

# Validating a Clinical Cancer Exome Assay: A real world perspective

Natalie Bir, Sean D. McGrath, Ben Kelly, Patrick Brennan, Elaine R. Mardis,  
Richard K. Wilson, Catherine Cottrell, Vincent Magrini

Institute for Genomic Medicine

## background

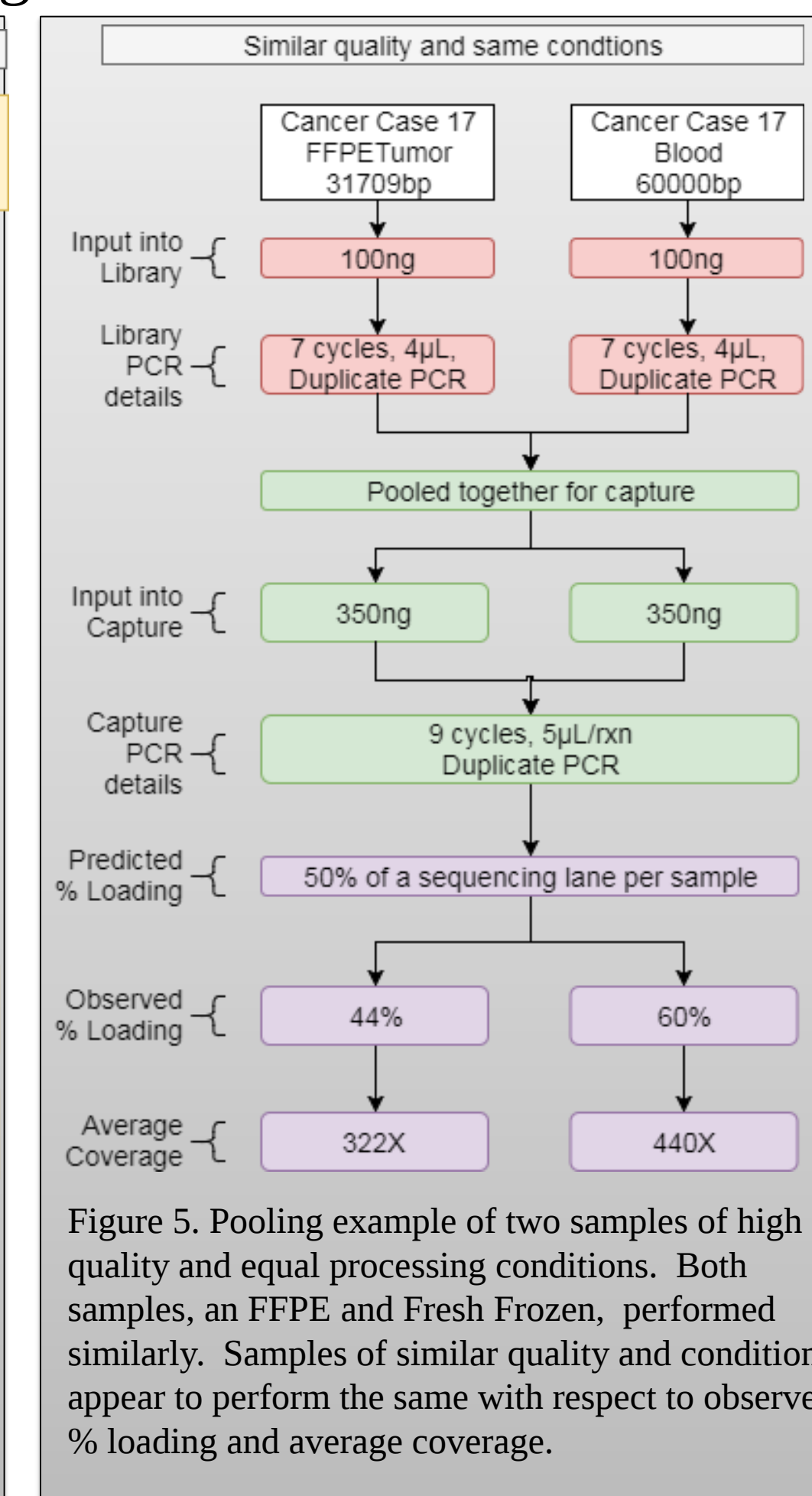
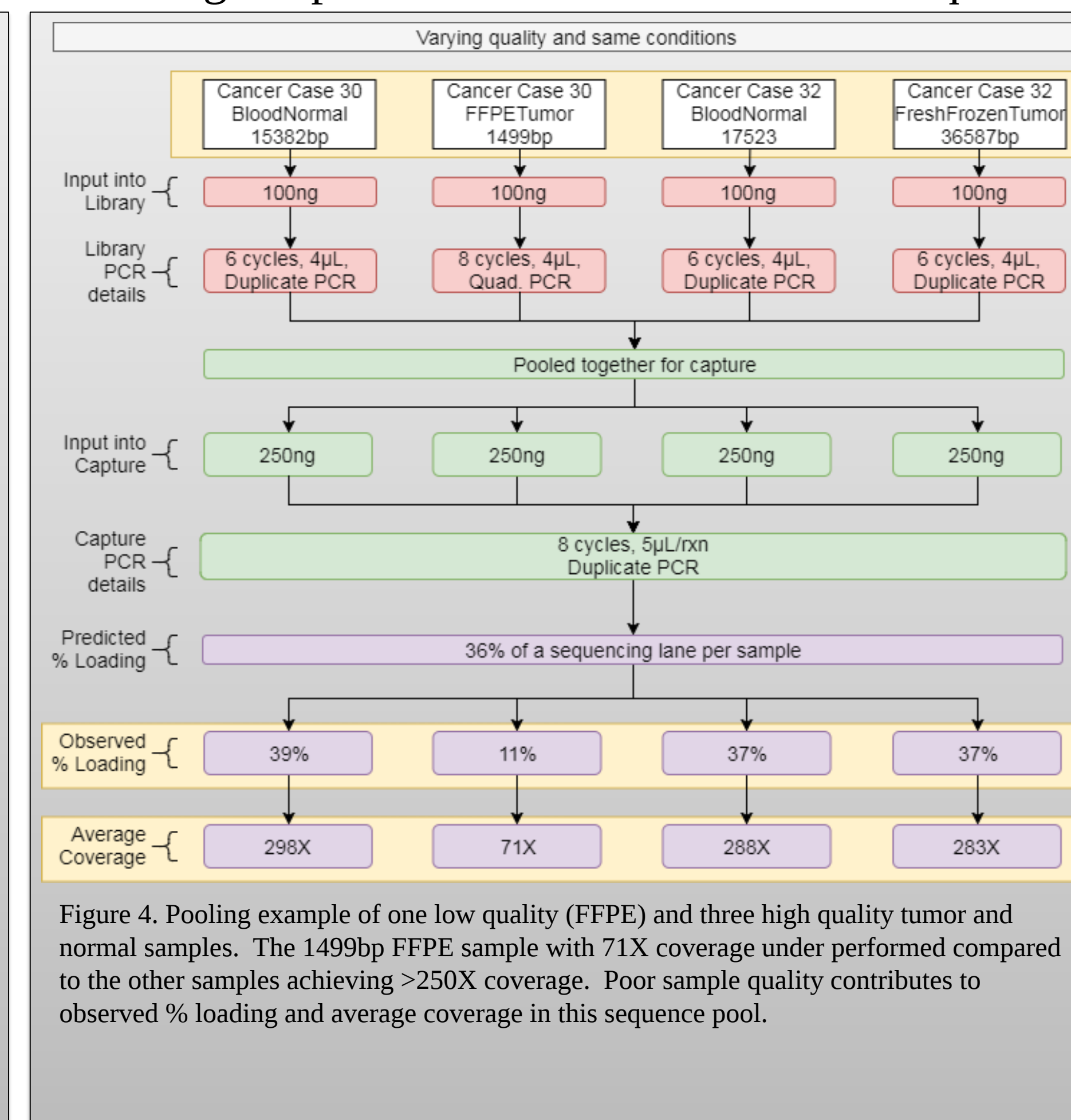
