

RESEARCH

Open Access



A systematic benchmark of high-accuracy PacBio long-read RNA sequencing for transcript-level quantification

David Wissel^{1,2,3†}, Madison M. Mehlferber^{4,5†}, Khue M. Nguyen⁶, Vasili Pavlenko⁵, Elizabeth Tseng⁷, Mark D. Robinson^{1,3*} and Gloria M. Sheynkman^{4,5,8*}

[†]David Wissel and Madison M. Mehlferber contributed equally to this work.

*Correspondence:
mark.robinson@mls.uzh.ch;
gs9yr@virginia.edu

¹ Department of Molecular Life Sciences, University of Zurich, Zurich, Switzerland

² Department of Computer Science, ETH Zurich, Zurich, Switzerland

³ SIB Swiss Institute of Bioinformatics, University of Zurich, Zurich, Switzerland

⁴ Department of Molecular Physiology and Biological Physics, University of Virginia, Charlottesville, VA, USA

⁵ Department of Biochemistry and Molecular Genetics, University of Virginia, Charlottesville, VA, USA

⁶ Center for Digital Health, Berlin Institute of Health (BIH) at Charité - Universitätsmedizin, Berlin, Germany

⁷ Pacific Biosciences, Menlo Park, CA, USA

⁸ UVA Cancer Center, Charlottesville, VA, USA

Abstract

Background: The assembly of fragmented RNA-sequencing reads into complete transcripts is error-prone, particularly for genes with complex splicing, resulting in ambiguity in transcript discovery and quantification. PacBio long-read RNA sequencing resolves transcripts with greater clarity than short-read technologies. PacBio Kinnex employs a cDNA concatenation approach that increases read yield on average by 8-fold relative to previous protocols. However, its quantitative performance remains under-evaluated at scale.

Results: Here, we benchmark the high-throughput PacBio Kinnex platform against Illumina short-read RNA-seq using matched, deeply sequenced datasets across a time course of endothelial cell differentiation. Compared to Illumina, Kinnex achieves comparable gene-level quantification and more accurate transcript discovery and transcript quantification. While Illumina detects more transcripts overall, many reflect potentially unstable or ambiguous estimates in complex genes. Kinnex largely avoids these issues, producing more reliable differential transcript expression calls, despite a mild bias against short transcripts (shorter than 1.25 kb). When correcting Illumina for inferential variability, Kinnex and Illumina quantifications are highly concordant, demonstrating equivalent performance. We also benchmark long-read tools, nominating Oarfish as the most efficient for our Kinnex data.

Conclusions: Together, our results establish Kinnex as a reliable platform for full-length transcript quantification.

Keywords: PacBio, Long-read RNA-seq, Quantification, Endothelial cells

Background

RNA-sequencing (RNA-seq) has enabled comprehensive measurements of the transcriptome, but standard short-read approaches face inherent limitations in resolving full-length transcripts. The assembly of fragmented reads into complete transcripts is error-prone, particularly for genes with complex splicing, resulting in ambiguity in



© The Author(s) 2026. **Open Access** This article is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if changes were made. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by/4.0/>.

transcript discovery and quantification [1]. PacBio's long-read RNA-seq (lrRNA-seq) platform addresses this limitation by sequencing entire RNA molecules. However, until recently, the low throughput (less than 10 M reads per sample) limited its utility for quantitative analyses, a finding confirmed by benchmarking efforts such as the Long-Read RNA-seq Genome annotation Assessment Consortium (LRGASP) [2].

The MAS-Iso-seq method, commercialized as PacBio Kinnex, employs a cDNA concatenation approach that increases read yield on average by 8-fold relative to previous protocols [3]. This throughput improvement could enable new experimental designs that combine the transcript resolution of long-reads with the quantitative depth required for transcriptome profiling. Nonetheless, a rigorous evaluation of Kinnex for transcript quantification has not yet been performed. Previous studies have focused primarily on earlier PacBio chemistries [4, 5] or have emphasized the potential for transcript discovery rather than quantification [6, 7]. Most existing benchmarking studies for long-read quantification more broadly have focused on Oxford Nanopore (ONT) data, which remains more error-prone than PacBio [5, 6]. As a result, it remains unclear how Kinnex compares to Illumina across complex biological samples, which tools best quantify its transcript abundances, and how it performs in differential expression analyses.

In this study, we evaluate the performance of PacBio Kinnex lrRNA-seq for transcript-level quantification, benchmarked against standard Illumina short-read RNA-seq. Using matched, deeply sequenced datasets collected across a five-day time course of human iPSC differentiation to endothelial cells, we systematically compare transcript discovery, abundance quantification, and downstream tasks such as differential gene and transcript expression. We also assess the performance of multiple long-read quantification tools and analyze sources of quantification discordance between platforms. Together, our results help to define the capabilities and limitations of Kinnex for transcript-resolved RNA-seq at scale.

Results

Kinnex lrRNA-seq is technically robust and enables effective biological downstream analysis

We collected RNA from human induced pluripotent stem cells (iPSCs) (WTC11) cells undergoing a five-day endothelial differentiation time course [8], and performed matched sequencing using both Kinnex (PacBio Revio) and Illumina (NovaSeq) platforms (see [Methods](#)). To assess technical performance, we spiked in SIRV controls (Fig. 1A). For each day of differentiation, we isolated RNA with three replicates for day zero and five, two replicates for day three, and one replicate for the remaining days (one, two, and four) (Fig. 1A). We excluded the single day-two replicate from further analysis due to quality issues (Additional file 1: Fig. S1A-J). Illumina libraries yielded approximately 67 M short reads per sample on average (Fig. 1B). In parallel, Kinnex produced approximately 54 M full-length reads per sample, on average, with a median length of ~1.55-kb—roughly 5.2 times longer than the 299 bp paired-end Illumina reads, and thus a higher total number of bases sequenced per sample (Fig. 1C, D). Both technologies achieved high alignment rates to the human genome (Illumina: ~95%, Kinnex: ~97%; Fig. 1B), along with high average base quality scores (Illumina: Phred~36, Kinnex: Phred~38; Fig. 1F). In terms of empirical error profiles, Kinnex showed more indel-related errors compared to Illumina

(Kinnex average edit distance (indels): ~ 0.00205 , Illumina average edit distance (indels): ~ 0.0000004), with SNVs showing the reverse behavior (Kinnex average edit distance (SNVs): ~ 0.00087 , Illumina average edit distance (SNVs): ~ 0.00301) (Fig. 1G). Both technologies showed a slight 3' bias for shorter transcripts (shorter than 1 kb), with this bias improving for transcripts of medium length with increasing length (between one and two kb, two and three kb, and three and six kb) (Fig. 1H). For long transcripts (longer than six kb), Illumina exhibited a moderate 5' bias, whereas Kinnex showed a U-shaped coverage, with distinctly lower coverage around the middle of the gene body. We note that the behavior of long-reads not fully covering very long transcripts is consistent with previous works and may be attributable to a combination of internal priming and the difficulty of RT for long transcripts [9].

Next, we leveraged inferential variability, a measure of the degree of uncertainty when estimating the abundance of a specific transcript. This uncertainty may become a major issue when short reads map ambiguously to multiple transcripts (see [Methods](#)). Using Salmon for the quantification of Illumina samples and Oarfish for the quantification of Kinnex samples, and adjusting for the inferential variability in Illumina samples using existing approaches [10], we observed decreasing correlations between transcript-level sample expression over the course of the differentiation time-course as well as high reproducibility of gene and transcript quantification for replicates of the same day (Spearman > 0.9) (Fig. 1I, J). Multidimensional scaling (MDS) and Spearman correlation analyses showed tight clustering among same-day replicates and a clear separation along the differentiation axis for both platforms (Fig. 1J, K). At full depth, on Day 0, Kinnex detected a total of 25,267 transcripts (where detection was defined by an expression of 1

(See figure on next page.)

