

Technical Support Center: Improving Read Length in Sequencing of Repetitive DNA

Author: BenchChem Technical Support Team. **Date:** April 2026

Compound of Interest

Compound Name: 7-Iodo-2',3'-Dideoxy-7-Deaza-Guanosine
Cat. No.: B10831770

[Get Quote](#)

This technical support center provides troubleshooting guides and frequently asked questions (FAQs) to help researchers, scientists, and drug development professionals overcome challenges associated with sequencing repetitive DNA regions.

Troubleshooting Guide

This guide addresses specific issues that can arise during the sequencing of repetitive DNA and offers potential causes and solutions.

Issue 1: My long-read sequencing run produced short average read lengths (low N50) for a sample with known repetitive elements.

- Potential Cause 1: Poor Quality High Molecular Weight (HMW) DNA

The success of long-read sequencing is highly dependent on the integrity of the input DNA. [1] Short fragments present in the DNA extract can be preferentially sequenced, leading to a lower average read length.[2]

Solution:

- **Assess DNA Integrity:** Before library preparation, it is crucial to assess the size distribution of your extracted DNA. Use methods like pulsed-field gel electrophoresis (PFGE) or automated electrophoresis systems (e.g., Agilent Femto Pulse) to ensure a high proportion of large DNA fragments (>50 kb).^{[2][3]} For optimal results with PacBio sequencing, a Genomic Quality Number (GQN) of >9 for fragments above 10 kb is recommended.^[2]
- **Optimize DNA Extraction:** Employ a gentle DNA extraction method specifically designed for HMW DNA. Avoid vigorous vortexing or pipetting, which can shear the DNA.^[1] Using a mortar and pestle for tissue grinding has been shown to yield longer reads compared to a tissue homogenizer.^[4]
- **Potential Cause 2: Suboptimal Library Preparation**

The library preparation protocol can introduce fragmentation or biases that shorten the resulting read lengths.

Solution:

- **Minimize Fragmentation:** Handle the DNA gently throughout the library preparation process. Use wide-bore pipette tips to reduce mechanical stress.
- **Size Selection:** Incorporate a size selection step using beads or gel electrophoresis to remove smaller DNA fragments before sequencing. This ensures that the library loaded onto the sequencer is enriched for long DNA molecules.
- **Follow Recommended Protocols:** Adhere to the specific library preparation protocols recommended by the sequencing platform provider (e.g., PacBio, Oxford Nanopore Technologies), as these have been optimized for long-read sequencing.^{[5][6]}
- **Potential Cause 3: Presence of Sequencing Inhibitors**

Contaminants from the DNA extraction process (e.g., salts, EDTA, phenols, polysaccharides) can inhibit the polymerase activity during sequencing, leading to premature termination of reads.^{[3][7]}

Solution:

- Assess DNA Purity: Use spectrophotometry (e.g., NanoDrop) to check the purity of your DNA sample. Aim for a 260/280 ratio of ~1.8 and a 260/230 ratio of 2.0-2.2.[3]
- DNA Purification: If contaminants are suspected, perform an additional cleanup step using a purification kit or bead-based methods.

Issue 2: The sequencing reads do not fully span the repetitive region of interest.

- Potential Cause 1: The repetitive region is longer than the achievable read length of the sequencing run.

Even with long-read technologies, some highly complex and long repetitive regions may exceed the length of individual reads.[8]

Solution:

- Optimize for Ultra-Long Reads (Oxford Nanopore): For Oxford Nanopore sequencing, specific protocols and library preparation kits are available to generate ultra-long reads (> 100 kb).[9] This often involves very careful HMW DNA extraction and handling.
 - Utilize Paired-End Sequencing (Short-Read/Hybrid approach): While not a long-read solution, paired-end short-read sequencing can sometimes help to scaffold contigs across shorter repetitive regions by providing information about the distance between the two ends of a fragment.[10]
 - Consider a Hybrid Assembly Approach: Combining the high accuracy of short reads with the length of long reads can help to resolve complex repeats.[11]
- Potential Cause 2: The repetitive region contains features that cause the polymerase to dissociate.

