



filter - DNBSEQ Eukaryotic Strand-specific Transcriptome Resequencing

Report ID: F25A430000584_HOMufgjR

Date: 2026.01.26



Table of Content	2
1. Result	2
1.1. Project Information	3
1.2. Data Production	3
1.3. Quality Control	4
1.3.1. The distribution of base percentage and qualities along reads after data filtering	4
2. Methods	7
2.1. Experimental Procedure	7
2.2. Bioinformatic Analysis Workflow	8
2.3. Parameters for Data Filtering	9
3. Help	10
3.1. FASTQ format description	10
4. References	11

1. Result

1.1 Project Information

The basic information of the project is shown below:

- Project ID: F25A430000584_HOMufgjR
- Product name: filter - DNBSEQ Eukaryotic Strand-specific Transcriptome Resequencing
- Sample size: 13
- Library type: DNBSEQ Eukaryotic Strand-specific mRNA library
- Sequencing Platform: DNBSEQ
- Sequencing read Length: PE100
- Clean fastq phred quality score encoding: Phred+33

1.2 Data Production

After sequencing, the raw reads were filtered. Data filtering includes removing adapter sequences, contamination and low-quality reads from raw reads.

▲ Table 1 Statistics of clean data ▼

Sample Name	Clean Reads	Clean Base	Read Length	Q20(%)	Q30(%)	GC(%)
C1	24,119,764	4,823,952,800	PE100	98.57	94.37	47.95
C2	24,081,685	4,816,337,000	PE100	98.59	94.46	48.14
C3	24,164,297	4,832,859,400	PE100	98.62	94.55	48.05
C4	24,239,468	4,847,893,600	PE100	98.68	94.76	48.12
C5	24,097,076	4,819,415,200	PE100	98.59	94.44	48.12
C6	24,091,360	4,818,272,000	PE100	98.56	94.31	48.32
C7	24,129,069	4,825,813,800	PE100	98.65	94.67	47.26
V1	23,954,942	4,790,988,400	PE100	98.45	94.08	50.26
V2	24,275,001	4,855,000,200	PE100	98.38	93.84	50.56
V3	24,085,859	4,817,171,800	PE100	98.35	93.67	49.32
V4	12,295,063	2,459,012,600	PE100	98.69	94.81	49.88
V5	24,091,790	4,818,358,000	PE100	98.53	94.42	48.96
V6	24,206,123	4,841,224,600	PE100	98.46	94.14	48.68

- Sample Name: Sample Name;
- Clean Reads: Clean Reads;

- Clean Base: Clean Bases;
- Read Length: Read Length;
- Q20(%): Proportion of Q20;
- Q30(%): Proportion of Q30;
- GC(%): Proportion of GC.
- For better view, datas are formatted.

1.3 Quality Control

The quality of data was examined after filtering.

1.3.1 The distribution of base percentage and qualities along reads after data filtering

In the left figure, x-axis represents base position along reads, y-axis represents base percentage at the position; each color represents a type of nucleotide. Under normal conditions, the sample does not have AT/GC separation. It is normal to see fluctuations in the first several bp positions, which is caused by random primer and the instability of enzyme-substrate binding at the beginning of the sequencing reaction. In the right figure, x-axis represents base position along reads, y-axis represents base quality; each dot represents the base quality of the corresponding position along reads, color intensity reflects the number of nucleotides, a more intense color along a quality value indicates a higher proportion of this quality in the sequencing data.

