

Some of which are not to non-standard concepts, but in any case.

So I think there's one site to pull from there. If we look at the units.

I would say that the sites are relatively consistent on units.

I guess all doesn't get too badly sparse apart in terms of measurements, at least from this query.

But I think if you look at the rates of measurements in units, that's maybe one of the more consistent fields. I think the biggest discrepancies are probably in terms of the value as concept ID field and then just generally the diversity in terms of how many distinct measurements you

Measuring with and I guess some of that is coming from

Some people are carrying

I think this against a cohort while others are carrying it against their overall data sets, and I think we talked a little bit about that last week as well, so

That's kind of my, let's say general take on this.

Let me also

Post this or we can see we can share this.

Go forward. Sure. Sorry

Williams, Andrew E. 17:36

You say you could post it.

Hughlaugh, Jewel 17:36

I was going

I was gonna shift to the other tabs, but this is a good transition to tie up in if you want.

Williams, Andrew E. 17:42

Just a brief one on the type concept ID.

It's the main way. As you're aware, the, as you're all aware, that a lot of Providence of where the data comes from are preserved because

Our project here together is really about multistate data and highlighting the value of getting data from different sources about the same thing

It's going to be increasingly important that we pay attention to how we are all collectively recording

That.

Resonance in using the type concept ID O NLP derived will be one you know.

They'll be instances where the same kind of thing might come either from you know, some algorithm we apply to a waveform, or an image, or something that comes out of a database that's automatically associated with the generation of those waveforms that might be complementary

We will want to be able to distinguish those and so

As we have in the past said, we're going to need a kind of review the vocabulary for things like C2 type. This is another area where

We have the coverage we need for type concept IDs for the different sources that are that we need to be able to distinguish, but we're gonna wanna enable queries that use that type concept ID to say why it's not only the thing and the setting and so on.

But the origin of the data as a as a feature of interest in doing work with the data after it's been captured.

To just a little note.

Hughlaugh, Jewel 18:10

One relevant thing that we've also we talked about a couple weeks ago in the measurement call is what we should be able to distinguish between ideally in tabular laboratory data like Low Inc.

Does that's using a table and create a control versus flow sheets. And so that's where you can have potentially the same measurement being represented in a low inc table as well as in a flow sheet.

But it's very important to make be able to make the distinction, especially if you're trying to define cohorts and you want one particular data.

Or another lab, or data, or gene expression.

So and I believe and I think where we had landed in terms of the standard concept for representing flow sheets was nursing reports.

But I can double check on that and I think I think we have discussion about this, but I think that was the type concept ID of data when we had talked about this prior.

I don't know Pauline or anybody else.

You remember.

Confirm that, but otherwise I can just find it offline.

OK. So that you're again for confirming.

I think this is already very interesting to see and what for me I find interesting is this has sites that have not yet or that I at least have it. Not yet that I have not ingested data centrally from your sites, but it also represents something that's maybe.

A different set of your measurements than what I have seen centrally or like the I have compared centrally.

I should say so.

Regularly receiving insights there.

Williams, Andrew E. 18:12

Last comment real quick.

Hughlaugh, Jewel 18:13

Yep.

Williams, Andrew E. 18:15

I think going back to the type concept ID, if you are aware that you have either state or National death index.

State registry statistics as we do or national death index data as a complementary source of mortality information that are being that are part of your own map.

It is that a really crucial one. I would say of the outcomes that matter in a lot of different studies.

That's a big one.

Completeness of data and accuracy of cause of death and other kinds of things vary a lot amount depending on whether you're drawing just on what's in your DHR or you've got a registry like that, a state or a national registry, and so if you're aware of that, just.

Be sure that that's marked in there. Thanks.

Hughlaugh, Jewel 18:40

All right.

And I'm just following up quickly on.

See on one of these I've had a comment in the chat.

So Heidi mentions.

Schmidt, Heidi E. 22:05

I.

I saw the section in red, so I thought I'd give context.

Hughlaugh, Jewel 22:06

Assume, yeah.