Certain DNA sequences, such as long homopolymer stretches (e.g., poly-A tails) or regions prone to forming strong secondary structures, can cause the sequencing polymerase to stall or fall off the template.[12][13]

Solution:

- Optimize Sequencing Chemistry: Some sequencing platforms offer alternative chemistries or enzymes that are better at processing difficult templates.[\[14\]](#) For example, using a dGTP analog like 7-deaza-dGTP during PCR amplification prior to sequencing can help to disrupt G-C rich secondary structures.[\[15\]](#)[\[16\]](#)
- Bioinformatic Approaches: Specialized bioinformatics tools are designed to handle and assemble reads from repetitive regions.[\[11\]](#)[\[17\]](#)

Issue 3: High error rates in the sequenced repetitive regions are making assembly difficult.

- Potential Cause: Inherent error rates of the sequencing technology.

While constantly improving, long-read sequencing technologies have historically had higher error rates than short-read platforms.[\[18\]](#) These errors can complicate the assembly of highly similar repeat units.

Solution:

- Use High-Fidelity Long Reads (PacBio HiFi): PacBio's HiFi sequencing, which utilizes a circular consensus sequencing (CCS) approach, generates long reads with high accuracy (>99.9%).[\[18\]](#)[\[19\]](#) This high accuracy is crucial for distinguishing between individual repeat units.
- Increase Sequencing Depth: Higher coverage allows for more robust error correction during the assembly process. The consensus sequence derived from multiple reads is more likely to be accurate.[\[10\]](#)
- Employ Polishing Steps in Bioinformatics: After an initial assembly, use "polishing" algorithms that leverage the raw sequencing data (or even supplementary short-read data) to correct errors in the assembled sequence. Tools like Pilon (for Illumina data) and Arrow or Nanopolish (for long-read data) are commonly used.[\[11\]](#)

Frequently Asked Questions (FAQs)

Q1: Why is it so difficult to sequence repetitive DNA with short-read technologies?

Short-read sequencing technologies, such as those from Illumina, generate reads that are typically 50-300 base pairs long.[9][20] When a repetitive region is longer than the read length, it becomes impossible to uniquely map the reads and determine their correct order and copy number.[21][22] This leads to fragmented assemblies with gaps in the repetitive regions.[20][23]

Q2: What are the main long-read sequencing technologies available for sequencing repetitive DNA?

The two predominant long-read sequencing technologies are:

- Single-Molecule Real-Time (SMRT) Sequencing from Pacific Biosciences (PacBio): This technology, particularly with its High-Fidelity (HiFi) reads, produces reads that are thousands of base pairs long with very high accuracy (>99.9%).[19][21]
- Nanopore Sequencing from Oxford Nanopore Technologies (ONT): This technology passes a single DNA molecule through a protein nanopore and measures the resulting changes in electrical current to determine the sequence.[19][24] It is capable of producing very long to ultra-long reads.[9]

Q3: What is a "hybrid assembly" and when should I use it?

A hybrid assembly combines the strengths of both short-read and long-read sequencing data.[11] Long reads are used to create a scaffold of the genome, correctly ordering and orienting the contigs, especially across repetitive regions.[25] The highly accurate short reads are then used to "polish" this scaffold, correcting errors and improving the overall sequence accuracy. This approach is particularly useful for large and complex genomes with a high proportion of repetitive DNA.[11]

Q4: Are there methods to sequence a specific repetitive region without having to sequence the entire genome?

Yes, targeted sequencing methods can be used to enrich for specific repetitive regions of interest. These approaches include:

- Amplicon Sequencing: This method uses PCR to amplify the target region. Long-range PCR can be used to generate amplicons that are several kilobases long, which can then be

sequenced using long-read technologies.[26]

- Hybridization Capture: This technique uses probes to capture and enrich for specific DNA sequences from a complex library.[27] The enriched library is then sequenced.
- CRISPR-Cas9 Based Enrichment: This method uses the CRISPR-Cas9 system to target and cut out specific genomic regions, which can then be isolated and sequenced. PacBio offers the PureTarget repeat expansion panel which uses this technology.[26][27]

Q5: What are some key bioinformatics tools for assembling genomes with repetitive elements from long reads?