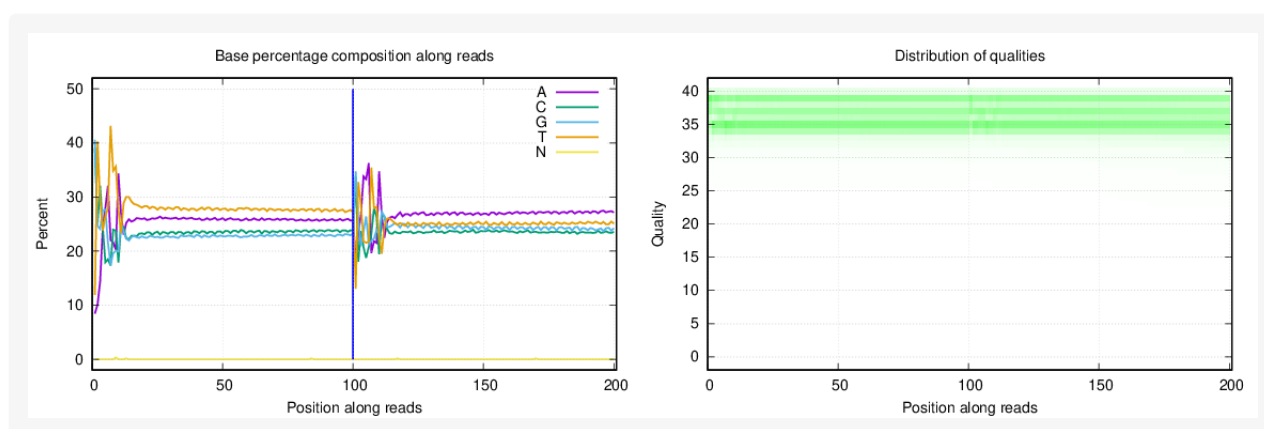


Figure 1-1 Distribution of base percentage and qualities of sample C7.

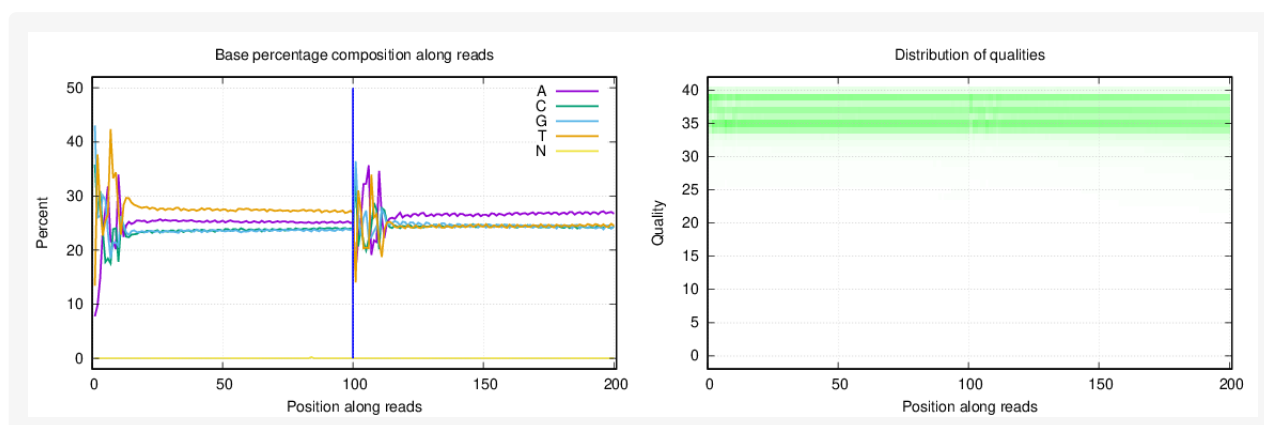


Figure 1-2 Distribution of base percentage and qualities of sample C6.

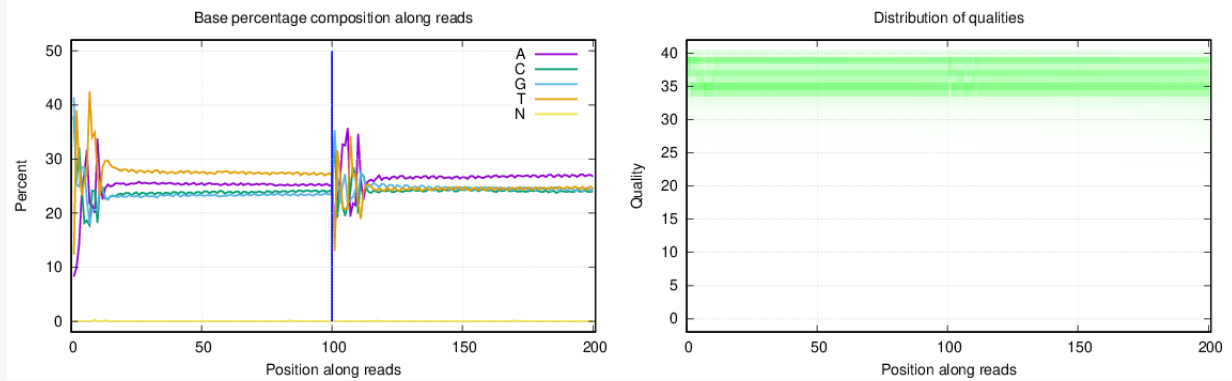


Figure 1-3 Distribution of base percentage and qualities of sample C5.