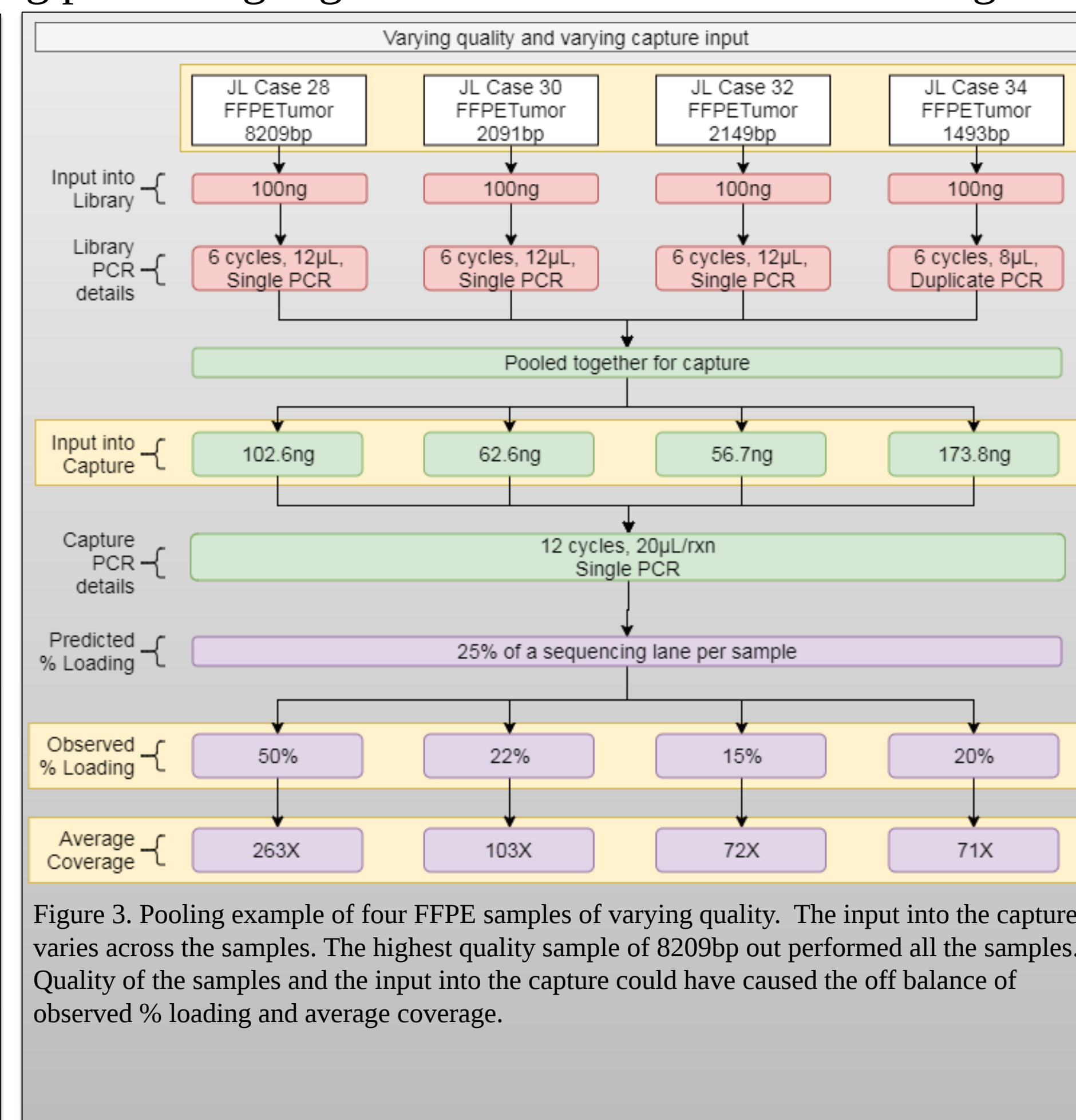
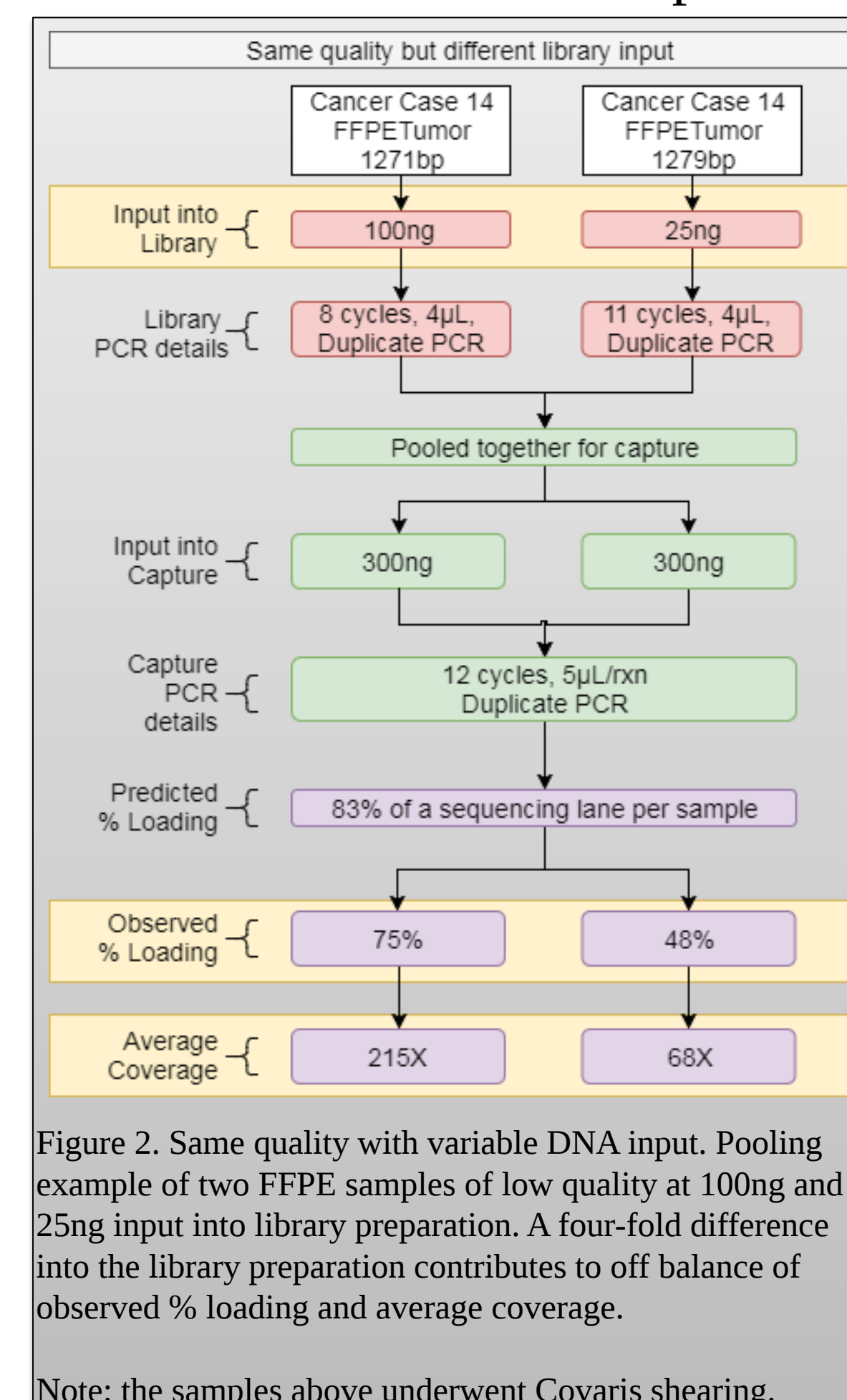
To date, IGM has processed 65 cancer cases using enhanced whole exome sequencing under a research protocol. Here, we report on the transition of the research protocol into a clinical assay to be performed in the CAP/CLIA laboratory. The results from 59 cancer samples highlight the need to evaluate sequencing metrics and examine trends associated with sample quality in the pursuit of a standardized clinical process and analytical pipeline with defined pass/fail metrics, while optimizing the sensitivity and specificity of variant detection.