Fig. 1 Dataset quality control, summary statistics and characterization. **A** Experimental design to collect a matched RNA-sequencing dataset of differentiated primordial endothelial cells to track transcript dynamics over time. The Kinnex protocol emphasizes the use of full-length cDNA, distinguishing it from short-read RNA sequencing sample preparation workflows. Kinnex preparation abbreviations (bc1 = barcode 1, bc2 = barcode 2, A, B, C and D denote Kinnex-specific primers, c1, c2 and c3 represent individual cDNA molecules. Respective transcript quantification and discovery methods employed for each technology are denoted. **B** Number of raw and aligned reads per technology, stratified by their source across all samples. **C** Read lengths of 1M randomly sampled reads aligning to the human genome, stratified by technology for Days 0 and 5. We highlight Days 0 and 5 due to their replication. For Illumina, the length was calculated across both ends. **D** Number of base pairs sequenced from 1M randomly sampled reads collected for each technology (same reads as **C**). **E** Number of covered junctions per read from 1M randomly sampled reads collected for each technology (same reads as **C**). **F** Average base quality of 1M randomly sampled reads aligning to the human genome, stratified by technology for Days 0 and 5 (same reads as **C**). The average base quality for Illumina was determined from both ends. **G** Average edit distance per read, stratified by indel vs SNV errors, of 1M randomly sampled reads collected for each technology (same reads as **C**). **H** Normalized coverage of reads aligning to 2,500 randomly sampled GENCODE transcripts per length group each across gene body percentiles, stratified by technology and transcript length, across all samples. **I** Spearman correlation heatmap of transcript-level quantifications highlighting sample similarities across the full differentiation for each technology. **J** Scatter plot for Day 0–1 and Day 0–2 replicates at the gene and transcript levels following quantification, stratified by technology. **K** Transcript-level quantification-based MDS plots highlighting sample differences across the full differentiation, stratified by technology. **L** Left: Upset plot of the transcripts detected on Day 0 (> 1 CPM in all three replicates) for both technologies. Right: Attributes of transcripts detected by both or only one of the technologies. GC content, Inferential variability as measured by the inferential relative variance on Day 0 Illumina samples, transcript length, and median expression in both of the technologies. **M** Precision and recall of de-novo transcript discovery on three Day 0 (E1 mix) SIRV replicates by technology. Both platforms were ran on downsampled read sets consisting of 2.5 M reads. **N** Mean inferential variability of SIRV transcripts on three Day 0 replicates by whether they were discovered by both platforms, only Illumina, only Kinnex, or neither platform. Figure 1A was created in Biorender, M. (2025) [Biorender](#)

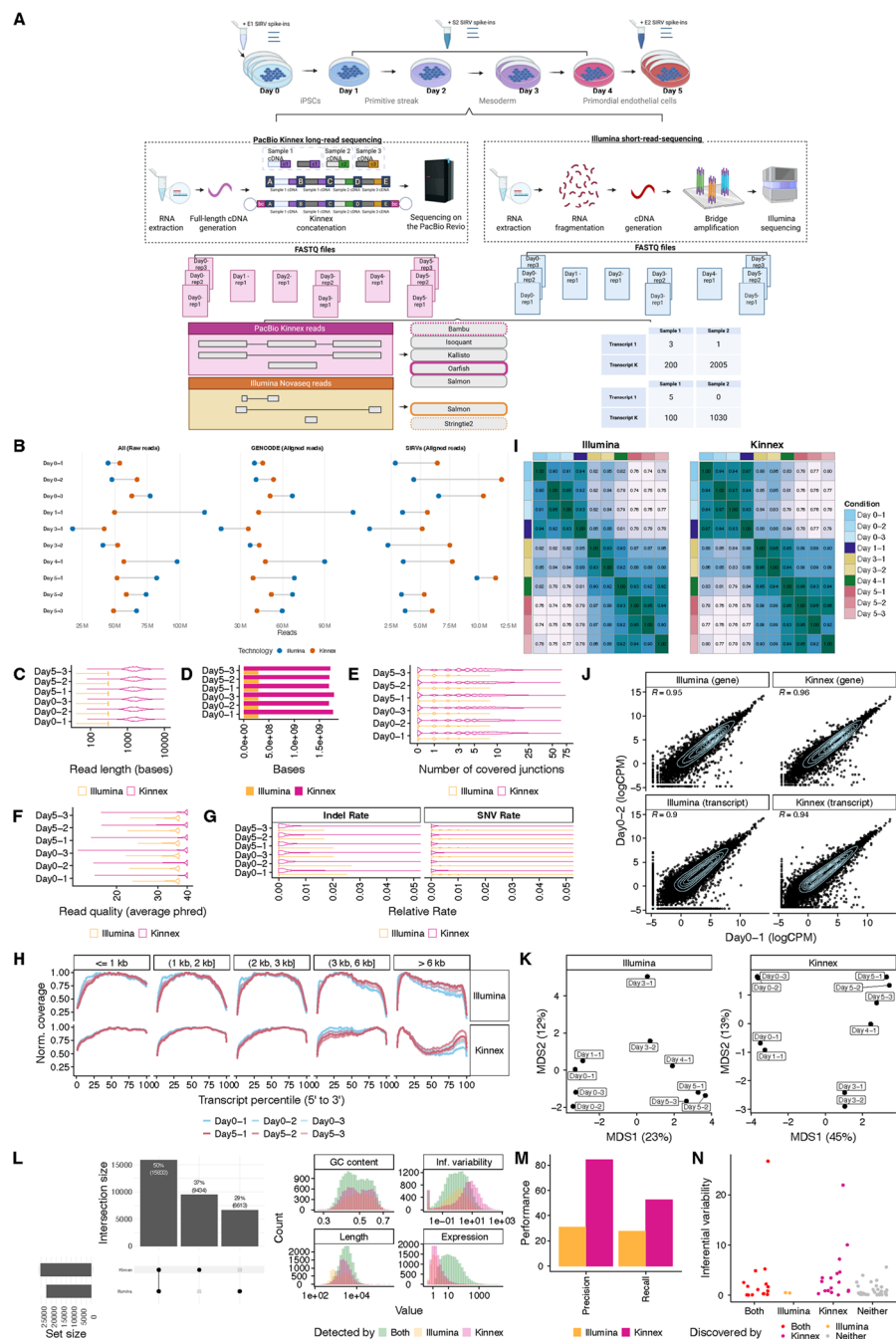


Fig. 1 (See legend on previous page.)

CPM or greater in all three replicates), 2,821 more compared to Illumina (Fig. 1L, left). The pattern of Kinnex detecting more transcripts than Illumina also held when defining detection using a simple read count threshold at lower, downsampled depths (Additional file 2: Table S1). The two main patterns in terms of which technology tended to detect which transcripts were the following: First, Illumina detected more transcripts that were relatively short (1 kb - 1.25 kb) compared to Kinnex (Fig. 1L, right). However, Illumina struggled with the detection of transcripts that exhibited high inferential variability

(Fig. 1L, right). In terms of expression and GC content, there were no clear patterns as far as differences in detection between the technologies went.

Differential Transcript Usage (DTU) analysis revealed prominent transcript switches across the differentiation, including within transcripts of the *IPO11* gene (Additional file 1: Fig. S2A, Additional file 3: Table S2). Marker gene expression matched biological expectations associated with the differentiation, including *NANOG* downregulation and endothelial cell (EC) marker induction (*KDR*, *ETV2*, *CD34*, *SOX7*, and *SOX17*) (Additional file 1: Fig. S2B and Additional file 4: Table S3) [11–15].

Collectively, these results demonstrate that Kinnex reads are comparable in quality to those from Illumina, are able to capture expected biological patterns, and are suitable for a technical comparison.

Kinnex enables the discovery of novel transcripts

We next assessed the potential for Kinnex and Illumina to discover transcripts (also known as transcript identification). We used the SIRV spike-ins, which provide a known set of sequences introduced during library preparation.

Using Bambu for the discovery of novel transcripts in Kinnex samples and StringTie2 for Illumina samples, both with default parameters, we found that Kinnex outperformed Illumina in both precision and recall, indicating more complete and accurate transcript discovery when using long-read data (Fig. 1M). To understand why Illumina-based transcript discovery was incomplete, we examined transcripts that were identified by Kinnex but missed by Illumina. We found that such transcripts tended to have more complex splicing patterns, as evidenced by higher inferential variability—such as a large number of alternative exons—which are challenging to resolve with short reads, although the pattern was not fully clear (Fig. 1N).