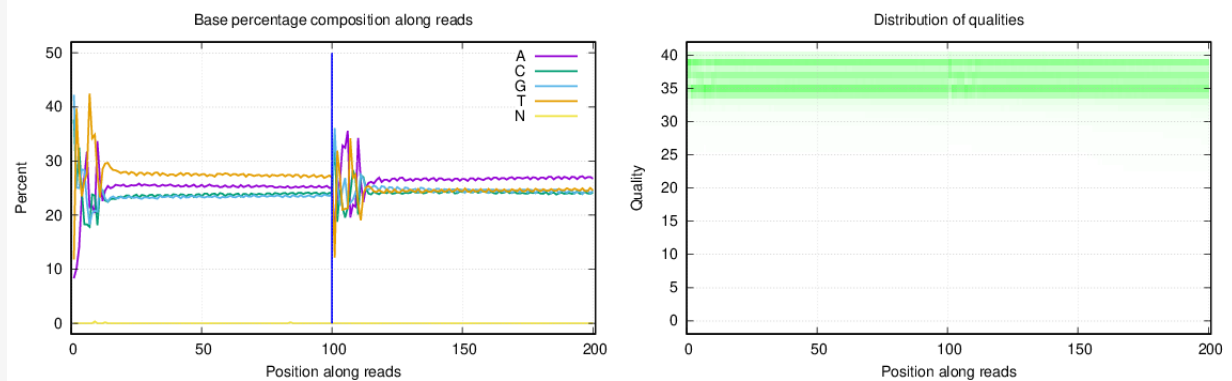


Figure 1-4 Distribution of base percentage and qualities of sample C4.

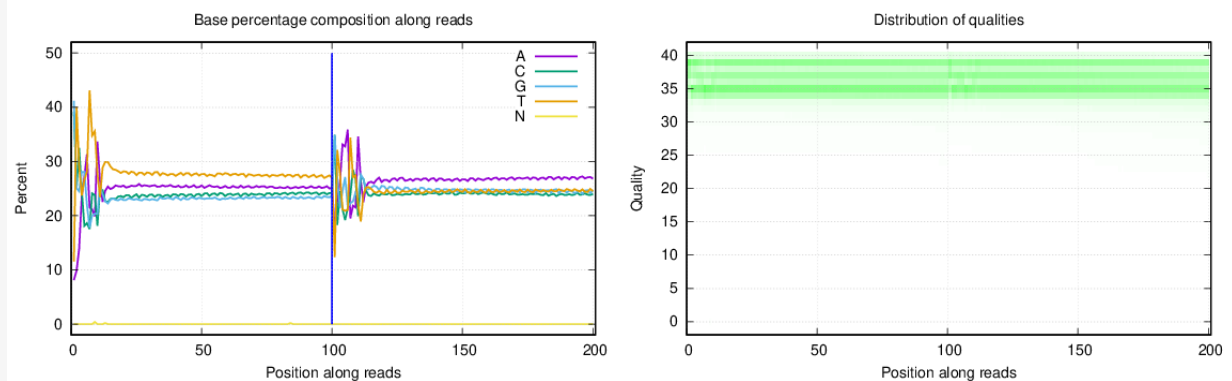


Figure 1-5 Distribution of base percentage and qualities of sample C3.

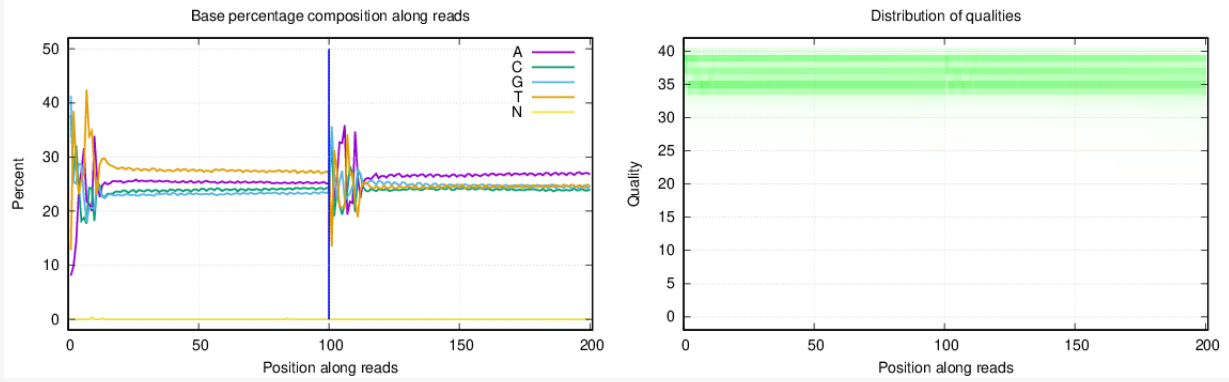


Figure 1-6 Distribution of base percentage and qualities of sample C2.