Thank you.

Yeah. So I guess my question then to you is in your measurement table, do you have both flow sheets and lookin mapping?

So like are you mapping flow sheet data in there?

Schmidt, Heidi E. 22:20

The majority of them are flowcharts and I'd say 27 million flow sheets.

Hughlaugh, Jewel 22:22

Yeah.

Schmidt, Heidi E. 22:28

And about how Millstone Low Inc.

And just that in the chat.

After we did our last.

Someone just said that people would have kind of like that high level context of.

What was coming from where?

As have from our site.

Help in that gap analysis.

So if we could cover flow sheets to be defined in order to get more information, then that would probably be something that would be handled either in searching with the tool or some sort of Delphi analysis I was thinking.

Hughlaugh, Jewel 22:33

Assume, Yep.

Thank.

OK. I'm gonna jump over here quickly to the value query.

To this was the initial query that I had sent, which is more or less trying to get a group about what are the most common.

Categorical values in your instance and I included four of the responses here. There is, I think one or two that I don't know if something got switched in the query or what, but I have excluded these.

In that MDC job of our Talk.

Despite of quick notes here.

So I'm roughly figure for you.

But not have a very large fraction of your values conceptually.

You're actually right to go.

I saw Mike is on the call.

I don't know Mike. If it's these are just by default zero is filling in your values concept ID or is this something that hasn't you haven't yet mapped, but you have like a source value that you're planning to map.

Kelley, Mike 22:33

Yeah. There are source values that were just working through the mapping for them right now.

Hughlaugh, Jewel 24:20

Assume. OK, then that's, that's perfect. Then I think that that's also helps you keep track of the the missing mapping essentially. So that's that's to be expected.

These ones, though, will need to be updated, so these are.

I did a quick check in Atlanta.

So there are numeric values represented as categorical.

I think sometimes you could use this like as a temporary response or something. I think at least for the top two I don't check all four, but they were both non-standard.

Take a quick look if you have a response.

That's like 24.4.

It's likely that that's gone into the value as number fields and your values concept ID will be empty because that is a numeric resp.

I'm guessing for responses I don't know what it actually for.

Wildly enough, there are non-standard concepts in OMA that are like numbers like this. So I understand why you would map it to that, but.

Yeah. If it's a numeric response, it's in the same, especially with the decimal and everything and it's, it's probably going to be going to be a volume number.

OK.

Quick note then for LP.

Of course, we talked about this in one of the prior calls, but.

There is a problem that you have identical concept names that are both standard concepts.

So what you have at LP is there's actually all of the negative response measurements like 1/3 of them are mapped to.

I think this is Low Inc code and 2/3 mapped to the State Inc version of negative.

I think we talked about this a couple weeks ago, but.

If it and I guess if your response that's like this, then definitely defaults to the low inc response.

And would by default if you have one of them in your set, I would try to consolidate and not have two different representations of negative. When you talk about.

Creating.

When down the line and you wanna have a very specific measurement with a specific response value.

It's important that these are consolidated and so part of the process of that, we're gonna try to do where we provide additional guidance with standardizing mappings is exactly this. If there are semantically very similar or even in this case, identical concepts that we guide you to choose.

The one that's gonna be consistent across the consortium.

So just a brief note there.

Interesting though.

Again, you see the Clango come scale.

So I think that that's some targeted mapping going in to Florida to get that.

These responses in particular.

Otherwise, always kind of interesting just to see what you're actually getting in this. In this values concept of a column.

Responses are also just to like the transcripts, but they also are percentages out of the total of the top 10.

So it's kind of a word process, but.

Just for reference there. OK then I will jump to units. And so this last one, it's a bit of an interesting one. Takes a look at

Your and I can maybe add an additional insert here.

Takes a look at your measurement concept.

OK.

So kind of your lab test in this case and it isolates the ones that have the most distinct units mapped.

So we're not talking about distinct unit source values now.

That are mapped like was talked about with pH. But we're talking about unit.

Measurement concept IDs and then associated distinct units.