56 samples were processed using the NEB UltraII FS (enzymatic shearing) library prep method and the IDT xGEN capture method using the IDT Research Exome Panel plus the IDT CNV backbone. Three samples were processed using Covaris fragmentation and NEBNext UltraII with the IDT xGEN capture method.

| Sample Type         | No. of Samples | Range of Genomic DNA Peak Size |
|---------------------|----------------|--------------------------------|
| Fresh Frozen Tissue | 21             | 10,777bp – 60,000bp            |
| Blood Normal        | 26             | 13,029bp – 60,000bp            |
| FFPE Tumor Tissue   | 12             | 1,271bp – 31,079bp             |

## various sample processing scenarios and pooling

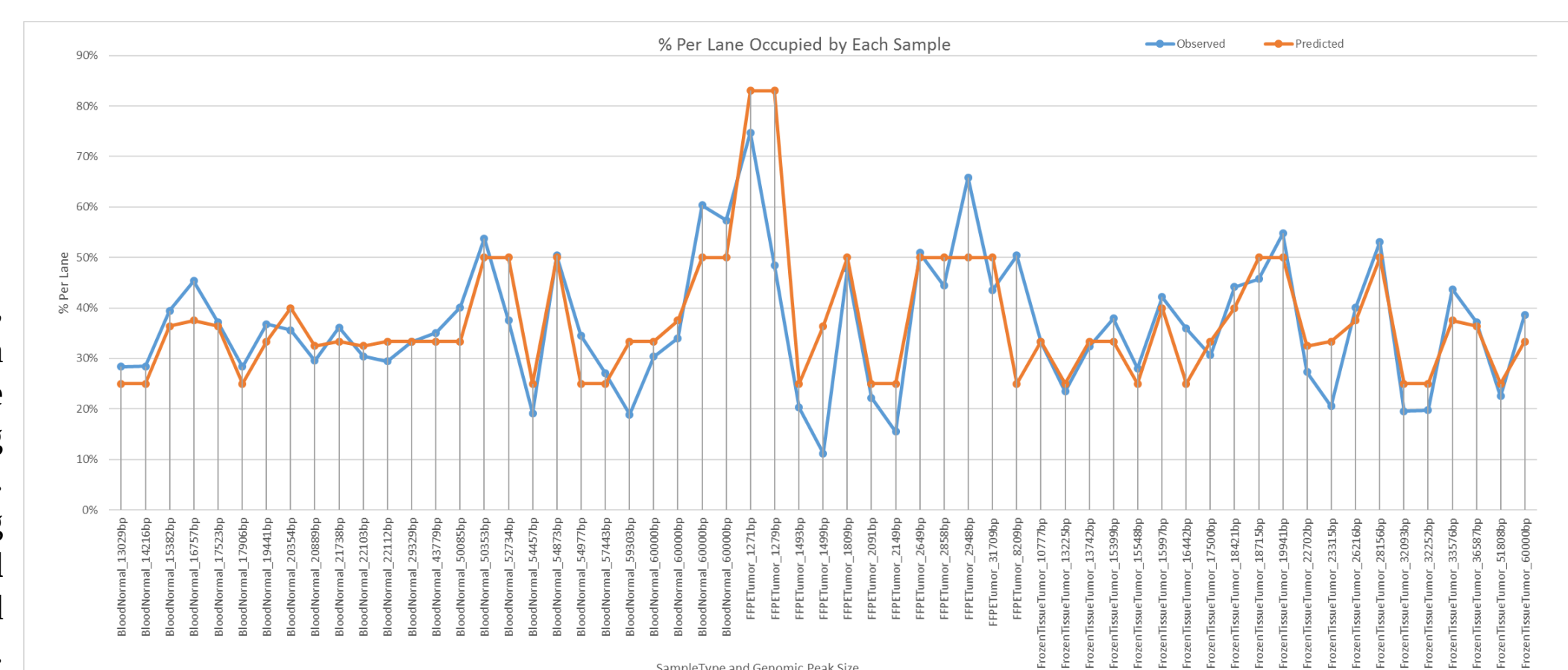
The processing profile highlighted below aid our understanding of standardizing the protocols for more consistent sequencing results.