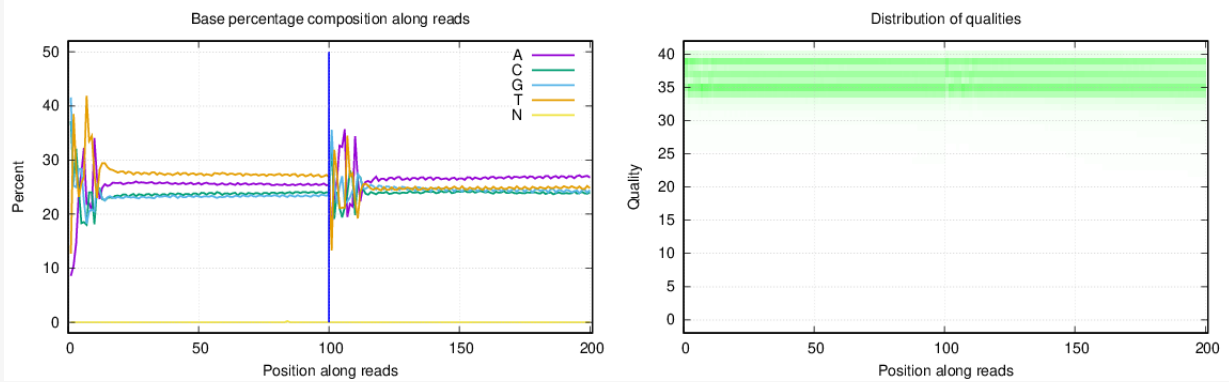


Figure 1-7 Distribution of base percentage and qualities of sample C1.

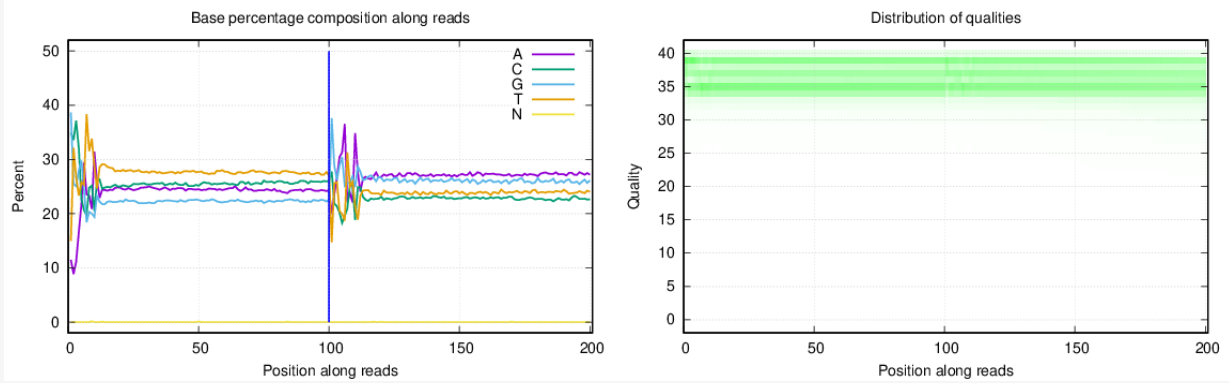


Figure 1-8 Distribution of base percentage and qualities of sample V6.