Given Kinnex's greater reliability for transcript discovery, we used it as the basis for identifying novel transcripts in the differentiation samples. We ran Bambu on the Kinnex-aligned reads and classified the transcripts using SQANTI3, leading to the identification of 12,341 novel transcripts and 515 novel genes (Additional file 1: Fig. S2C; Additional file 5: Table S4) [2, 16]. Of the 12,341 novel transcripts identified using Bambu and classified with SQANTI3, the majority (~89%) were lowly expressed (less than 10 CPM). However, 1,337 transcripts (10.8%) were at least moderately expressed (10 CPM or greater) and multi-exonic (SQANTI3 categories: ISM, NIC, or NNC; Additional file 1: Fig. S2C) [2, 17]. Among these, 1,097 (8.9% of all novel transcripts) were predicted to be protein-coding based on ORFanage. Overall, 77.5% (9,565/12,341) of all novel transcripts were predicted to be protein-coding. As an illustrative example, BambuTx5158 is a novel transcript with moderate expression (34 CPM), overlapping the *RP11* gene locus (ENSG00000254202), that has strong read support from both Illumina and Kinnex (Additional file 1: Fig. S2D), but is predicted to be non-coding.

Benchmarking of Kinnex quantification tools shows strong performance of Oarfish

Although Kinnex long-read sequencing excelled at discovering novel transcripts, its ability to measure the abundances of transcripts remained uncertain. Therefore, we next benchmarked the performance of Kinnex quantification tasks across several facets of transcriptome quantification, in comparison to Illumina.

To perform an in-depth characterization of transcript quantification, we first attempted to determine suitable quantification tools for Kinnex data, as long-read analysis methods are still in the process of refinement, while methods for the quantification of short-read RNA-seq (srRNA-seq) are comparatively well established [18]. We conducted a systematic comparison of five leading tools—Oarfish, lr-kallisto, Isoquant, Bambu, and Salmon—using read-depth-matched Kinnex datasets and SIRV spike-ins as ground truth [16, 19–22].

First, across the five tools, we compared the computational resource usage requirements needed by running different numbers of subsampled reads on the same hardware. In these efficiency evaluations, Oarfish and lr-kallisto stood out for their fast runtime and low memory usage across read depths (Fig. 2A).

Next, we proceeded to evaluate transcript quantification performance across three tasks: Differential Transcript Expression (DTE), fold-change estimation, and absolute quantification. First, we asked how well each tool detects Differential Transcript Expression (DTE) between the SIRV E1 and E2 mixes. Isoquant and lr-kallisto performed best, with accuracy comparable to Illumina (Fig. 2B). Oarfish and Bambu showed lower DTE accuracy, but this was largely due to an outlier transcript (SIRV303), which was challenging to resolve because of its sequence similarity to other SIRVs (Fig. 2C, top). The SIRV303 transcript is fully contained within the SIRV301 transcript. While the expression of SIRV303 stayed constant between the SIRV mixes used for days 0 and 5, the expression of SIRV301 increased 16-fold. Consequently, many of our tested long-read methods attributed most or even all reads of the SIRV303 transcript to the SIRV301 transcript, given the lack of unique alignments to the SIRV303 transcript (Fig. 2C, bottom).

Second, we examined the accuracy of the resulting fold-change estimates. Oarfish and Bambu were again affected by the SIRV303 outlier, and Salmon showed consistently poor performance across all metrics (Fig. 2D). The results of the remaining methods were concordant with the DTE results.

Third, we compared absolute abundance estimates (log-CPM) to ground truth SIRV concentrations. All tools showed moderate agreement (Pearson $r \sim 0.70$ overall), while Illumina displayed higher accuracy ($r \sim 0.88$ overall) (Fig. 2E). We also observed a moderate bias across all long-read methods against shorter transcripts (less than 1.25 kb), likely reflecting a limitation in Kinnex library preparation, which we believe to be related to the ratio of SPRI beads used for cleanup during the cDNA synthesis, having picked bead concentrations that would skew towards longer transcripts (Fig. 2D, E).

While SIRV spike-ins provide a ground truth, they do not accurately reflect the full complexity of the human transcriptome, so we next assessed tool performance for endogenous transcript quantifications against the full GENCODE annotation. Since a ground truth of endogenous transcripts is not available, as knowledge of the transcriptome is rapidly evolving, we ranked the five tools on two proxy metrics: relative quantification efficiency (fraction of counts over inputted raw reads) and reproducibility of Differential Transcript Expression (DTE) calls across tools (both Kinnex and Illumina Salmon analyses). The relative quantification efficiency measures how efficiently a method turns raw reads into counts. Assuming equally accurate quantification results between methods, this directly correlates with total counts and thus power. The number