And so what we have here is going for both LP and MDC, pretty consistent.

There's not a lot of mapping or at least unit diversity amongst your

Your measurements.

But when we get down into Columbia.

So for you, I see that Columbia has a good chunk. I think we've been over

is 1 000 million hundred million measurements or so mapped to O and you are expected within that kind of?

Group you have a bunch of unit mapping that make sense.

But from one that I wanted to mention here was referral lab results.

So we also.

See this. At least I don't know if Columbia on topic.

I can't remember, but we see this referral lab at our in our Talk data as well.

We have chosen to exclude.

These data points because at least as they are represented in Epic, we don't have any additional information about the lab that was taken.

My last here.

It shows up in it, in our case as epic as like this is something.

Unknown lab that the patient had and it was a new entry, but that's it.

So we see don't have this, just to know there's a lot of units associated with it. If you have somehow additional information that you that they could somehow tie this into a legitimate lab that would have then a response that you could use.

Or if this has with some other concept. I don't know the Columbia IDs or how this is working.

Maybe worth taking a look.

It's a a sizable number of measurements.

But it's something that I've seen as well and maybe worth just a quick investigation on.

Units.

Williams, Andrew E. 18:10

Comment real briefly.

I think in general when you have something like this at any site, regardless of your DHR.

Hughlaugh, Jewel 18:18

Yeah.

Williams, Andrew E. 18:28

There are probably several, sometimes more than one at your organization, who are whose job it is to know about this stuff.

People who are in charge of laboratory information systems or in charge of pharmacy data or in charge of pack system or whatever it is, and I think.

Sometimes there is a person who is just looking at your data and what you need. That's the question right be who would the organization could not find? And if you don't already have access to this, who do we need to help me arrange a meeting?

With them, just to be really practical about how you might investigate something like this, sometimes it's not something you do with code. Sometimes it's something you do with a person who has the intimate knowledge of the of the relevant data source.

And you, you all know just kind of covering that in general. Like what?

What some of the options are and then it's.

Hughlaugh, Jewel 18:27

Thank you.

All right.

And so then the next set of highlighted calls were a pretty

Large number of distinct units, although there are a large number of

Accounts of these, but these are a lot.

A lot of units associated with these different call count types.

And so there is in that low inc code, I don't think I think a spreadsheet that just transcribed, but.

They are a number per volume and so it could be the case that it's a lot of different volume size like number per microliter, number per milliliter, things like that.

But to have to with that like prescribed unit, I would double check and make sure and that at your unit source values if that is really the case. Maybe you already have and this is a most point, but anything typically above honestly above like a 3 or 4.

Distinct units I would not recommend take look and see where there's something coming from.

Because that's that's a. That's a pretty big diversity for a single measurement.

So just something to note there for the Columbia lab.

And in the last bit was kind of what I'll touch on a little bit where.

When you are trying to investigate.

What those might be?

What you do and what I have done actually adds with the Talk data is you can use like an aggregating like depending on the the SQL, because you're using in data bricks you have a function called collect set that's going to aggregate all of the values in.

A mass into the distinct values available.

It's like an array for in unique values.

And so I've done this for these top 10.

Like measurements at Talk have a lot of distinct units, and what you'll see is it don't.

Maybe I can zoom in here a little bit.

Maybe it's probably kind of weird.

So in our case you see that there's, OK, man, what's going on with my email? I don't.

It's just get over there for you here.

There, diversity is like the way of representing, you have like 1000 per microliter like calls per UL, per UL, and then you have some stuff that's clearly not a unit.

So this is that that's coming out of flow sheets.

And again, my apologies for the rest of this.

But we have the values in our unit field or the stuff that's clearly not not of 10.

So you have stuff that ends up in the unit field of flow sheets. That is a very much not intended to be there and so.

I'm guessing as you go through and do this exercise of trying to understand and collect the especially for these measurements that do you do see a lot of diversity in the units?

Yeah, but some of this type of unexpected data as well, so hopefully that was clear what I just described. But when I'm doing I'm guessing my table in this case by this measurement. So I'm saying everything that is this measurement I'm collapsing.