Several assemblers are specifically designed to handle long-read data and resolve repetitive regions. Some popular tools include:

- Canu: An assembler designed for high-error long-read data that uses an Overlap-Layout-Consensus (OLC) algorithm.[11]
- Flye: An assembler optimized for speed and noisy long-read data, which includes algorithms for resolving repeats.[11]
- MaSuRCA: A hybrid assembler that can combine Illumina short reads with PacBio or Nanopore long reads.[11]
- TandemTools: A suite of tools specifically designed for mapping long reads and assessing assembly quality in extra-long tandem repeats.[17]

Quantitative Data Summary

Table 1: Comparison of Sequencing Technologies for Repetitive DNA

Feature	Short-Read Sequencing (e.g., Illumina)	Long-Read Sequencing (PacBio HiFi)	Long-Read Sequencing (Oxford Nanopore)
Read Length	50 - 300 bp[9][20]	10 - 25 kb[9]	1 kb - 4+ Mb[9][28]
Accuracy	>99.9%[21]	>99.9%[19]	~95% (raw reads, can be improved with consensus)[18]
Throughput	Very High	High	High, scalable
Ability to Resolve Repeats	Poor, if repeat is longer than read length[21][22]	Excellent[29]	Excellent[19]
Primary Challenge	Inability to span long repeats[20]	Higher initial DNA input requirement	Higher raw read error rate[18]

Table 2: General Recommendations for HMW DNA Quality Control

QC Metric	Recommended Value	Purpose
DNA Integrity	Majority of fragments > 50 kb	Ensures long reads can be generated.[1]
260/280 Ratio	~1.8	Indicates purity from protein contamination.[3]
260/230 Ratio	2.0 - 2.2	Indicates purity from salt and organic solvent contamination. [3]