## predicted vs. observed % per lane

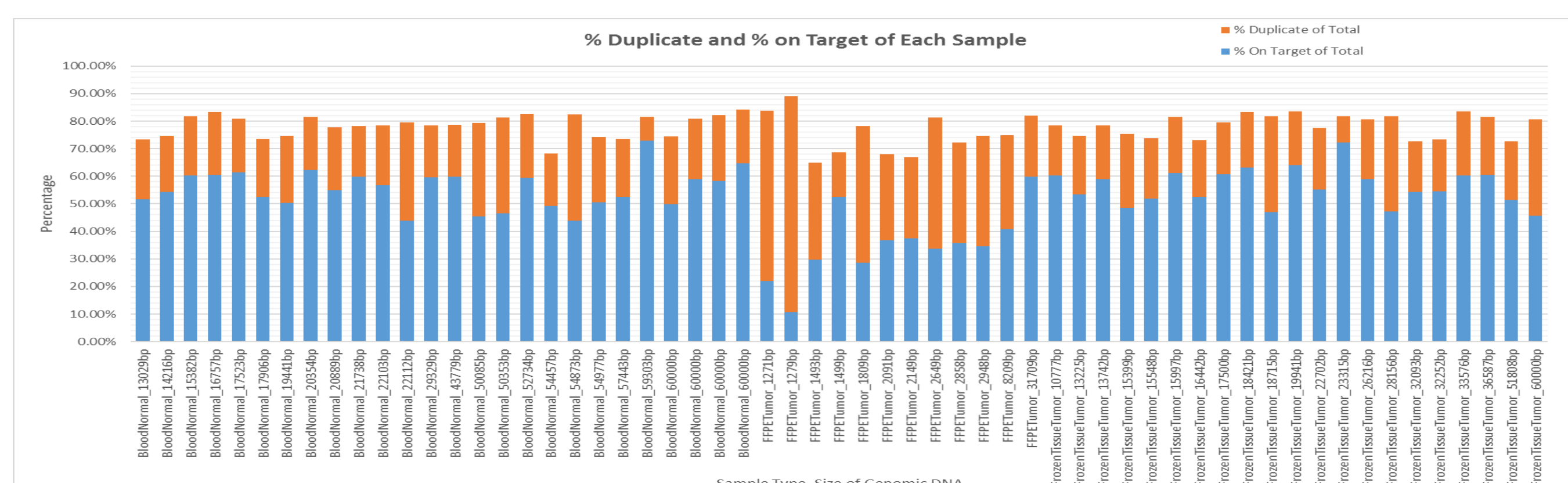
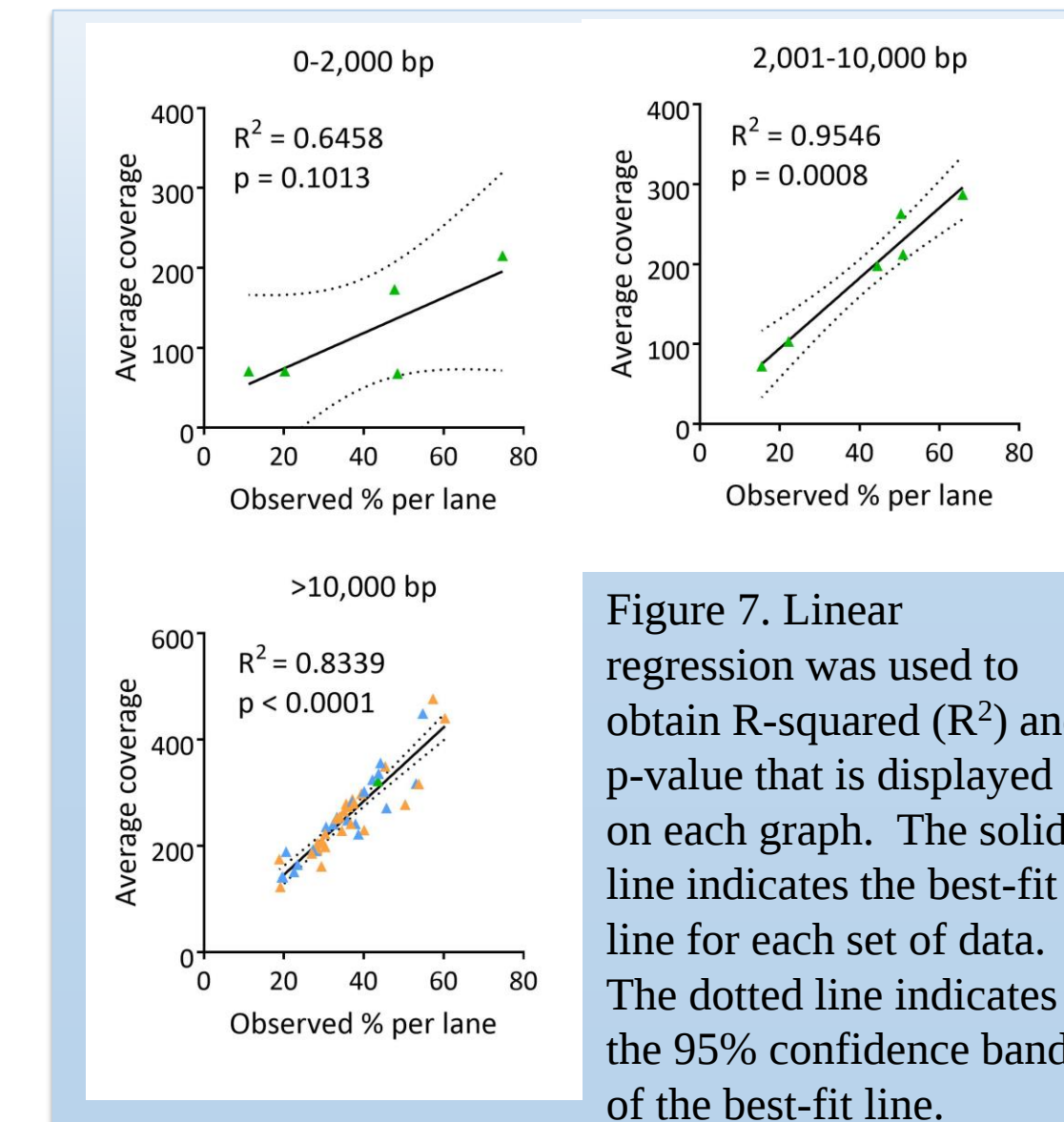
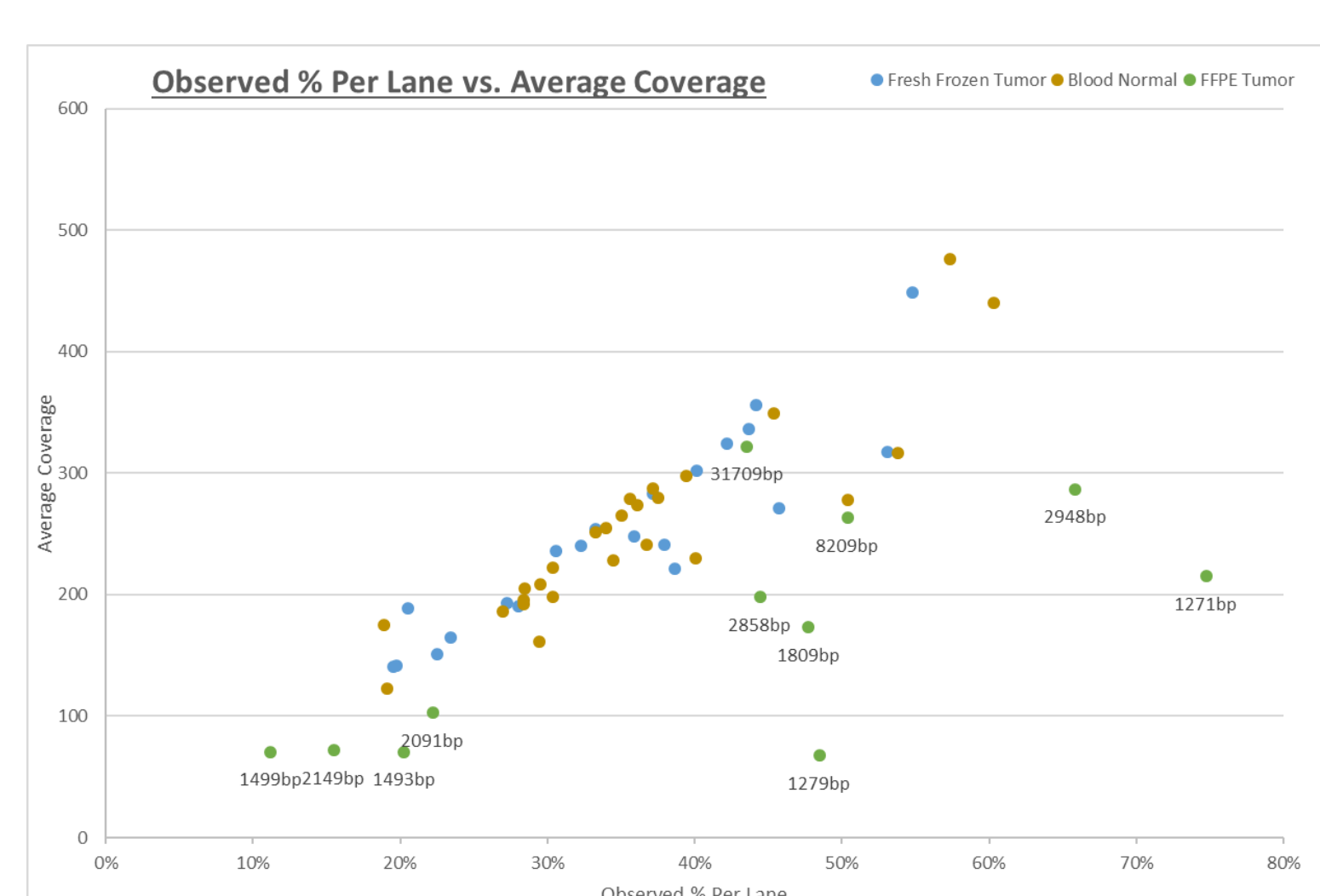
Samples may be run across multiple lanes to achieve higher coverage. The predicted % per lane of a sample is calculated by adding up the % each sample occupies within **each** lane. The observed % per lane of a sample is dependent on actual total number of reads. Below illustrates loading a set of samples on the Illumina HiSeq4000 and the predicted and observed % per lane calculation.