I'm counting how often it occurs.

I'm counting the distinct mapped units and here that I'm aggregating all of my distinct unit source values and so like, Yeah, you have a question mark. You have numbers you have, like, you're all over the place with this, with this particular measurement.

So.

The responses there are the just the data is messy, especially the flow sheet data that we have seen or that we work with with regard to the unit entry is messy. It is not well constrained and you end up with stuff that that you wouldn't expect to be there in.

Yeah, I have the query so I you can make an addendum to the query that I had sent you that does this.

At least the database, version of the most other flows will have some sort of empty aggregating function.

I don't know. Yeah, if you know offhand, collect set.

Is that also in Strimvelis?

I thought it might have been but.

Schmidt, Heidi E. 35:54
It's aggregate.
MSD is, yeah.

Hughes, David 35:57
Aggregates always takes the unique.

Schmidt, Heidi E. 36:00
I think so.
I did some of that when I had to do the lateral and split.

Hughes, David 36:10
OK.

Schmidt, Heidi E. 36:17
So.
But didn't have to then take column F and use it I could match it up to.
A specific unit somewhere.

That's I think that would be splitting and doing some sort of pattern matching.

Hughes, David 36:35
OK.

But yeah, we'll go on this on this last example, this is one where you do see.

So there's three units that we're supposed to have all.

Are they?

One would think that milligrams per deciliter is of course not molar per volume.

That's, but we're still mapping it as its own unit.

There's an element here if you have a measurement that specifies the type of units that you expect.

And the units do not fall under that.

Like what do you get?

And I think there's a kind of a case by case dependency here, but if you see something like that where?

All right, I have 3 distinct units and one missing, one of which follow my molar per volume expectation based on the low ink measurements and one of which does not.

But it's a valid unit, so to speak.

What do I do about that? I think that's.

Kinda going back to what Andrew was saying, there's kind of follow-ups like this where you can investigate further with.

The like the subject matter expertise in your area that know the lab systems in our lab and can tell you, like, This is why we have an inconsistent unit with our measurements because this is not the only case where that happens. I can tell you.

And even amongst the Olsby community, there is, I've seen multiple threads in the the Olsby forum talking about whether you should accept.

Measurements that are verified for measurements that they conflict with.

So sorry, go for it.

Schmidt, Heidi E. 38:20

My eyes were opened when?

When someone in our group pointed me to bower.org. I'm blanking on names because I'm so focused on one of the entities that you pointed out. And so like, if it even though it's a flow sheet.

It's a free account and I was just about to type that.

I was speaking, but this is the URL and so.

It's pretty much match the flow sheet name.

Then when you click on.

Create mass volume in urine. It could be low ink. Code 249-3, for example, and you scroll down to the example units it will have.

For that.

It will have.

This unit.

For.

That.

So that I know it's.

I didn't play around with like doing it as an API, but if it helps kind of with phishing, I thought I'd mention it.

Hughes, David 39:53

Yeah, thanks for pointing that link for those who haven't browsed in all the sites you think you have to create an account, you can do so for free and then you can search through the link relationships and kind of understanding.

It's a data structure or data model of of low ink ontology.

So thank you for that.

And we were specifically what Heidi was talking about was this one.

This one.

Which I think this is the one that you were referring to, right?

I don't remember.

That's exactly the one.

Or not any one.

Schmidt, Heidi E. 40:28

I guess that it might be similar.

Hughes, David 40:30

Yeah.

Schmidt, Heidi E. 40:39

There are two entries with two different low ink codes, but this that one says that it's milligrams per deciliter.

Hughes, David 40:39

Course that would be a mass ratio then.

So I think it's probably the other one that that you're referring to. It's OK.

Schmidt, Heidi E. 40:43

OK.

Yeah, yeah. Mass volume, it says.

Sorry, but the that's what I read and verify.

Hughes, David 40:52

Assume.

Go for volume.