Experimental Protocols

Protocol 1: High-Level Workflow for HMW DNA Extraction and QC

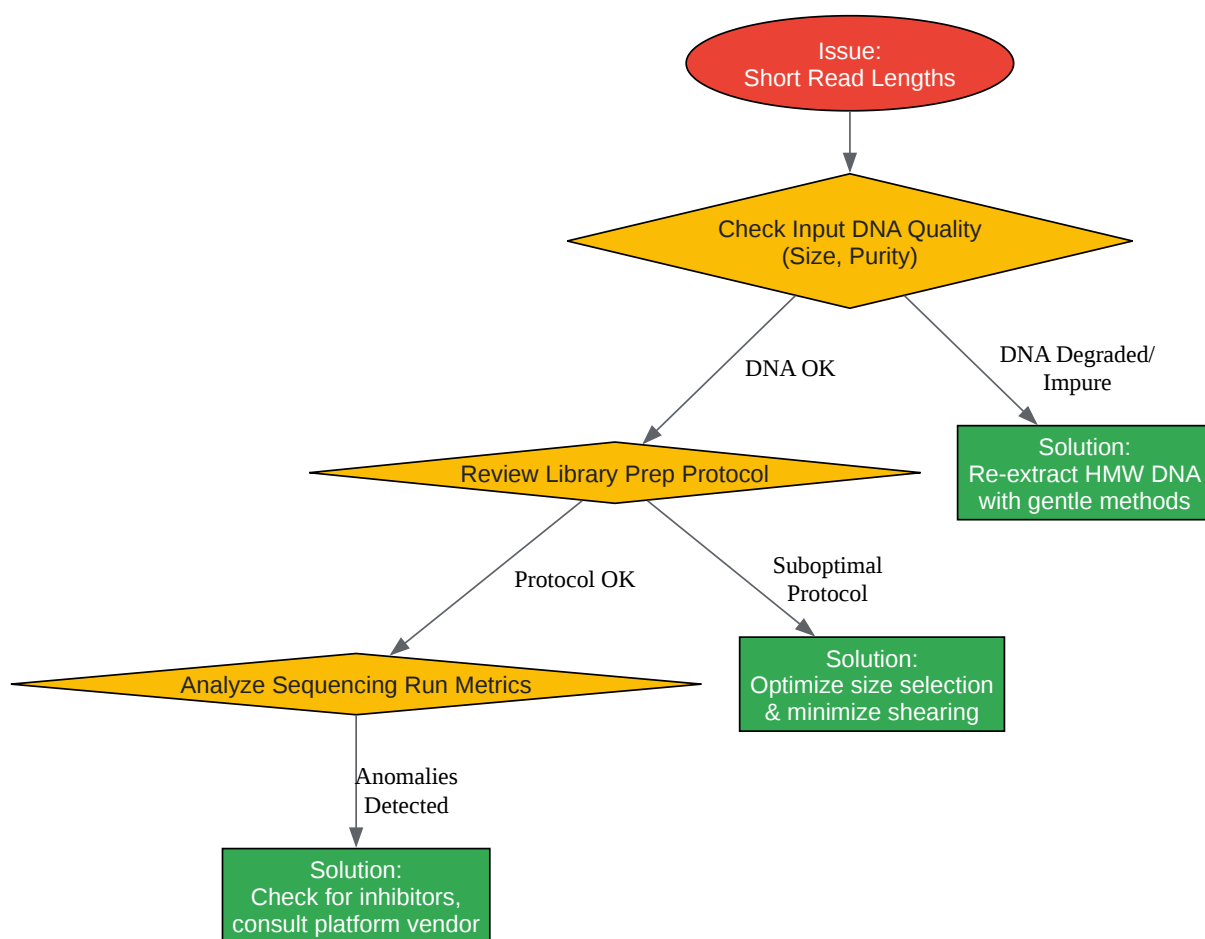
- Tissue Disruption: Gently disrupt the tissue using a method that minimizes DNA shearing, such as grinding in liquid nitrogen with a mortar and pestle.[4]

- Cell Lysis: Use a lysis buffer with gentle mixing to break open the cells and release the DNA.
- DNA Purification: Purify the DNA from other cellular components using a method suitable for HMW DNA, such as phenol-chloroform extraction followed by ethanol precipitation or a specialized kit.
- DNA Quality Control:
 - Assess DNA concentration using a fluorometric method (e.g., Qubit).[3]
 - Evaluate DNA purity using a spectrophotometer (e.g., NanoDrop) to check 260/280 and 260/230 ratios.[3]
 - Determine the DNA fragment size distribution using pulsed-field gel electrophoresis or an automated system like the Agilent Femto Pulse.[2][3]

Protocol 2: General Steps for Long-Read Library Preparation

- DNA Fragmentation (Optional/Tunable): Depending on the desired read length, the HMW DNA may be fragmented to a specific size range using enzymatic or mechanical methods.[30][31]
- End Repair and A-tailing: The ends of the DNA fragments are repaired and a single adenine base is added to prepare them for adapter ligation.[30][31]
- Adapter Ligation: Sequencing adapters, which are necessary for the sequencing instrument to recognize the DNA, are ligated to the ends of the fragments.[30][31]
- Cleanup and Size Selection: The library is cleaned to remove any unincorporated adapters and enzymes. A size selection step may be performed to enrich for fragments of a desired length.[31]
- Library Quantification and QC: The final library is quantified and its size distribution is checked before loading it onto the sequencer.[31]

Visualizations



[Click to download full resolution via product page](#)

Need Custom Synthesis?

BenchChem offers custom synthesis for rare earth carbides and specific isotopic labeling.

Email: info@benchchem.com or [Request Quote Online](#).