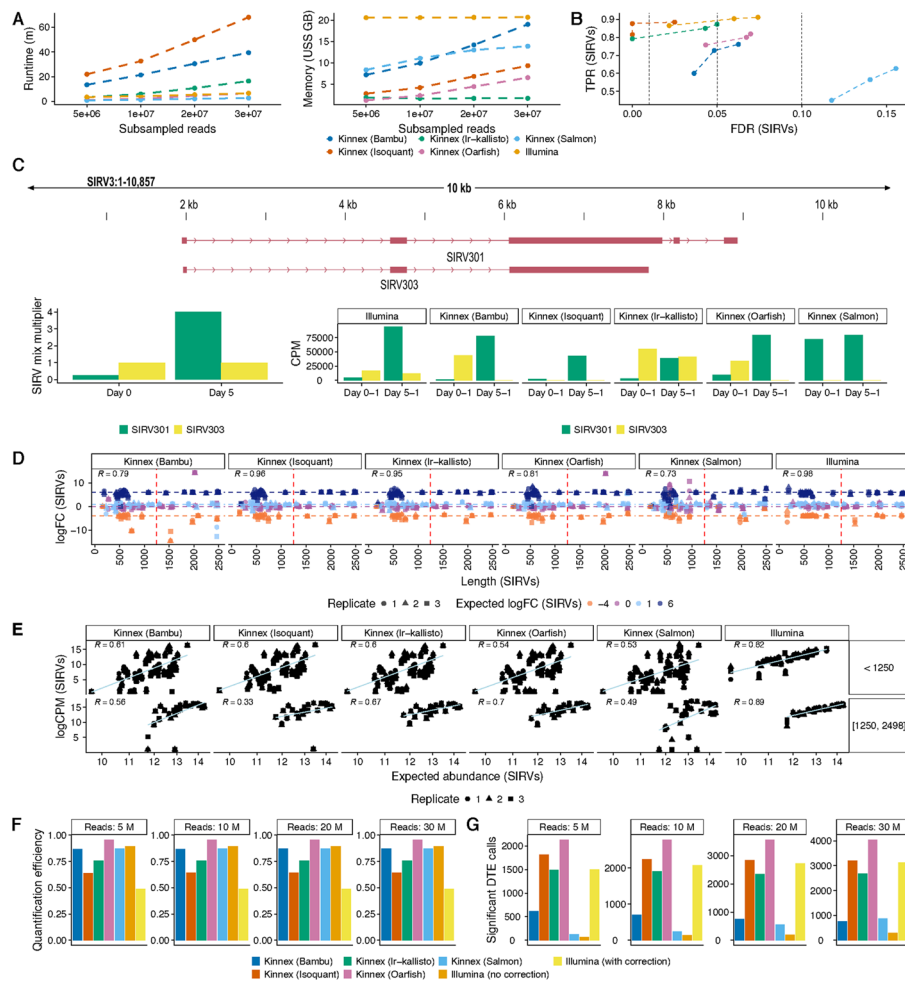


Fig. 2 Comparison between different Kinnex quantification methods. **A** Computational requirements of quantification methods for Kinnex and Illumina in terms of runtime (left) and memory (right) across all days and replicates of differentiation after downsampling the GENCODE-aligned data of each day to a fixed read depth (see [Methods](#)). Alignment and index creation were not included within computational requirements. **B** True positive rates (TPRs) and false discovery rates (FDRs) for Kinnex and Illumina quantification methods on Differential Transcript Expression (DTE) between Day 0 (E1 mix) and Day 5 (E2 mix) on the SIRV spike-ins when downsampling to a fixed depth of 1M reads for all samples (see [Methods](#)). **C** Top: Browser track highlighting an example of a pair of problematic SIRV transcripts that may cause outliers for some Kinnex lRNA-seq quantification methods. Bottom: Highlight of how the SIRV301 and SIRV303 SIRV transcript changes between the mixes utilized for Day 0 and Day 5 samples. While the SIRV303 transcript stays constant in expression between mixes, the SIRV301 transcript increases its expression drastically (16-fold), leading to most of the Kinnex methods attributing all expression to the SIRV301 sample in the Day 5 mix, since the SIRV303 transcript is fully contained within the SIRV301 transcript. **D** Relative quantification results between the same replicate on Day 0 (E1 mix) versus Day 5 (E2 mix) of four quantification methods for Kinnex lRNA-seq and Illumina as measured by similarity of observed and expected log-fold changes between SIRV spike-in mixes across SIRV transcript length when downsampling to a fixed depth of 1M reads for all samples (see [Methods](#)). **E** Absolute quantification of four quantification methods for Kinnex lRNA-seq and Illumina as measured by concordance of observed and expected absolute abundances for each replicate on Day 0 (E1 mix) when downsampling to a fixed depth of 1M reads for all samples (see [Methods](#)), stratified by SIRV transcript length. R denotes Pearson correlation. **F** Relative quantification efficiency, computed as number of counts divided by number of raw reads, of four quantification methods for Kinnex lRNA-seq and Illumina, stratified by the downsampling number of input reads. **G** Number of significant DTE calls that overlap with at least one other method at the same depth of four quantification methods for Kinnex lRNA-seq and Illumina, stratified by the downsampling number of input reads

of significant DTE calls aims to measure power in downstream tasks directly. Since our dataset does not have ground truth for DTE, we aim to remove false positive calls by only counting calls that are shared between two quantification methods (either within Kinnex or between Kinnex and Illumina).

For these metrics, Oarfish outperformed all other tools, including Illumina, achieving strong efficiency and the highest number of reproducible DTE calls (Fig. 2F, G). Together with its favorable run-time/memory profile and solid SIRV performance (the only primary deficit being the SIRV303 outlier), we selected Oarfish for all downstream Kinnex analyses.

Kinnex outperforms Illumina for DTE by resolving transcript-level read ambiguity

Using Oarfish for Kinnex and Salmon for Illumina, we compared gene-level differential expression (DGE) between platforms. Calls were highly concordant between Illumina and Kinnex across matched read depths (Fig. 3A), as expected, and previously reported for other long-read technologies [6].

(See figure on next page.)

Fig. 3 Quantification comparison between Kinnex and Illumina. **A** Overlap of DGE calls between Day 0 and Day 5 based on Oarfish Kinnex and Salmon Illumina when downsampling GENCODE-aligned data to a fixed read depth for each replicate. **B** Overlap of DTE calls between Day 0 and Day 5 based on Oarfish Kinnex IrRNA-seq and Salmon Illumina RNA-seq when downsampling GENCODE-aligned data to a fixed read depth for each replicate. **C** Example of a gene with a ‘flip’, characteristic of transcripts with high inferential variability in Illumina, where Illumina assigns moderate CPM values to a transcript in one replicate, while for another replicate, that transcript has close to zero CPM. **D** Transcript-model corresponding to the ‘flip’ example (**C**). **E** Inferential variability (see [Methods](#)) for GENCODE transcripts according to the number of transcripts per gene, stratified by technology. The y-axis indicates mean inferential relative variance across samples. Per sample, the inferential relative variance measures overdispersion; a value of 100 would indicate that the inferential sample variance is 101 times the inferential sample mean, indicating strong overdispersion. **F** Overlap of DTE calls between Day 0 and Day 5 based on Oarfish Kinnex IrRNA-seq and Salmon Illumina RNA-seq when downsampling GENCODE-aligned data to a fixed read depth for each replicate, with Illumina counts corrected for inferential variability (see [Methods](#)) [10]. **G** Number of detected transcripts and genes for each replicate of Day 0 and Day 5 when downsampling to 30 M reads, stratified by technology. **H** Example of a ‘division’, where Illumina evenly divides count mass between transcripts with high sequence similarity, whereas Kinnex assigns most count mass to a singular transcript. **I** Transcript-model corresponding to the ‘division’ example (**H**), with read support for both Illumina and Kinnex, highlighting that both technologies primarily show read support for the ENST00000263440 transcript. **J** Mean inferential variability per gene in Illumina by the $-\log_{10}(q)$ value of testing for multinomial discordance between Illumina and Kinnex, stratified by whether the two platforms detected an equal number of transcripts (equal detection) for the gene in question, or Illumina detected more transcripts (Illumina overdetected) or Kinnex detected more transcripts (Kinnex overdetected). **K** Number of transcripts per gene by the $-\log_{10}(q)$ value of testing for multinomial discordance between Illumina and Kinnex, stratified by whether the two platforms detected an equal number of transcripts (equal detection) for the gene in question, or Illumina detected more transcripts (Illumina overdetected) or Kinnex detected more transcripts (Kinnex overdetected). **L** Distribution of transcript length for transcripts that are detected by both technologies and transcripts detected only by Kinnex or Illumina, highlighting that Kinnex shows a clear length bias for transcripts shorter than 1.25 kb. **M** Number of transcripts missed by each technology and possible reasons for their non-detection, stratified by average CPM in the technology in which they were detected. **N** Relative quantification (log-fold change) concordance between Kinnex and Illumina on the gene and transcript level, for the first replicate, when filtered for transcripts and genes that show low inferential variability in Illumina and are detected by both platforms. **O** Absolute quantification (log CPM) concordance between Kinnex and Illumina on the gene and transcript level, for the first replicate, when filtered for transcripts and genes that show low inferential variability in Illumina and are detected by both platforms. Unless otherwise noted, Kinnex was downsampled to 30 M reads and quantified against GENCODE V45 using Oarfish. Illumina was also downsampled to 30 M reads per sample and quantified against the same transcriptome using Salmon (see [Methods](#)). R denotes Pearson correlation

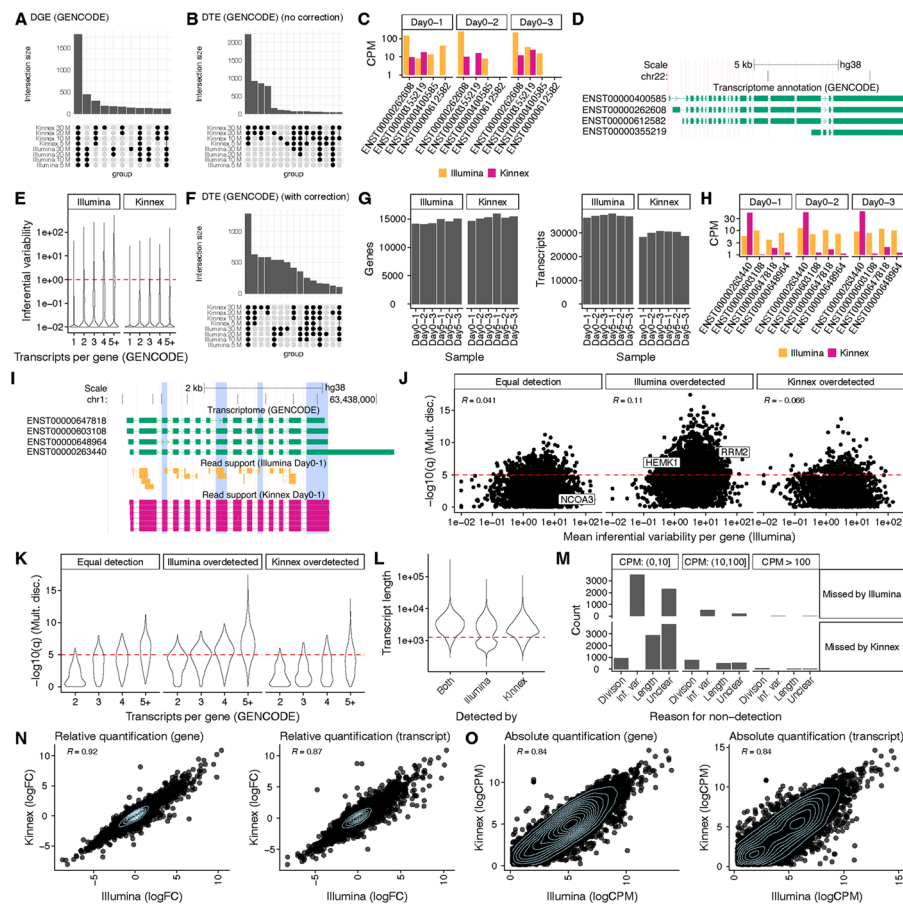


Fig. 3 (See legend on previous page.)