Williams, Andrew E. 40:58

Wondering if you or Paula had advice for people who want to do any kind of browsing of linked whether the it's worth the effort to, you know, download mains and use it, or you get most of what you need without it. Just if anybody's kind of follow.

Optimal ways to browse for stuff outside of I guess what's provided via Athena for relevant mappings for bio.

Do you have any thoughts on that?

Talpin, Paula 41:02

Yeah, it's possible to download the link sources directly into the database and work with those sources there.

Another option is allowing search, but you should be organized in the link system.

Also.

So it's the two best ways I recommend.

I would recommend the first option where you just leverage the sources of flowing that they provide.

As on their website.

There is download link option. You can get this for instance.

I just can put it into the chat.

Williams, Andrew E. 42:10

Really, you when you say 'is that what you're talking about?

Talpin, Paula 42:10

I didn't get main.

I'm just not the logic source files.

Williams, Andrew E. 42:24

I see.

Talpin, Paula 42:25

What? Yeah.

And it's quite useful for work within the database when you don't need to search on the web.

Williams, Andrew E. 42:34

I see. Thanks.

Hughes, David 42:40

Right, but yeah, what I would recommend checking out.

So I found it to be helpful to go kind of to the source origins, both in the case of Snow Med and low ink.

And other ontologies to browse and kind of interrogate them as well, especially for very particular use cases like investigating the units of a precise link measurement on.

Definitely worth knowing that those resources are there and taking a look at them if you if you need.

OK.

That's so that the extent, I just also wanted to give a little bit of an example on.

diversity that you do see in pop up in these unit fields.

I've already mentioned we talked a little bit about inferring units.

You'll see that I do that in a couple of these cases.

So this is a.

This is one where a percent is expected, and so the ones that are missing, I infer that it's percent.

You'll notice this is one that I need to probably follow up more on.

There are actually almost the same concept, one of which is specifically through HSC.

One of which is not specified.

This one actually has a possible unit of milligrams per deciliter and it also has a missing unit.

Which I'm guessing is coming out of our legacy data.

Remember I would be interested.

But this is one that I will need to take a look at in more detail but.

Yeah, I guess that all have in my about.

About units for now, unless there are any questions on the sheet or other otherwise.

And then I can go to the last topic if I not.

Schmidt, Heidi E. 44:02

Do you have the logic for the collect set? Even if it's not gonna translate into?

All right, maybe.

Hughes, David 44:38

Yeah, I'll put it in the chat.

It's the same query actually that I used before.

but actually no, it's slightly different.

It's, I'll send this to you because I did break it apart.

So just for speed reasons.

So that's that's the query.

Can the collect set give.

Actually that one.

Sorry that has any AC which you don't wanna run.

Let me modify that for you real quick.

That was an old version.

So.

I think the most difficult translation will be to like SQL Server probably, but I don't know how they handle.

Then the handling of array aggregation on SQL Server is always slightly different but.

breakable should be pretty quick.

OK.

So I'm gonna start my sharing, so maybe I'll take a quick pause and say.

I can go through.

Really, this like of a gap analysis and a dashboard and my thoughts on that one that would show Epic specific data and epic tables and names that we would then have to exclude entirely from this recording.

I think we could chop the recording probably at this point.

I don't know though.

I'm looking a little bit toward them, kind of showing Andrew and others at curatorial here with this. Would this be valuable or should we?

Should we keep on to this and try to do with something more full-fledged in the future office hours?

Williams, Andrew E. 45:17

You are the best judge of that.

I think it's really exciting.

And as long as people don't.

Take it to be anything other than what you show them.

I think it's really OK.

I think when we've hesitated in the past to show things that are in progress is because people might.

Not understand the full implications of being in progress, and you're the best person to kind of judge that risk and what's what's there is.

What? What's your thought?

I'm, I'm browsing the curatorial back to you.

Hughes, David 47:09

I think it's.

I think it could be useful, so I would I would.

I would error on the side of showing it and discussing the gap analysis, so I will, I'll go for it, OK.

Williams, Andrew E. 47:19

OK.

