form. Note, in particular, that the fact that we chose identity maps is not important: any non-zero scalar multiple would have worked just as well, and does not change the isomorphism class of the representation (because it can be absorbed into the change of basis).

I should point out that it is a very special property of this quiver which means that it is so easy to describe the indecomposable representations. In fact, a quiver has only finite many indecomposable representations if and only if it is an orientation of a simply-laced Dynkin diagram (*ADE*). Further, in this case, the indecomposable representations are in bijection with the positive roots by the map sending an indecomposable representation X to $\sum_i \dim X_i e_i$.

Representations of a quiver with relations. We can complicate matters for ourselves by considering quivers with relations instead of only quivers. A relation is a linear combination of paths, of length at least 2, all having the same starting and ending point. A representation of a quiver with relations is a representation in the usual sense in which the corresponding relations hold.

All the other definitions work the same way in this more general setting. It turns out that every category of finite-dimensional representations of a finite-dimensional algebra, over an algebraically closed ground field, is equivalent to a category of representations of some quiver with relations. So this is in fact one perspective on a very general topic.

We will be particularly interested in the representations of the quiver with relations $(\overline{A}_n, R)$ which is obtained by doubling an A_n Dynkin diagram, so that each edge is replaced by a pair of arrows going in opposite directions, and imposing the relations that the composition of each pair of opposite arrows is zero (composed in either order).

Observe that any indecomposable representation of A_n with any orientation defines an indecomposable representation of $(\overline{A}_n, R)$, where we assign the zero map to each arrow that is not in the orientation of the quiver from which we started.

Example 7. *Some examples.*

In fact, every indecomposable representation of $(\overline{A}_n, R)$ is isomorphic to one of this form. Again, this is special to the fact that we started with an A_n quiver.

We notice that the number of indecomposable representations of $(\overline{A}_n, R)$ corresponding to the root $e_i + e_{i+1} + \cdots + e_{j-1}$ is 2^{j-i-1} . Hmm.

Stability conditions. For Q a quiver, we write $K_0(Q)$ for the Grothendieck group of Q , which is by definition a vector space spanned by $e_i = [S_i]$, where S_i is the simple representation supported at vertex i . A representation X of Q has a corresponding class $[X] \in K_0(Q)$, namely $[X] = \sum_{i \in Q_0} e_i \dim X_i$. It will be convenient to work with $K_0(Q)_{\mathbb{R}} = K_0(Q) \otimes_{\mathbb{Z}} \mathbb{R}$.

Write $K_0(Q)_{\mathbb{R}}^*$ for the dual space, i.e., $\text{Hom}_{\mathbb{R}}(K_0(Q)_{\mathbb{R}}, \mathbb{R})$.

Fix $\theta \in K_0(Q)_{\mathbb{R}}^*$. A representation X of Q is semistable with respect to θ if

- $\theta([X]) = 0$, and
- $\theta([Y]) \leq 0$ for Y any subrepresentation of X .

The notion of semistability comes from geometric invariant theory. For quivers like the one we consider, the original motivation will not be very relevant. (It plays more of a role when there are moduli of representations.)

Let X be an indecomposable representation of $(\overline{A}_n, R)$ supported on vertices $i, i+1, \dots, j$. It will be important to understand which are the subrepresentations of X . It is convenient to draw the *mountainscape* for X : that is to say, we interpret the non-zero arrows of X as cover relations in a poset. A poset of this form is called a fence poset. To choose a subrepresentation of X , we must choose a subspace of each vector spaces, but since the vector spaces of X are at most one-dimensional, we must simply, for each i in the support, decide whether to choose the corresponding vector space or not. The condition to be a subrepresentation (that the linear maps of X should restrict to define a representation on the chosen subspaces) means exactly that the subset of vertices where we select the vector space rather than zero, form an order ideal.

Now we can address the question: given an indecomposable representation X of $(\overline{A}_n, R)$, where is it semistable?