Next, we compared DTE calls between Kinnex and Illumina. Kinnex reported many more significant DTE events than Illumina (Fig. 3B), consistent with previous work comparing other long-read technologies with Illumina [10]. To explain this discrepancy, we manually inspected several genes in which Kinnex and Illumina quantification differed. In such cases, there was a tendency for the associated genes to express complex transcripts with highly similar sequences. For example, the *CECR2* gene, has four transcripts of highly similar sequence differing only by short exons and alternative acceptor sites (Fig. 3C, D). For many such cases, Illumina quantifications exhibited replicate-to-replicate fluctuations in abundances, sometimes dramatically assigning abundances that alternated between near-zero and mid-range values across replicates (“transcript flips”). Kinnex data, on the other hand, did not suffer from this problem, rather assigning expression to one or multiple dominant transcripts in a consistent manner (Fig. 3C, D).

To test whether the pattern observed in *CECR2* was an isolated case or a transcriptome-wide phenomenon, we quantified transcript measurement instability across the Kinnex and Illumina results using inferential variability [23]. Consistent with previous findings, we found that transcripts quantified by Illumina short reads exhibited higher inferential variability, as compared to those quantified by long-read technologies such

as Kinnex, in our case [6, 10] (Fig. 3E). This was particularly true for genes with a higher number of annotated transcripts (Fig. 3E).

The consequence of high inferential variability in Illumina data is that it leads to inflated dispersion estimates, ultimately diminishing statistical power to detect true DTE events [10]. This reduced power explains Illumina's initially small number of DTE calls. Indeed, transcripts with high inferential variability were disproportionately represented among those DTEs missed by Illumina but detected by Kinnex (Additional file 1: Fig. S3A).

Correcting for inferential variability improves Illumina quantification performance

To focus the benchmarking on transcripts that Illumina can quantify reliably, we first corrected Illumina-associated transcript counts for inferential variability, specifically by down-weighting transcripts with highly uncertain quantification [10]. After correction, the overlap between Illumina and Kinnex DTE calls improved substantially, both in the number of events captured and q-value distributions (Fig. 3F, Additional file 1: Fig. S3A). Accordingly, we employed inferential variability correction on all datasets in this study (including Fig. 1 [10]), except where variability itself was the subject of investigation (Fig. 3C, H, and Additional file 1: Fig. S3F-H).

Illumina's apparent transcript sensitivity reflects ambiguity, not true transcript diversity

Next, we evaluated the sensitivity with which each platform could detect transcripts. For each sample replicate, we determined the number of detected genes and transcripts, using 30 M downsampled reads (a transcript or gene was counted as detected if it had one CPM or higher in all samples of the respective day, see [Methods](#)). Kinnex detected a comparable number of genes on average, compared to Illumina (15,253 vs. 14,506 for Illumina), but notably fewer transcripts (29,721 vs. 37,163; Fig. 3G).

Upon closer investigation, we noticed examples in which the additional transcripts detected by Illumina may have been spurious. For the *ALG6* gene, for example, Kinnex identified a single dominant transcript (ENST00000263440), whereas Illumina reported expression that was divided more evenly across all four transcripts (Fig. 3H), despite limited unique read support for any of these transcripts (Fig. 3I). This suggests that transcripts detected by Illumina could be false positives caused by short-read ambiguity rather than genuine transcript diversity.

In these cases, the overall distribution of transcript proportions of genes reported by Illumina would differ from that reported by Kinnex. To quantify these discrepancies, we employed a multinomial discordance test, which assesses differences in transcript usage between platforms (i.e., DTU, see [Methods](#)). We found that gene-level discordance (q-value) showed no correlation with the mean inferential variability [23] (Fig. 3J).

Similarly to what we previously found with inferential variability, multinomial discordance tended to increase with the number of transcripts per gene, especially for genes for which Illumina detected more transcripts than Kinnex (Fig. 3K). Despite these overall trends, we still observed genes with high multinomial discordance but low inferential variability, and vice-versa, implying that transcript divisions and inferential noise are related but distinct sources of error (Additional file 1: Fig. S3B-D). Finally, consistent

with the earlier SIRV findings, Kinnex showed a slight length bias that was independent of the transcript-division issue (Fig. 3L).

Together, our analyses suggest at least two reasons why quantification differed between platforms. First, Illumina's short-read limitations led to unreliable quantifications for complex genes, manifested either as transcript flips across replicates or transcript division of expression among multiple similar transcripts. Second, Kinnex exhibited reduced detection efficiency for shorter transcripts, as seen in both spike-in controls and endogenous genes. Despite these biases, transcripts detected by only one platform tended to be lowly expressed, especially when no clear technical explanation was evident (Fig. 3M).

Kinnex and Illumina show strong concordance for reliably detected transcripts

Having established the sources of platform-specific transcript detection differences, we then evaluated how well Illumina and Kinnex quantifications agreed for the subset of transcripts reliably detected by both platforms (see [Methods](#)). This left 8,797 transcripts and 12,197 genes for comparison (note that the lower number of transcripts than genes is due to filtering for inferential variability). Overall, the two platforms agreed closely. Relative quantification (fold change estimates) was highly correlated with Pearson correlation values exceeding 0.9 at the gene level and approaching 0.9 at the transcript level (Fig. 3N). Absolute quantifications (logcounts per million (CPM) values) were somewhat less correlated, but still exceeded values of 0.8 at both the gene and transcript level (Fig. 3O). Taken together, within this unbiased subset, Kinnex quantified transcript expression as reliably as Illumina, and did so consistently across the range of transcript lengths (Additional file 1: Fig. S3E-F).

Kinnex's quantification performance remains good at approximately base-normalized depths

Lastly, we investigated whether the good performance of Kinnex data was solely due to the increase in bases relative to Illumina, at equal depth (Fig. 1D). For this, we repeated both the transcript discovery experiment on the SIRVs and the quantification comparisons between Illumina and Kinnex. While the median read lengths of Kinnex are ~5.2 times longer than Illumina, we performed the experiments at one tenth of the depth for the SIRV data and one sixth of the depth for Illumina data. This yielded 0.25 M reads for Kinnex and 2.5 M for Illumina for the SIRVs and 5 M reads for Kinnex and 30 M reads for Illumina for GENCODE. Overall, results were concordant with higher read numbers. In particular, Kinnex continued to outperform Illumina in transcript discovery (Additional file 1: Fig. S3G), and still achieved comparable correlation numbers in both absolute and relative quantification when compared to Illumina at a higher depth (Additional file 1: Fig. S3H-I). We believe the benefit of Kinnex, and long-read data more broadly, to lie in the lowering of inferential variability, relative to Illumina (Fig. 3E). This comes at a cost, given that long-reads are frequently sequenced at a lower depth than Illumina, which may especially affect downstream tasks such as DTE. Despite this, we did not find this to be the case here, given that Kinnex continued to perform well in DTE and DGE, even at a sixth of the depth of Illumina (Fig. 3A, F).