Hughes, David 47:23

Then let me share my screen for those of you who have already seen this.

The data dashboard that we have, I should update just a second here that we just together for on the central cloud.

This should look familiar to you.

But I'll kind of give a very brief overview of what we've been doing here.

Some of the questions that that I get a lot is or that's kind of comes up as you're trying to set up a study.

Are you trying to investigate a set of samples that we're looking for?

Let's say CEMO data and we know that we're capturing it in a hospital, but we have nothing in camp.

One, why is that?

Where is it and how can we make sure that we get it into camp?

And so in that case, right.

Like I can get my device domain and I'm already that I have a very low device coverage, at least at Tufts, but we're capturing. We're capturing a small percentage of our devices at the moment.

Working on that.

This kind of gives you a hint if you look at the number of rows of potential devices and you look at the number of actual device exposure data, how much are you missing?

And then that's contextualized a little bit by looking at how many unique concepts you have between camp and epic.

How many unique persons you have and how many units you're actually capturing on.

In this case, we have a lot of people in and a lot of exits, so it's the issue is not coming from just missing people, it's actually coming from missing mappings and what you see if you go to something like ECOM3 I know I can search that. Now in this table and I have a list of data and again.

This is not.

This is maybe a bit misleading in the sense that I'm looking at flow sheets and these are not direct device entries.

They also indicate that a device is being used.

It's like a secondary way of understanding a device.

But we have, at least in our Taha epic instance, we have a list of flow sheet data related to ECOM3.

But you'll have if I could show them, there's, there's nothing with using.

So if like this data ID we don't have any ECOM3 data yet mapped and that's just one example like.

If you were some it's not useful for something you can do so here and you can see where but then exists in your setup versus where it exists in your source.

And you can do the same. What I've also done is the second table I create.

I actually just took all of the unique terms in the vocabulary in this case the device domain, but per domain and I've then mapped like how often those terms appear in epic versus in setup.

Oh, you can actually search this table for like, I don't know if proteins genome be in here, but like, uh, prosthetic or something like this. And you can look and and compare at a term by term level how often something appears.

And then this is just broken down by the various domains.

So you can do the same for drugs.

So when it's interesting actually in the drug domain is we are actually have additional granularity beyond what Epic has in this whole load and just actually it's filtered. So that's why it's not loading.

And this is kind of a weird artifact from the way that handles its medication mapping to RX norm.

So we still have a lot of drug at hand by far.

So we can investigate and there's potential rationale for that.

So that we can though you could then search things there and you could use like.

If you do something like this, you'll see a lot of different options from EPIC, but then I can.

Well we can strip entries and you can again these three tables are you with.

And again, I did the same thing here where I have different drug terms of interest and you can compare those between Epic and setup.

Some for each domain so.

This dashboard is intended to serve as a little bit of like a first step into understanding.

On either a term level or a condition based level.

The the disparities or the gaps that you have between your underlying source data and your setup instance.

And it does so by essentially pre-populating a set of three tables per domain and it uses in our case other staging tables that we create during our ETL, or really the raw epic clarity tables.

To get into.

What when you're looking for?

Like kind of this general overview that you can search through and identify where things might be missing.

So you and again, this is the first version, but the intent is to better understand what things were missing and identify why we're missing them, or how we can potentially map them and what part of the logic that exists. So we have like references to the log right where this actually is drawn from or where that code would need to be changed.

So I see this is a another tool to help update and upgrade an ETL to capture more elements of interest.

Williams, Andrew E 10:13

Benjamin Jend.

It's my first time seeing that.

That's so exciting to read you.

Make things happen so well and so quickly.

It's kind of takes my breath away.

Some of the content is that we're a site that is looking to be as comprehensive as possible with respect to the available data, and we also played this role in helping to expand the vocabulary to capture things that currently can't be mapped.

Both of these things are probably considerations for people who are thinking about.

The value of this tool for your work, if you're at a site that's primarily trying to capture what's easy to get in the broad domains and then the specific things that we're requiring for bridge to AI.