Proposition 4. *An indecomposable representation X supported at vertices $i, i+1, \dots, j-1$, is semistable at those $\theta \in H_{e_i+\dots+e_{j-1}}$ such that, for $i \leq k \leq j-2$, we have $\theta(e_i + \dots + e_k) \leq 0$ if $X_{k \leftarrow k+1}$ is non-zero and $\theta(e_i + \dots + e_k) \geq 0$ if $X_{k \rightarrow k+1}$ is non-zero.*

Proof. First we check that if X is θ -semistable then θ must be contained in the region described in the proposition. This is pretty immediate from the definition of semistability. By the first condition for semistability, θ must be in $H_{e_i+\dots+e_{j-1}}$. Then we observe that if $X_{k \leftarrow k+1}$ is non-zero, then X has a subrepresentation supported on vertices i to k , so $\theta(\alpha_i + \dots + \alpha_k) \leq 0$, and similarly if $X_{k \rightarrow k+1}$ is non-zero, then X has a subrepresentation supported on vertices $k+1$ to $j-1$, so $\theta(e_{k+1} + \dots + e_{j-1}) \leq 0$, and $\theta(e_i + \dots + e_k) = \theta(e_i + \dots + e_{j-1}) - \theta(e_{k+1} + \dots + e_{j-1}) = -\theta(e_{k+1} + \dots + e_{j-1}) \geq 0$.

The converse is (slightly) more complicated, because X may have many more subrepresentations beyond the ones we thought about in proving the first direction, which (in principle) could lead to further conditions on θ . But it does not. Suppose that θ is in the set described by the statement of the proposition. And suppose that Y is a subrepresentation. It is sufficient to consider the case that Y is indecomposable, since if $Y = Y_1 \oplus Y_2 \oplus \dots \oplus Y_r$ with the Y_k indecomposable then $\theta([Y]) = \theta([Y_1]) + \theta([Y_2]) + \dots + \theta([Y_r])$, and each of the Y_i would itself be an indecomposable subobject of X , so we would have $\theta([Y_k]) \leq 0$ for all k , which would establish what we want.

Now consider the case that Y is an indecomposable subobject. Suppose that Y is supported on the vertices $k, k+1, \dots, \ell$. The cases where $k = i$ or $\ell = j-1$ have already been considered above: the conditions in the statement of the proposition are exactly what arises from the definition of semistability. So consider the case that $i < k$ and $\ell < j-1$. In this case, because Y is a subobject, we have that $X_{k-1 \rightarrow k}$ is nonzero and $X_{\ell \leftarrow \ell+1}$ is nonzero. It follows that $\theta(e_i + \dots + e_{k-1}) \geq 0$ and $\theta(e_{\ell+1} + \dots + e_{j-1}) \geq 0$. Thus $\theta([Y]) = \theta([X]) - \theta(e_i + \dots + e_{k-1}) - \theta(e_{\ell+1} + \dots + e_{j-1}) \leq 0$, as desired. $\square$

We have proved:

Theorem 4. *There is a bijection between the shards for A_n and the indecomposable representations of $(\overline{A}_n, R)$, where the shard corresponding to the indecomposable X is the locus where it is semistable.*

4.1. Doubled quiver and preprojective relations. In order to carry out the same kind of plan for other finite Coxeter groups, we “just” need to find an appropriate algebra. This turns out to be fairly easy for the other simply-laced types (D, E) . For the non-simply-laced cases (B, C, F, G) , there are some technical complications which I would rather not go into, but the same ideas can be made to go through. For H_3, H_4 , and $I_2(m)$, the non-crystallographic types, the situation is quite different: because the Coxeter group is not a Weyl group (i.e., there is no way to choose a root system lying in a lattice $\mathbb{Z}^n$), while algebras will necessarily correspond to a Weyl group, it is much more difficult to address these cases using representation theory. (Nonetheless, there is a kind of “non-crystallographic folding” which relates $I_2(5)$ to A_4 , H_3 to D_6 , and H_4 to E_8 , so it is perhaps not totally hopeless. I am not aware of any attempt to pursue this, though.)

Now consider any type ADE . Let $\bar{Q}$ be the double of the Dynkin quiver, of type Q , and let Π be the relations that at each vertex i that the sum of all the paths of length two from i to i gives zero.

$(\bar{A}_n, \Pi)$ has many more indecomposable representations than $(\bar{A}_n, R)$, and for almost all types, infinitely many indecomposable representations. However, we can restrict our attention to the representations which are *bricks*. A brick is a representation all of whose non-zero endomorphisms are invertible.