Discussion

Overall, our results provide evidence of strong transcript-level quantification performance for Kinnex. Similarly to other long-read technologies, Kinnex largely seemed to avoid problems of inferential variability and detected more DTE events than Illumina. Even with extensive correction and filtering, Illumina suffered from issues likely caused by short read mapping uncertainties, leading to phenomena such as transcript fluctuations and “flips” as well as artifactual “division” of expression across transcripts, which may erode DTE power. For transcripts reliably quantified across both platforms, Kinnex and Illumina showed high concordance, supporting Kinnex as an effective alternative for transcript-level studies. Nonetheless, Kinnex exhibited some limitations, particularly showing a bias against shorter transcripts and exhibiting a slightly lower performance in absolute compared to relative quantification. In addition, Kinnex and similar lrRNA-seq technologies are typically more expensive than Illumina for the same sequencing depth, though rapid improvements in chemistry and instrumentation are reducing this gap.

Our work extends several recent works focused on evaluating the quantification and transcript discovery capability of lrRNA-seq methods [4–7]. In contrast to previous efforts, we introduce a large-scale Kinnex dataset, sample-matched with Illumina, which newly enables the benchmarking of PacBio data at a similar depth to ONT and Illumina. In addition, we evaluate the quantification performance of Kinnex and highlight its strengths and weaknesses, in particular its good performance at the transcript level through avoiding inferential variability-related pitfalls and its length bias.

Conclusions

These results position Kinnex as a practical, reliable choice for projects that must resolve complex transcript architecture and perform differential expression transcript analyses. Future work should add matched ONT data, deeper coverage datasets across more diverse biological systems, and emerging single-cell long-read methods.

Methods

Dataset generation

Stem cell culture

Undifferentiated human-induced pluripotent stem cells (hiPSCs, WTC11, NIGMS Repository Number GM25256) were obtained through Coriell via a Materials Transfer Agreement with the University of Virginia. Mycoplasma testing was performed using Bioluminescent assay for mycoplasma detection, and the result was negative. GM25256 (WTC11) was purchased from the Coriell Cell Repository and characterized—including karyotypic analysis, species verification by LINE assay, surface antigen expression, alkaline phosphatase activity, pluripotency via embryoid body formation, and mRNA expression profiling—while UVA did not perform additional cell line identification.

WTC11 cells were thawed from Cryostor (Stem Cell Technologies) and cultured using mTESR Plus Media (Stem Cell Technologies) on Matrigel (Corning) coated plates with media changes following manufacturer recommendations. Cells were cultured to maintain undifferentiated cell populations and passaged when confluent

using ReLeSR (Stem Cell Technologies) and replated at desired densities per manufacturer guidelines on Matrigel-coated dishes.

Stem cell (WTC11) derived primordial endothelial cells

WTC11 cells were seeded with mTESR Plus Media (Stem Cell Technologies) onto 6-well dishes or 10 cm Matrigel-coated plates for primordial endothelial cell differentiation as described previously [8]. Twenty-four hours after plating, on Day 0 media was aspirated and replaced with Diff APEL 2 (Stem Cell Technologies) differentiation media containing 5 μ M GSK3i (CHIR99021, Reprocell). On Day 1 media was aspirated and replaced with differentiation media containing 50 ng/mL essential fibroblast growth factor (bFGF, R&D Systems). On Days 2, 3, and 4 media was aspirated and replaced with differentiation media containing 25 ng/mL of bone morphogenetic protein 4 (BMP4, R&D Systems) and 50 ng/mL of vascular endothelial growth factor VEGF (Fisher Scientific). Cells were collected from Days 0–5 of the differentiation protocol using Accutase (ThermoFisher) per manufacturer guidelines and pelleted. Three biological replicates collected on Day 0, with Day 0–1 and Day 0–2 collected from a 10 cm dish and Day 0–3 from a 6-well dish. Three biological replicates were collected as described above for Day 5, with Day 5-1 and Day 5-2 derived from a 10 cm dish and Day 5-3 from a 6-well plate. Two biological replicates were obtained from Day 3 from 10 cm dishes. Days 1, 2, and 4 were all collected from 10 cm dishes. For each sample, RNA was extracted at approximately 1 million cells using an RNeasy (Qiagen) kit, with quality assessment via Bioanalyzer. Spike-In RNA Variants (SIRVs, Lexogen) were reconstituted per manufacture guidelines and added to samples with Day 0 biological replicates receiving E1 SIRVs, Day 5 E2, and all other Days receiving S4. Master mixes of RNA and SIRVs were split into 4 tubes for parallel long-read Pacific Biosciences (PB) Revio sequencing and short-read sequencing on the NovaSeq at 150 bp.

qPCR gene marker analysis

Transcriptomes were collected from cultured WTC11 cells at every day of the differentiation protocol using the RNeasyPlus Micro Kit (QIAGEN). Resulting cDNA libraries were assessed via qPCR using published PCR primers (Integrated DNA Technologies, based on [8, 24]) and miScript SYBR Green PCR Kit (QIAGEN) to validate gene expression patterns associated with the intermediate states.

Short-read RNA-sequencing

Aliquots of total RNA samples created from the master mix preparation as described above, were used as input for sequencing to Kapa RNA HyperProKits (polyA-enriched) for short-read RNA sequencing library preparation. The resulting libraries were sequenced on the NovaSeq at 150bp.

Kinnex long-read RNA sequencing library preparation and long-read RNA sequencing run

Aliquots of total RNA samples created from the master mix preparation as described above were used as input for PB Kinnex. From the RNA, cDNA was synthesized using the Iso-Seq Express Kit (PB). Approximately 300 ng of cDNA from each sample was barcoded and equally pooled. All replicates in one WTC11 sample were made into one Kinnex full-length library and sequenced on one SMRTCell (i.e., WTC11 Day 0 and WTC11 Day 5 replicates were sequenced in one run).

Data analysis***Preprocessing***

To prepare raw sequencing files for alignment, we performed different preprocessing steps per technology.

Illumina. We ran Trim Galore 0.6.2 in paired-end mode with a minimum length of 20 and a minimum quality of 20 using the `--phred33` option.

Kinnex. Kinnex data were preprocessed on-instrument. Basecalling was performed using basecaller version 5.0 and otherwise standard parameters. CCS reads were generated using ccs 7.0.0, followed by skera demultiplexing of Kinnex S-reads using skera 1.0.99, lima 2.8.99, and Isoseq refine 3.99.99, at which point reads were ready for alignment.

Separating FASTQ into spike-in and non-spike-in data

As an initial step, both for QC and since spike-ins may require different alignment options, we separated all FASTQ files into reads aligning to either one of the artificial spike-in genomes or to the human genome.

Alignment for Illumina-generated RNA-sequencing files Data was aligned to a concatenation of the human genome, the full SIRV genome, and the ERCC genome using STAR 2.7.11b [25]. Afterward, BAM files were separated by whether they aligned to the human or one of the artificial genomes and converted to FASTQ using samtools 1.19.2 [26].

Alignment for Kinnex-generated RNA-sequencing files Data was aligned to a concatenation of the human genome, the full SIRV genome, and the ERCC genome using minimap2 2.28-0 with the `splice:hq -uf` options [27]. Afterward, BAM files were separated by whether they aligned to the human or one of the artificial genomes and converted to FASTQ using samtools 1.19.2.

Genome alignment (Kinnex)

Genome alignment was performed using minimap2 2.28-0.28.28.28 with the `splice:hq -uf` options and provided annotated splice sites using `-junc-bed`. For the alignment of SIRV reads, we also applied the `--splice-flank` option.

Transcriptome alignment (Kinnex)

Transcriptome alignment was performed using minimap2 2.28-0 with the `map-hifi --eqx` options, allowing for a maximum of 100 secondary alignments by using the `-N 100` option.

Genome alignment (Illumina)

Genome alignment was performed using STAR 2.7.11b using the `--outSAMstrandField intronMotif` argument and otherwise default parameters. Indices for STAR were generated using the `--sjdbGTFfile` option to which appropriate transcriptome GTF files were passed. The `--sjdbOverhang` option was set to 150.

Transcriptome alignment (Illumina)

Transcriptome alignment was performed using STAR 2.7.11b using default parameters. Indices for STAR were generated using default parameters.