|                                                     | Lane 1                                                                                            | Lane 2                              |
|-----------------------------------------------------|---------------------------------------------------------------------------------------------------|-------------------------------------|
|                                                     | Total of 360 Million Clusters                                                                     | Total of 360 Million Clusters       |
| Sample 1                                            | 120 Million Clusters<br>33% of lane                                                               | 180 Million Clusters<br>50% of lane |
| Sample 2                                            | 120 Million Clusters<br>33% of lane                                                               | 180 Million Clusters<br>50% of lane |
| Sample 3                                            | 120 Million Clusters<br>33% of lane                                                               |                                     |
| <b>PREDICTED % Per Lane of a Sample</b>             | <b>OBSERVED % Per Lane of a Sample</b>                                                            |                                     |
| Total of Two Lanes<br>Total of 720 Million Clusters | Based on Total Reads from All Lanes<br>Total Reads/2 = 1 Cluster<br>360 Million Clusters / 1 Lane |                                     |
| Sample 1<br>300 Million Clusters<br>83% of lane     | Sample 1<br>538,205,644 Total Reads<br>75% of lane                                                |                                     |
| Sample 2<br>300 Million Clusters<br>83% of lane     | Sample 2<br>349,020,246 Total Reads<br>48% of lane                                                |                                     |
| Sample 3<br>120 Million Clusters<br>33% of lane     | Sample 3<br>432,025,233 Total Reads<br>60% of lane                                                |                                     |