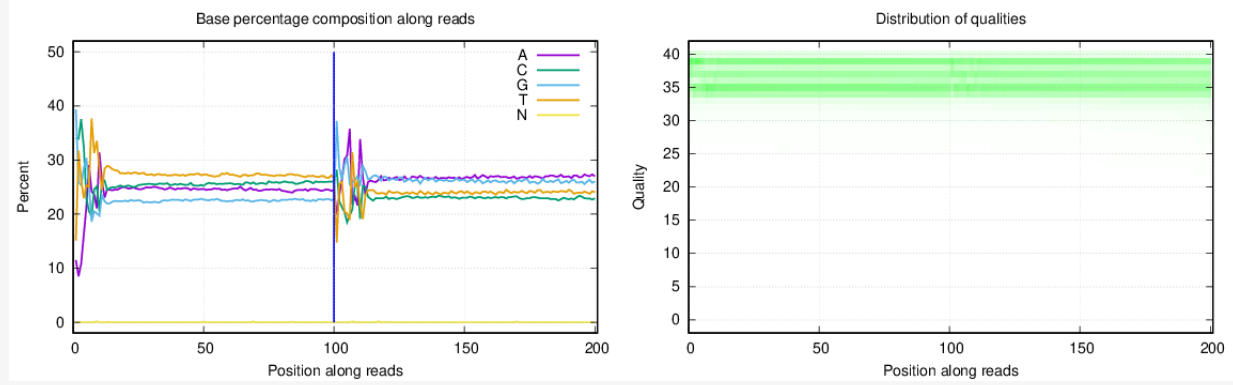


Figure 1-9 Distribution of base percentage and qualities of sample V5.

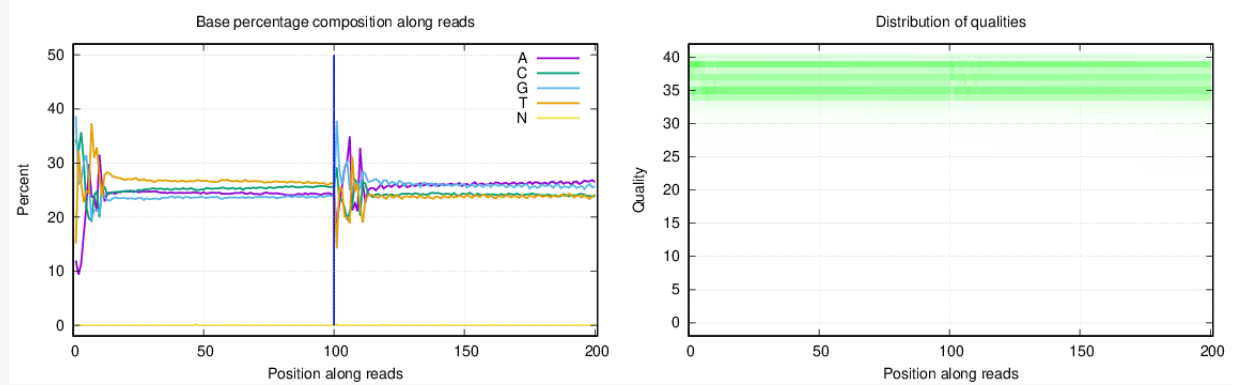


Figure 1-10 Distribution of base percentage and qualities of sample V4.

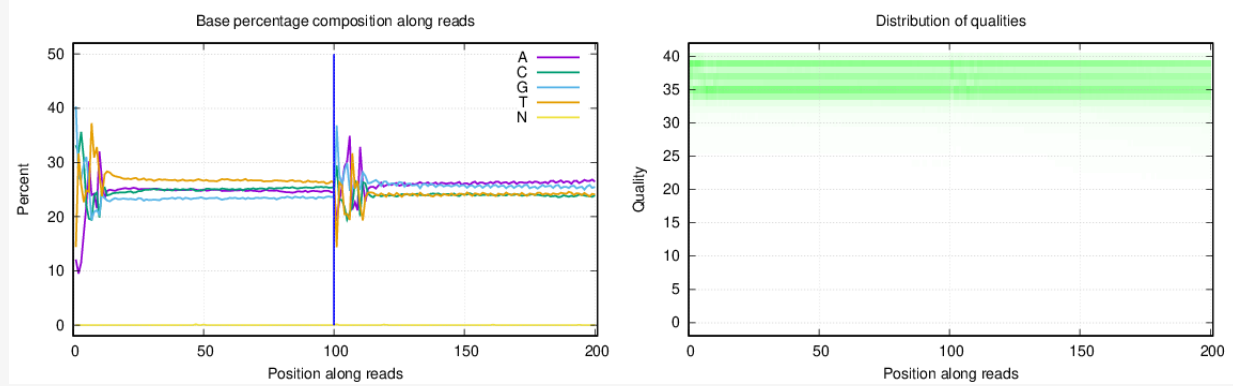


Figure 1-11 Distribution of base percentage and qualities of sample V3.