There's less probably of an interest.

Once, comprehensive representation possible, and it's a non-trivial amount of work to get after things.

It's like the case like of things that people are interested in capturing but they're not.

Identifying new gaps in the setup vocabulary that need to be filled in order to represent them.

And then there's an additional piece of work that we can work we can pursue together on things like.

Rubalshilov, Ruben 10:21

Jend.

This is to really what I'm working towards here at UCSF.

So it's it's pretty close to my heart and we've struggled with it because of the domain transformations that happen as you pull data into setup.

So that's my my really nice question is you don't use provenance information for this. You use broad concept level like when you populate those percentages.

I'm based on the broad concept, like a medication ID or like an epic ID.

To strip concept and and and like I've tried to do it.

At all with the provenance information in place, because it's just very difficult to capture all of the filters that the ETL actually accounts for. The missing data and things like that. To do it at the concept level and the fact that you move across domains as you.

Transform from one table to like it in tables in setup. But I would be very excited to talk to you about this.

So I see this is like on top of my list as we improve our instance. So I really wanna hear your thoughts about it.

Hughall, Jend 10:23

Yeah, too, absolutely.

I think you're like the real on the head with the the biggest issue, right?

And that's why I need needed to use a lot of our intermediate like our staging tables, because the way that our ETL is structured is the majority of the filters, with the exception of the mappings themselves, happen at the staging table level.

So like we have to make sure that we have an encounter that that's good. And so also these numbers are in context like.

We're assuming that we've captured the encounters appropriately in OMA.

And we're saying that the encounters that make it encounters that make it into our staging table have already been vetted upstream.

So like like the long filter and we don't have all of these other things that like, was the record disqualified or like was the record?

I don't know invalid for whatever reason.

And these like get chopped out or at it in the entry to the staging table.

So there is some.

Upstream logic that populates this in the the numbers that you're seeing are.

Not total counts of like perfect provenance.

Everything in the diagnostic table.

It's like what exists in our staging table that then we're able to capture assuming that the staging table is more or less like our ground truth.

So there's additional of interest needed for us, but you're actually right. It's very difficult to.

Take a look at just an epic table which has all of its discrepancies and columns that that reggie things missing dates and all this, and then be able to then compare that with something like an empty table where there's a lot more constraints.

On certain table versus the like the missing and things like that.

So absolutely would love to talk to you more about this offline.

I think you're missed. Andrew.

Williams, Andrew E 10:25

In context how many sites do something?

How in a table that's persisted.

So you don't have to re-query the original table multiple times depending on which part of your ETL you're working in.

You want, you want yourself a lot of re-querying of the same data by creating that first match up table which I think is the staging table. You're you're referring to in general.

There's other approaches that don't like that and.

More than one, but that's a least category of the saving a lot of requiring work.

And I can just wondering how many.

Sites do their ETL that way, if you're doing it yourself and you're not getting it from someone else, how many sites do that initial mapping table and then draw from that?

You do do know at UCSF or other stuff.

Hughall, Jend 10:26

Might does as well, at least to some extent, with the staging tables in EDC.

Williams, Andrew E 10:26

It's.

It does save both on processing and potential to do something that's meant to be equivalent in an unreasonably different way.

Would you you have to be the same algorithm in the same way or the same source or like there's some slight difference in the code because you're you're not using the same code in the same source.

So there's a couple of different problems that it solves, but I know other people feel like it's optimized in other ways, so.

Anyhow, then, however?

Hughall, Jend 10:27

Yeah, so the the the way that we've structured the dashboard.

It's a single Python file. Will run the dashboard for you, the data to actually support the dashboard. This is all just created with SQL queries and the like.

It's just.

There's a tables that have a particular data structure, but you can populate them however you want, and you can. The intent is that it's flexible to accommodate not just specific epic tables or epic staging tables, but you could.

You could give.

Content from other source system tables as you as you would like, so hopefully it's fairly extensible, but I can share the code both for the dashboard and for the.

Dashboard for sure.