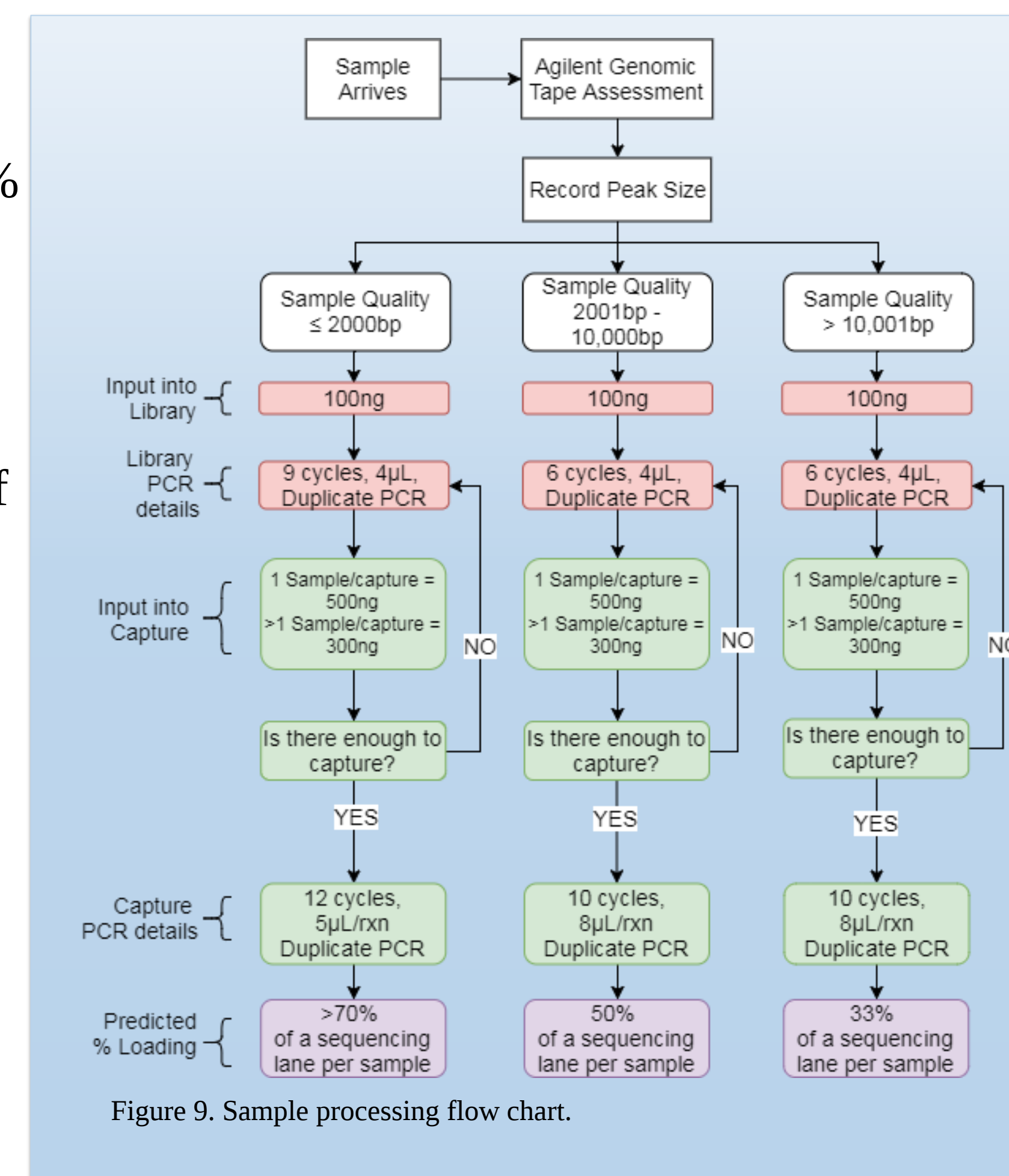
## data

We observe a strong correlation with observed % loading and average coverage as seen in Figure 6. If you break down the samples by sample quality we identify three different categories (Figure 7). Although this is a small sample size, this is a good start into understanding how the samples need to be loaded onto the sequencer.



## results

Sample quality, genomic DNA input into the library preparation, library input into the capture, and % loading per lane on the sequencer all are important factors. Figure 9 is a recommended sample processing flow chart. Cycle conditions are based on averages of what we have done in the past. We will adjust the processing parameters as our sample types (FFPE/frozen tissue/blood), tissue sources, and numbers increase. Optimizing the protocol is an ongoing task.



## acknowledgements

Thank you Dr. Katherine Miller for aiding in the statistical analysis of the samples. Thanks to the Genomic Services group for processing the samples and generating the data.