References

We used the GRCh38 genome (GCA_000001405.15_GRCh38_no_alt_analysis_set.fna.gz) as the human reference genome and the primary assembly of GENCODE V45 as the human reference transcriptome [28]. For the spike-ins, we used the genome and transcriptome corresponding to SIRV set four, including short SIRVs, long SIRVs, and ERCCs.

Quantification (Kinnex)

Isoquant. isoquant 3.6.3 was used to quantify genome alignments with precomputed DB files, never requiring a polyA tail, data type `pacbio_ccs`, and otherwise default parameters [20].

lr-kallisto. kallisto 0.51.1 was used in conjunction with bustools 0.44.1 to quantify FASTQ files using an index k-mer setting of 63, a bus threshold of 0.8, and PacBio as the platform (`-P PacBio`) [21].

Oarfish. oarfish 0.7.1 was run to quantify transcriptome alignments with model coverage applied, no filters, and otherwise default parameters [22].

Bambu. Bambu 3.4.0 was run to quantify genome alignments with default parameters [16].

Salmon. Salmon 1.10.3 was run to quantify transcriptome alignments with the `--ont` flag and otherwise default parameters.

Quantification (Illumina)

For Illumina data, we used Salmon 1.10.3 to quantify FASTQ files using an index k-mer setting of 31, a fragment length prior mean of 250, a fragment length standard deviation of 25, the `--validateMappings --gcBias --seqBias` flags, and otherwise default parameters [19].

Inferential variability correction

For Illumina, we corrected for inferential variability using the `catchSalmon` function in edgeR 4.4.0. Kinnex quantifications were not corrected for inferential variability [10, 29].

Gene-level quantifications

Gene-level quantifications were generated by summing all transcript-level quantifications for all transcripts belonging to a gene, according to a particular annotation.

Quality control

Read count numbers. Read count numbers were calculated using samtools 1.19.2. For aligned data, we counted primary alignments after separation into SIRV and GENCODE data.

Read length and number of bases. Read length was calculated using bioawk 1.0 as the number of nucleotides in each read for Kinnex and the sum of nucleotides in each paired-end read for Illumina. Number of bases was calculated as the sum of read lengths for each technology.

Number of covered junctions. Covered junctions were calculated as the number of introns of length 20 or longer in the cigar string of the genome alignments of each technology using pysam 0.23.0.

Read quality and relative edit distance. Read quality was calculated as the average Phred value of all bases of that read for Kinnex or the average Phred of all bases in each paired-end read for Illumina using bioawk 1.0.

Relative edit distance, SNV rate, and indel rate were calculated separately from BAM files using pysam 0.23.0. For Kinnex, edit distance was calculated as the NM tag divided by the aligned length per read, SNV rate as the number of mismatches from the MD tag divided by the aligned length, and indel rate as the total length of insertions and deletions from the CIGAR string divided by the aligned length. For Illumina, edit distance was calculated as the sum of NM tags from both paired-end reads divided by the sum of aligned lengths, SNV rate as the sum of mismatches from both MD tags divided by the sum of aligned lengths, and indel rate as the sum of insertion and deletion lengths from both CIGAR strings divided by the sum of aligned lengths.

Gene body coverage. Gene body coverage was calculated on genome-aligned BAM files using randomly subsampled GENCODE transcripts using RSeQC 5.0.3 [30]. Gene body coverage was calculated per transcript and within size bins. For each size bin, we randomly sampled 2,500 GENCODE transcripts.

Fragment size. For Illumina data, fragment size was calculated on transcriptome-aligned BAM files by using the `TLEN` attribute of the BAM file using pysam 0.23.0.

Transcript and gene filtering

Depending on the analysis, we performed different types of transcript or gene filtering. SIRV transcripts were never filtered.

Replicability (Fig. 1I, J). Only transcripts or genes having at least one CPM in at least three samples in both platforms were kept.

Detection differences between technologies (Figs. 1L, 3J–M). A transcript or gene was considered as detected if it had at least one CPM in all three replicates of Day 0.

QC detection (Fig. 3G). A transcript or gene was considered as detected if it had at least one CPM in the sample in question.

Differential analyses. For all differential analyses, we used the `filterByExpr` function of edgeR 4.4.0 with default parameters.

Technology detection (Fig. 3J–M). Transcripts were considered as detected if they had at least one CPM in all three replicates of day zero.

Relative quantification concordance (Fig. 3N). Transcripts were kept if they had inferential variability less than or equal to one for the Illumina data, and they had at least one CPM in all three replicates of Day 0 in both technologies. Requirements for genes were identical except that no inferential variability filter was applied.

Absolute quantification concordance (Fig. 3O). Transcripts were kept if they had inferential variability less than or equal to one for the Illumina data, and they had at least one CPM in all three replicates of Day 0 in both technologies. Requirements for genes were identical except that no inferential variability filter was applied.

Transcript discovery (GENCODE)

Bambu 3.4.0 was run to perform transcript discovery using Kinnex genome alignments, selecting the Bambu-suggested NDR and keeping ISM novel transcripts (`remove.subsetTx = FALSE`).

Transcript discovery (SIRVs)

Transcript discovery on the SIRVs was run on SIRV genome alignments, downsampled to 2.5 M reads for both platforms, unless otherwise specified. Genome alignments were generated without providing the location of existing SIRV splice junctions to the aligner.

Discovery for Illumina was performed using stringtie 2.2.3 with default parameters and without a reference transcriptome. Discovery for Kinnex was performed using Bambu 3.4.0 in de-novo mode, with an NDR of 1.0 and keeping ISM novel transcripts (`remove.subsetTx = FALSE`) [31].

Performance was calculated using gffcompare 0.12.6 by comparing the predicted transcriptomes of both technologies to the ground truth transcriptome of the short SIRVs [32]. Performance metrics were reported at the intron chain level.

Annotation of novel transcripts (GENCODE)

Novel transcripts were annotated using SQANTI3 5.1.2. For determining whether a particular novel transcript was coding or non-coding, we ran ORFAnage 1.1.0 in BEST mode, using GENCODE V45 CDS as a reference, and determined a novel transcript as coding if it was assigned a CDS by ORFAnage [2, 17].

Inferential variability

We used fishpond 2.12.0 for all analyses regarding inferential variability. In particular, we used computeInRV to calculate inferential variability as the mean inferential relative variance of each transcript across samples [23].

Quantification of transcript mass division between Illumina and Kinnex

To quantify differences in the division of transcript mass for a particular gene between Illumina and Kinnex, we performed a test for multinomial discordance between Illumina and Kinnex. In particular, DTU was performed with edgeR's glmQLFit and diffSpliceDGE across all replicates of day zero of both Illumina and Kinnex, accounting for both technology and replicate. We tested significance on the technology coefficient. Transcript-level q-values were either used directly or aggregated to the gene level using Simes' adjustment.

Definition of categories leading to lack of detection (Fig. 3M)

Transcripts missed by Kinnex.

1. *Length*: Transcripts that were shorter than 1,250 nucleotides
2. *Division*: Transcripts longer than 1,250 nucleotides but having a $-\log_{10}(q)$ of multinomial discordance between Illumina and Kinnex greater than five
3. *Unclear*: All other transcripts not detected by Kinnex but detected by Illumina

Transcripts missed by Illumina.

1. *Inf. var.*: Transcripts with a mean inferential variability value greater than one
2. *Unclear*: All other transcripts not detected by Illumina but detected by Kinnex

Downsampling

For optimal comparability in some analyses, we downsampled read FASTQ or alignment BAM files without replacement.

SIRV files were downsampled to 0.25 M, 0.5 M, 1 M, and 2.5 M, while GENCODE files were downsampled to 5, 10, 20, and 30 M, respectively. For the SIRV downsamplings, we performed five downsamplings per target depth with different random seeds, while for GENCODE only one downsampling was performed.

BAM files for genome and transcriptome alignments were downsampled using a script leveraging Rsamtools 2.18.0 that kept all alignments corresponding to a source read.

FASTQ files were downsampled using seqtk 1.4, ensuring contiguity for paired-end Illumina reads.

Differential analyses (GENCODE)

Differential Gene Expression (DGE) was performed between all replicates of Day 0 and Day 5, using edgeR's `glmQLFit`, accounting for both day and replicate. We tested significance on the day coefficient and called genes as significant that had a q-value less than 0.01.