For the underlying logic that goes to epic, I could maybe put it on epic user web or do something like that to do something that so.

The dashboard shouldn't have any epic-specific stuff like.

So we should be good there.

Schmidt, Heidi E 10:27

Yeah, I mean if the queries that are being run underneath are.

Epic user web.

Then I'm assuming that like we have to plug in whatever SQL logic.

Would be anyway.

And then if someone didn't have an epic site, they'd have to plug in that their source setup.

Hughall, Jend 10:27

You just have to be able to to create three different CSV files for the domain with a particular structure, and it's more or less like what is the table that you're talking about.

What is the the the the the you're talking about and the distinct concept in that and then the counts of people and visits and the concept like counts of concept level?

Maria's Kathleen Nesterov 10:28

Play me a song.

Hughall, Jend 10:28

That's really the main thing.

And then if you can create a setup query that generates that then.

Then the dashboard will show it for you essentially.

Williams, Andrew E 10:29

It would be wonderful.

I don't know that we wanna argue this.

It would be great.

Some of the metadata about.

The data set that we end up producing from the project, if we could estimate the overall coverage of what got mapped in the relevant domains, in addition to trying you know what we did to validate those mappings and ensure their quality and so on so forth just to say.

Because it's possible to have a ton of things not really cover a fairly small amount of what was potentially available in a given area, and I think that's often the case when you're doing a big ETL. And so having some characterization of that would.

be fantastic from as many sites as we could get it and this would really enable us to understand that.

So it's a really.

Jend, have you ever, you know, you had a whole bunch of notes about Rogers original ETL?

How you were added those or communicated those to him?

Hughall, Jend 10:29

To yeah, I need to.

I need to actually send an e-mail with that I was hoping to.

Kind of refine my my notes on that, I did want to to Sarah I think.

I think Sarah should have been in Seattle and I've maybe distributed them to some other people like individually, but haven't put them in a format that I don't know, like a slightly less blunt.

I think I was maybe a bit crude in my in my evaluation of certain things and I want to make it a little more nuanced so.

I think that's once I can I can do to refine those notes, then I will send them to him and provide that probably.

So.

Williams, Andrew E 10:29

It's great I think.

Yeah, so other people are using it a fair amount. And so where there's knowledge to be gained about things to be alert to and perhaps correct.

Be good to get that.

Maybe I don't know anybody else. You've shared it with wants to kind of help in wordsmithing some of it is really descriptors associated with the things that you note that you want to or is it?

Hughall, Jend 10:29

I just kind of reported the set of issues that I had created.

In our GitHub repository that were just like, there wasn't a lot of context around the issues and it like, I think having additional because now also with the updates that I've made some those issues were created and resolved.

I think that there's additional information that I can add and kind of contextualize them more.

So I think what you want was nice, but there's more to be said.

Williams, Andrew E 10:30

Yeah, I guess if there's a way anybody who's using.

Rogers ETL as a with to someone. I. Whether it's building on what you did or independent from it, that's a good thing I guess for us.

Alright to do.

And I'm great that you've you've done that, but anybody else who's done it either.

Building on what you did or independent of it. I think it would be great for us to encourage that.

Hughall, Jend 10:30

Awesome. Well, thank you all very much.

Hopefully this was a useful useful session.

And has focused is talking to you all.

Actually, I'm gonna be probably not next week, but hopefully in a couple weeks I'll be back.

Yeah, thank you all very much.

Williams, Andrew E 10:40

You, Jend.

Hughall, Jend 10:40

Here's a nice rest of your day and.

Williams, Andrew E 10:40

Hope you have a great great week.

Hughall, Jend 10:40

Thank you. And Katari, please reach out.

Let's let's chat more about this. I would love to get your thoughts and also love to collaborate more on this because it's a work in progress and I think there's a lot of potential to do some cool stuff on.

Rubalshilov, Ruben 10:40

That's great.

Thank you, Jeff.

This was great.

Hughall, Jend 10:40

Alright.

Yep, thanks all.

Alvarez, Marta stopped transcription

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK

BK