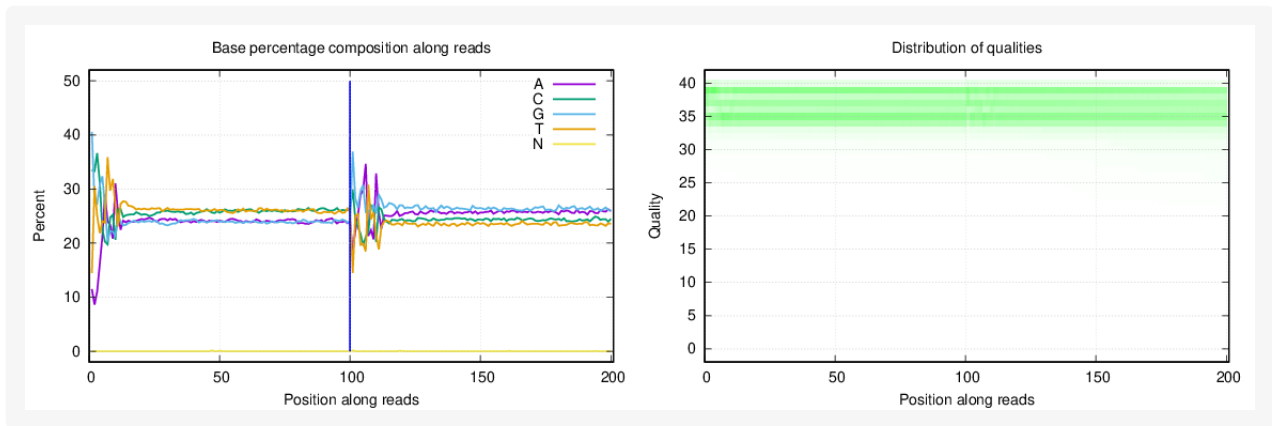


Figure 1-12 Distribution of base percentage and qualities of sample V2.

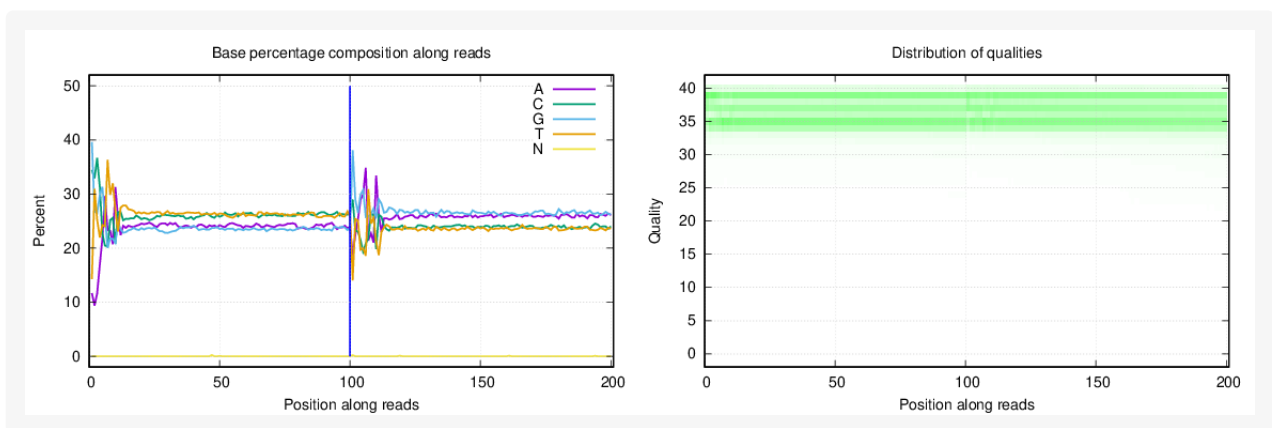


Figure 1-13 Distribution of base percentage and qualities of sample V1.