References

- [1. neb.com \[neb.com\]](#)
- [2. Evaluation of controls, quality control assays, and protocol optimisations for PacBio HiFi sequencing on diverse and challenging samples - PMC \[pmc.ncbi.nlm.nih.gov\]](#)
- [3. nanoporetech.com \[nanoporetech.com\]](#)
- [4. Balancing read length and sequencing depth: Optimizing Nanopore long-read sequencing for monocots with an emphasis on the Liliales - PMC \[pmc.ncbi.nlm.nih.gov\]](#)
- [5. twistbioscience.com \[twistbioscience.com\]](#)
- [6. Sample and Library Preparation for PacBio Long-Read Sequencing in Grapevine - PubMed \[pubmed.ncbi.nlm.nih.gov\]](#)
- [7. MGH DNA Core \[dnacore.mgh.harvard.edu\]](#)
- [8. bioinformatics.stackexchange.com \[bioinformatics.stackexchange.com\]](#)
- [9. nanoporetech.com \[nanoporetech.com\]](#)
- [10. irepertoire.com \[irepertoire.com\]](#)
- [11. Genome Assembly: Overview of the Tools - CD Genomics \[cd-genomics.com\]](#)
- [12. MU Genomics Technology Core \[mugenomicscore.missouri.edu\]](#)
- [13. genomique.irc.ca \[genomique.irc.ca\]](#)
- [14. base4.co.uk \[base4.co.uk\]](#)
- [15. Short Sequencing Trace Read Lengths \[nucleics.com\]](#)
- [16. bitesizebio.com \[bitesizebio.com\]](#)
- [17. academic.oup.com \[academic.oup.com\]](#)
- [18. Probably Correct: Rescuing Repeats with Short and Long Reads - PMC \[pmc.ncbi.nlm.nih.gov\]](#)

- [19. Nexco | The Power and Promise of Long-Read Sequencing Technologies \[nexco.ch\]](#)
- [20. Short-read sequencing — Knowledge Hub \[genomicseducation.hee.nhs.uk\]](#)
- [21. frontlinegenomics.com \[frontlinegenomics.com\]](#)
- [22. researchgate.net \[researchgate.net\]](#)
- [23. Long-Read DNA Sequencing \[genome.gov\]](#)
- [24. The Complete Overview of Long-Read Sequencing in 2024 - CD Genomics \[cd-genomics.com\]](#)
- [25. Read \(biology\) - Wikipedia \[en.wikipedia.org\]](#)
- [26. pacb.com \[pacb.com\]](#)
- [27. pacb.com \[pacb.com\]](#)
- [28. A Hitchhiker's Guide to long-read genomic analysis - PMC \[pmc.ncbi.nlm.nih.gov\]](#)
- [29. pacb.com \[pacb.com\]](#)
- [30. idtdna.com \[idtdna.com\]](#)
- [31. NGS Library Preparation: Proven Ways to Optimize Sequencing Workflow - CD Genomics \[cd-genomics.com\]](#)
- [To cite this document: BenchChem. \[Technical Support Center: Improving Read Length in Sequencing of Repetitive DNA\]. BenchChem, \[2026\]. \[Online PDF\]. Available at: \[https://www.benchchem.com/product/b10831770/docs#technical-support-center-improving-read-length-in-sequencing-of-repetitive-dna\]](#)

Disclaimer & Data Validity:

The information provided in this document is for Research Use Only (RUO) and is strictly not intended for diagnostic or therapeutic procedures. While BenchChem strives to provide accurate protocols, we make no warranties, express or implied, regarding the fitness of this product for every specific experimental setup.

Technical Support: The protocols provided are for reference purposes. Unsure if this reagent suits your experiment?

Need Industrial/Bulk Grade? [Request Custom Synthesis Quote](#)

BenchChem

Our mission is to be the trusted global source of essential and advanced chemicals, empowering scientists and researchers to drive progress in science and industry.

Contact

Address: 3281 E Guasti Rd

Ontario, CA 91761, United States

Phone: (601) 213-4426

Email: info@benchchem.com

[Contact our Ph.D. Support Team for a compatibility check](#)