

Running a Pipeline

Saved and imported pipelines can be applied to different data sets in other projects. Prior to running a saved or imported pipeline, the data needs to be imported and the sample attribute need to be specified under the [Data](#) tab. Pipelines that include tasks requiring sample attributes (e.g. GSA or ANOVA) will not run unless the sample attributes have been specified beforehand. All saved and imported pipelines are available for all users on a Partek Flow instance to run. To run a pipeline:

1. Click on a circular data node under the [Analyses](#) tab and expand the **Pipelines** section from the menu on the right (Figure 1). The [context-sensitive menu](#) will only display pipelines that can be applied to the data type of the selected data node

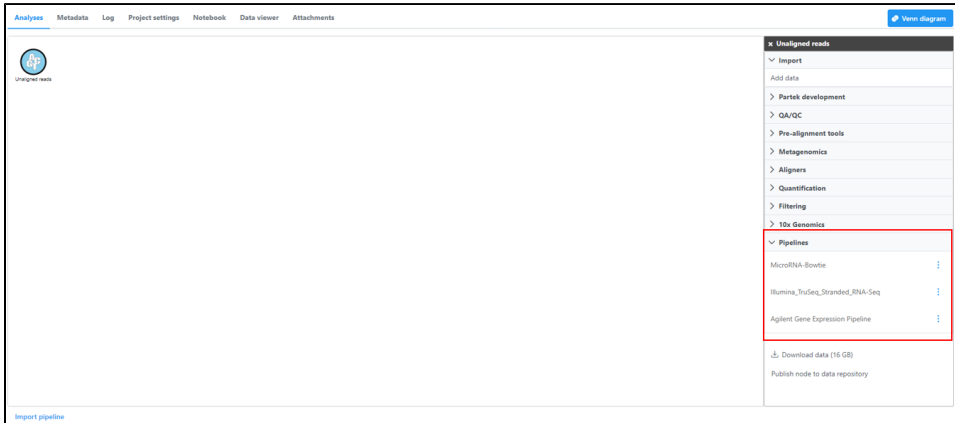


Figure 1. Loading a pipeline. In this example, the context-sensitive menu is showing all pipelines that can be applied to Aligned reads

2. Click on the pipeline name from the menu. Note that hovering the mouse over the pipeline name will show the description (if one was added) in a pop-out balloon
3. If further settings need to be specified for any of the tasks, you will be redirected to a task-specific page. For example, if the chosen pipeline includes a STAR alignment task, you need to specify the species and STAR aligner index (Figure 2). Other tasks that require additional settings include *Quantification to annotation model*, *Differential gene expression*, *Variant detection* and others. For each task, specify the requested settings and click **Next**. Other task settings that were specified when the pipeline was saved (e.g. alignment parameters) will be applied automatically.

Figure 2. Additional settings may be requested for certain tasks when loading a pipeline. In this example, the STAR alignment task is requesting an index to align to

The additional settings requested for certain tasks allows for flexibility in how the pipelines are used. For example, a pipeline that was initially created from a project on human data can be reused on data from another species by specifying a different *Assembly* (Figure 2). Other examples include performing quantification using different annotation models and customizing statistical models for different study designs.

Once additional settings have been specified for each task, all tasks will be queued and the jobs will run sequentially (Figure 3). The status of queued tasks can be monitored under the [Log](#) tab. If you have set up email notifications (see [Personal](#)), you will receive an email when the pipeline finishes.

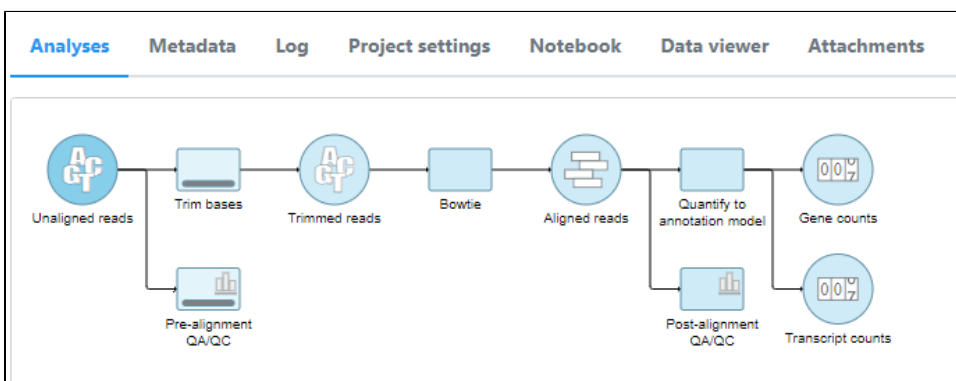


Figure 3. All tasks in the chosen pipeline have been queued

Additional Assistance

If you need additional assistance, please visit [our support page](#) to submit a help ticket or find phone numbers for regional support.

[« Making a Pipeline](#) [Downloading and Sharing a Pipeline](#) »



Your Rating: Results: 29 rates