2. Methods

2.1 Experimental Procedure

The library construction method and sequencing process are carried out according to the following steps:

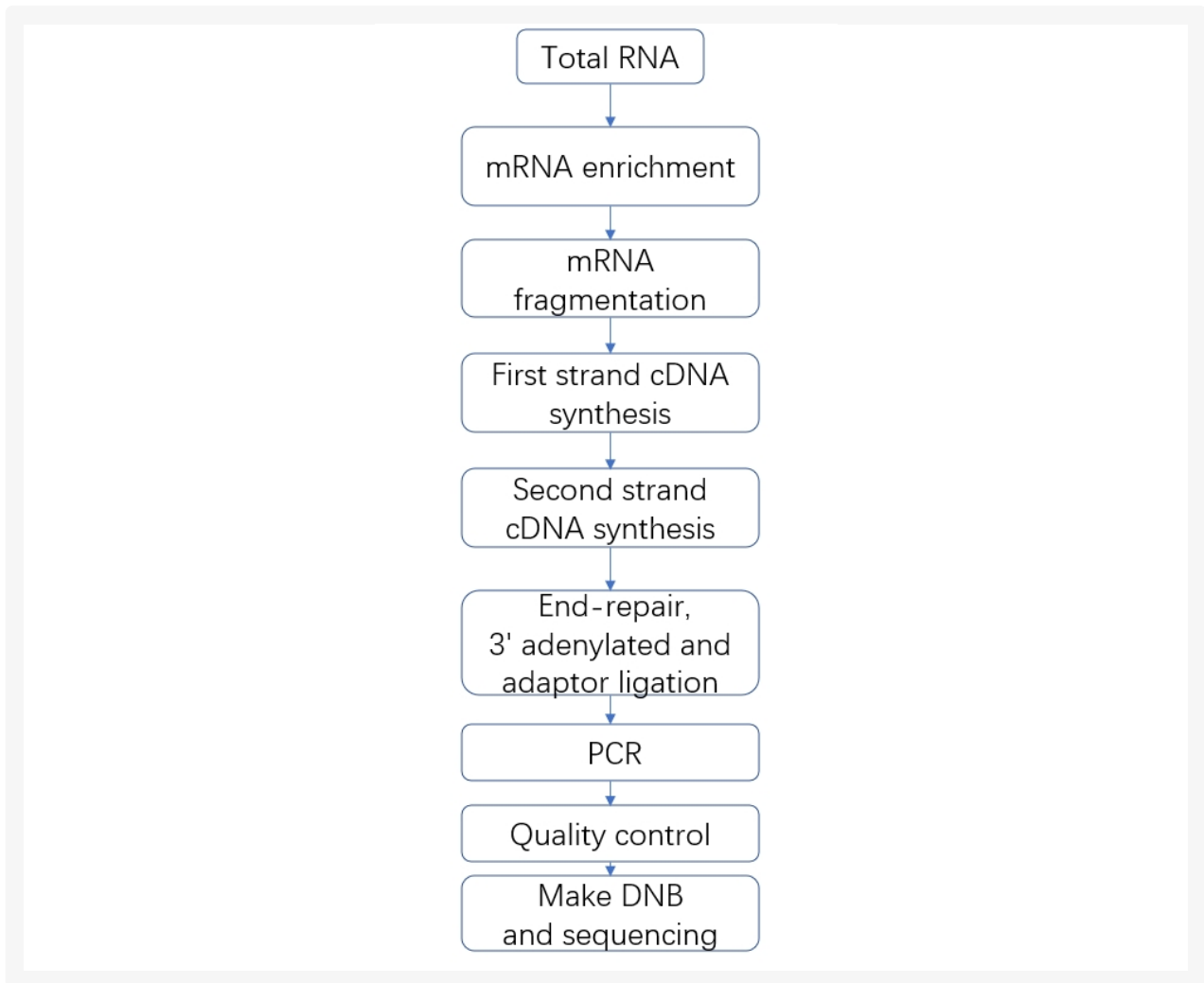


Figure 2 Workflow of experiment.

1. Take a certain amount of total RNA samples, and use oligo dT beads to enrich mRNA with poly A tail;
2. mRNA molecules were fragmented into small pieces;
3. The fragmented mRNA was synthesized into first strand cDNA using random primers;
4. The second strand cDNA was synthesized with dUTP instead of dTTP;
5. The synthesized cDNA was subjected to end-repair and 3' adenylated. Adaptors were ligated to the ends of these 3' adenylated cDNA fragments;
6. Perform PCR amplification;
7. Library quality control;
8. Library circularization;
9. The library was amplified to make DNA nanoball (DNB);
10. Sequencing on DNBSEQ (DNBSEQ Technology) platform.

2.2 Bioinformatic Analysis Workflow



Figure 3 Bioinformatic analysis workflow.

2.3 Parameters for Data Filtering

Raw data with adapter sequences or low-quality sequences was filtered. We first went through a series of data processing to remove contamination and obtain valid data. This step was completed by SOAPnuke [\[1\]](#) developed by BGI.