Proposition 5. *The indecomposable representations of $(\bar{A}_n, R)$ are exactly the brick representations of $(\bar{A}_n, \Pi)$*

Partial proof. The relations of R imply that all the paths of length 2 from i to i are zero for any i . Thus, in particular, their sum is zero, so a representation of $(\bar{A}_n, R)$ is a representation of $(\bar{A}_n, \Pi)$. Further, it is easy to check that if X is an indecomposable representation of $(\bar{A}_n, R)$ over the field K , then $\text{Hom}(X, X) = K$: the only endomorphisms are multiples of the identity.

It remains to show that there are no additional brick representations of $(\bar{A}_n, \Pi)$, which is not too hard. $\square$

The general statement is proved in [T].

Theorem 5. *For simply-laced Dynkin types, there is a bijection between brick representations of the preprojective quiver with relations (Π_Q, T) and the shards of the Weyl group of the same type; the shard corresponding to a brick B is exactly the set of θ such that B is θ -semistable.*

Unlike in type A_n , these bricks do not all have the property that there is some orientation of a Dynkin quiver of which they are a representation.

Shard polytopes via representation theory. Comparing the definition of the shard polytope (in the form I gave it) to the definition of the representation-theoretic shard, the following is clear:

Proposition 6. *The shard polytope corresponding to a shard module X is the convex hull of the dimension vectors of the sub-representations of X .*

The convex hull of the subrepresentations of X is called its Harder–Narasimhan polytope. They were first studied by Baumann–Kamnitzer–Tingley [BKT], and more recently by Fei [F1, F2] and Aoki et al. [AH+].

Fei and Aoki et al. both prove that if A is a finite-dimensional algebra with finitely many bricks, then the Minkowski sum of the Harder–Narasimhan polytopes

of all the bricks of A yields a polytope dual to the g -vector fan of A . Thanks to work of Mizuno [M], the g -vector fan of A can be described as the fan whose codimension 1 faces are the union of the domains of semistability of the bricks of A . There is a natural lattice on the set of vertices of the polytope dual to this fan (the lattice of *torsion classes* of the algebra).

In other words, this is a very general result which guarantees realizability of (finite) lattices of torsion classes. This includes, for example, Cambrian lattices. This result, by itself, does not constitute a proof of the result of [PPR], because there are quotients of the lattice which do not themselves come from any algebra. (This already happens in type A_3 .) However, using further facts about Harder–Narasimhan modules, it is pretty easy to recover the result of [PPR]. Extending it to other types turns out to be harder; I am working on it.

REFERENCES

- [AB+] Abram, Bastidas, Dequêne, Morales, Park, Thomas. Flows on graphs with cycles, locally gentle algebras, and the Mutoperhedron. arXiv:2601.08150.
- [AH+] Aoki, Higashitani, Iyama, Kase, Mizuno. Fans and polytopes in tilting theory I: Foundations. arXiv:2203.15213.
- [BKT] Baumann, Kamnitzer, Tingley. Affine Mirković–Vilonen polytopes. Publ. Math. Inst. Hautes Études Sci. 120 (2014), 113–205.
- [BEZ] Björner, Edelman, Ziegler. Hyperplane arrangements with a lattice of regions. Discrete and Computational Geometry 5 (1990), 263–288.
- [E] Edelman. A partial order on the regions of $\mathbb{R}^n$ dissected by hyperplanes. Trans AMS 283 (1984), 617–631.
- [F1] Jiarui Fei, Combinatorics of F -polynomials. IMRN 2023, 7578–7615.
- [F2] Jiarui Fei, Tropical F -polynomials and general presentations, J. LMS 107 (2023), 2079–2120.
- [M] Mizuno. Shard theory for g -fans. IMRN 2024, 13106–13126.
- [PPR] Padrol, Pilaud, Ritter. Shard polytopes. IMRN 2023, 7686–7796.
- [PS] Pilaud, Santos. Quotientopes, Bull. LMS 51 (2019), 406–420.
- [R1] Nathan Reading, Lattice and order properties of the poset of regions in a hyperplane arrangement. Algebra Universalis 50 (2003), 179–205.
- [R2] Nathan Reading, Lattice theory of the poset of regions. Chapter 9 of Lattice Theory: Special Topics and Applications, Volume 2, ed. G. Grätzer and F. Wehrung, Birkhäuser, 2016.
- [T] Hugh Thomas, Stability, shards, and preprojective algebras. Representations of algebras: proceedings of ICRA 2016. AMS, 2018, pp. 251–262.