SOAPnuke software filter parameters: "-n 0.01 -l 20 -q 0.4 --adaMR 0.25 --polyX 50 --minReadLen 100", steps of filtering:

1. Filter adapter: if the sequencing read matches 25.0% or more of the adapter sequence (maximum 2 base mismatches are allowed), remove the entire read;
2. Filter read length: if the length of the sequencing read is less than 100 bp, discard the entire read;
3. Remove N: if the N content in the sequencing read accounts for 1.0% more of the entire read, discard the entire read;
4. Remove polyX: if the length of polyX (X can be A, T, G or C) in the sequencing read exceeds 50bp, discard the entire read;
5. Filter low-quality data: if the bases with a quality value of less than 20 in the sequencing read account for 40.0% or more of the entire read, discard the entire read;
6. Obtain Clean reads: the output read quality value system is set to Phred+33.

3. Help

3.1 FASTQ format description

Images generated by sequencers are converted by base calling into nucleotide sequences, which are called raw data or raw reads and are stored in FASTQ format. FASTQ files are text files that store both reads sequences and their corresponding quality scores. Each read is described in four lines as follows:

```
@V300029258L2C001R0010000017/1
TTTTCTGCTCCTTTTGATGCTATTAACAATTGCTTCAAGTTCAAGGGCACCTGCCTCAAAGTCCCTTTCTCCAGACAAAATCTGAGA
GTTTCATCTTT
+
=,DDE@EFFF=DFDEFCCFDEFEGFEEAFDFE=FFCFEEEEFDEEEFDF8FFEFFEFF:FFEDF=EFDGE<1FDCEFFFFDFEDCF
FFFFFFFF9GCF
```

The first line is the sequence identifier and related description information, starting with '@'; the second line is the base sequence information; the third line starts with '+', followed by the sequence identifier, description information, or nothing; The four lines are quality information, which corresponds to the sequence in the second line. Each base has a quality score. Depending on the scoring system, each character represents a different quality value.

The figure below shows the concise correspondence between the sequencing error rate and the sequencing quality value. Specifically, if the sequencing error rate is denoted by E and the base quality value is denoted by SQ , there are the following relationships:

$$SQ = -10 * (\log \frac{E}{1-E}) / \log 10$$

In which:

$$E = \frac{Y}{1+Y}$$

$$Y = \frac{SQ}{e^{-10 * \log 10}}$$

1) For the quality system data with a sequencing quality value of 33: the sequencing quality value of the base = the ASCII value corresponding to the quality information character -33, for example, the ASCII value corresponding to A is 65, then the corresponding base quality value is 65-33 = 32. The base quality value of the DNBSEQ sequencing platform ranges from 2 to 42.

2) For quality system data with a sequencing quality value of 64: the sequencing quality value of the base = the ASCII value corresponding to the quality information character -64, for example, the corresponding ASCII value of c is 99, then the corresponding base quality value is 99-64 = 35. The base quality value of the DNBSEQ sequencing platform ranges from 2 to 43.

Table 2 A summary table of the relationship between sequencing error rate and sequencing quality

Sequencing error rate	Sequencing quality value	Character(Phred64)	Character
5%	13	M	.
1%	20	T	5
0.1%	30	^	?

- Sequencing error rate: Sequencing error rate
- Sequencing quality value: Sequencing quality value
- Character(Phred64): ASCII code under Phred+64
- Character(Phred33): ASCII code under Phred+33

4. References

1. Chen, Y., Chen, Y., Shi, C., Huang, Z., Zhang, Y., Li, S., Li, Y., Ye, J., Yu, C., Li, Z., Zhang, X., Wang, J., Yang, H., Fang, L., & Chen, Q. (2018). SOAPnuke: a MapReduce acceleration-supported software for integrated quality control and preprocessing of high-throughput sequencing data. GigaScience, 7(1), 1-6.
<https://doi.org/10.1093/gigascience/gix120> ↵

The logo for BGI Tech, featuring the letters 'BGI' in a bold, sans-serif font, followed by a dot and the word 'Tech' in a similar font. The background of the slide is a blue world map with a grid pattern, and a solid orange rectangle is positioned to the right of the logo.

BGI·Tech

Omics For All

| Contact us

Website: www.bgi.com

E-mail: info@bgi.com

For Research Use Only. Not for use in diagnostic procedures.

© 2023 BGI Genomics Co.,Ltd. All right reserved. All trademarks are the property of BGI, or their respective owners.