DTE was performed between all replicates of Day 0 and Day 5, using edgeR's `glmQLFit`, accounting for both day and replicate. We tested significance on the day coefficient and called transcripts as significant that had a q-value less than 0.01.

DTU was performed across all replicates across the timecourse (excepting Day2-1) using a cubic spline with edgeR's `glmQLFit` and `diffSpliceDGE`, accounting for both day and replicate. We tested significance on all spline coefficients jointly and called genes as significant that had a q-value less than 0.01, based on Simes adjustment.

Differential analyses (SIRVs)

We performed DTE on the SIRV datasets between all replicates of Day 0 and Day 5, corresponding to SIRV mixes E1 and E2, respectively. In particular, we used the true log-fold changes between SIRV mixes E1 and E2 as the ground-truth, yielding either 0 (no change between mixes), -1 (decrease from E1 to E2), or 1 (increase from E1 to E2).

We then performed DTE using edgeR's `glmQLFit`, with a design matrix accounting for both the mix and the replicate, testing the significance of the mix coefficient. Predictions were extracted using limma's `decideTests`, using an FDR cut-off of 1% and no log-fold change cutoff [33]. The FDR was calculated as the percentage of SIRVs found as -1 or 1 that were either 0 or the opposing sign in the ground-truth (i.e., false discoveries) over all discoveries (i.e., all SIRVs that had a non-zero value assigned). The TPR was calculated as the percentage of SIRVs that had a ground-truth non-zero value, were deemed significant, and had the same sign as the ground-truth over all ground-truth non-zero values.

All quantifications were downsampled to a depth of one million reads for fairness. FDR and TPR values correspond to mean values calculated across all five downsamplings.

Relative quantification efficiency

Relative quantification efficiency was assessed for each quantification method by calculating the average proportion of quantified counts to raw input reads at several downsampled depths for the GENCODE data (5 M, 10 M, 20 M, 30 M).

Definition of significant DTE calls (Fig. 2G)

We considered any call that had an FDR less than 0.01 and was called by at least one other quantification method at the same depth as significant.

Computational requirements

Computational requirements were measured using the Snakemake benchmark functionality and are reported over each of the three replicates for days zero and five. Since computational requirement variability was minimal for reruns, we do not report variability within replicate and day. All methods were run with access to 12 cores.

Reproducibility

All bioinformatics experiments were performed on an Ubuntu 22.04.5 LTS server with AMD EPYC 7742 CPUs. Experiments were run with Snakemake 8.2.3 within rule-appropriate mamba environments within Singularity containers [34]. All experiments are reproducible up to numerical stability via standard Snakemake commands documented on our Github (see Data availability).

All analyses were performed in respective command line tools or in R 4.4.3 [35]. Figures were generated using ggplot2 3.5.2 and ggpubfigs 1.1.0 [36, 37].

Supplementary Information

The online version contains supplementary material available at <https://doi.org/10.1186/s13059-026-03988-1>.

Additional file 1: Fig. S1. Quality control of all samples in our dataset highlights quality issues in Day2-1 for Kinnex. Fig. S2. Biological exploration on the iPSC-EC differentiation dataset reveals differential splicing, differential expression and a moderate amount of novel isoforms. Fig. S3. Performance of Kinnex relative to Illumina on DTE, and relative and absolute quantification, including when downsampled to approximately base-normalized depth.

Additional file 2: Table S1. Number of expressed genes and transcripts across downsampled sequencing depths (5M, 10M, 20M, and 30M reads), calculated at a constant count threshold of 5. Illumina detection was performed on downsampled counts to allow for normalized comparison with Kinnex.

Additional file 3: Table S2. Differential Transcript Usage (DTU) results based on full-depth Oarfish quantifications of the GENCODE V45 transcriptome plus novel transcripts identified by Bambu.

Additional file 4: Table S3. Differential Gene Expression (DGE) results based on full-depth Oarfish quantifications of the GENCODE V45 transcriptome plus novel transcripts identified by Bambu.

Additional file 5: Table S4. SQANTI annotation and characterization of full-depth GENCODE V45 plus novel transcripts identified by Oarfish. This dataset includes structural classifications, quality attributes, and expression levels normalized as Counts Per Million (CPM).

Acknowledgements

We would like to acknowledge Pacific Biosciences of California, Inc., for providing sequencing of PB Kinnex lrrNA-seq on the Revo free of charge. Early results based on the dataset introduced here were presented in a webinar organized by Pacific Biosciences of California, Inc. We are grateful to webinar attendees for their helpful feedback and discussions. Figure 1A was created with Biorender.com (publication license available on demand).

Peer review information

Andrew Cosgrove and Claudia Feng were the primary editors of this article and managed its editorial process and peer review in collaboration with the rest of the editorial team. The peer-review history is available in the online version of this article.

Authors' contributions

G.M.S. and M.D.R. conceived of the study design, and M.M. implemented the cell-based experiments. M.M. and V.P. performed experimental work and cell collection. M.M. and G.M.S. coordinated sequencing of the Illumina RNA-seq samples. E.T. coordinated the sequencing of PB Kinnex lrrNA-seq samples. D.W., M.M., and K.N. wrote the analysis code. D.W. wrote the Snakemake workflow. D.W. and M.M. analyzed the results. M.M. and D.W. drafted the original manuscript. E.T., M.D.R., and G.M.S. co-supervised the work. All authors reviewed, edited, and approved the final manuscript.

Funding

This work was supported by a Swiss National Science Foundation (SNSF) project grant 204869 to M.D.R., the Robert M. Berne Cardiovascular Research Center Training Program (T32HL007284), the Wagner Fellowship (UVA) to M.M.M., and a National Institute of General Medical Sciences (NIGMS), National Institutes of Health (NIH) grant R35GM142647 to G.M.S.

Data availability

For access to the raw sequencing files for Illumina and Kinnex, see University of Virginia (UVA), Integrated Translational Health Research Institute of Virginia (iTHRIV). Researchers who wish to gain access to the source data can contact the Sheynkman Lab to initiate a data request (gs9yr@virginia.edu). Before sharing, both groups will seek approval from the

NIGMS repository by completing the Statement of Research Intent forms. Once approved (typical timeline: days-weeks), a contract can be initiated between institutions for the transfer and reuse of the source data in this study. Count tables, separated into spike-in and non-spike-in data for the GENCODE reference transcriptome, and the GENCODE reference transcriptome augmented with Kinnex lRNA-seq-derived novel transcripts for both Illumina and four quantification methods for PB Kinnex (Methods), along with related paper outputs have been deposited to Zenodo [38]. All code to reproduce experiments, figures, and tables is available on Github and Zenodo under an MIT license [39, 40]. Undifferentiated human-induced pluripotent stem cells (hiPSCs, WTC11, NIGMS Repository Number GM25256) were obtained through Coriell via a Materials Transfer Agreement with the University of Virginia [41]. Spike-in controls were obtained from Lexogen (SIRV Set 4) [42]. The reference transcriptome was obtained from GENCODE (V45) [43] and the reference genome from NCBI [44].

Declarations

Ethics approval and consent to participate

The study was conducted in accordance with the Declaration of Helsinki and was reviewed by the University of Virginia Institutional Review Board for Health Sciences Research (IRB-HSR) and received an exempt determination (Study #301733).

Consent for publication

Not applicable.

Competing interests

E.T. is employed by and holds stock in Pacific Biosciences of California, Inc. G.M.S. is Chief Scientific Officer and holds equity in NeoSplice Therapeutics, Inc., and serves on the scientific advisory board of Quantum-Si Incorporated and holds stock in Quantum-Si Incorporated. Pacific Biosciences of California, Inc. ran sequencing of PB Kinnex lRNA-seq on the Revio free of charge.

Received: 8 July 2025 Accepted: 29 January 2026

Published online: 25 February 2026

References

- Monzó C, Liu T, Conesa A. Transcriptomics in the era of long-read sequencing. *Nat Rev Genet.* 2025;1–21.
- Pardo-Palacios FJ, Arzalluz-Luque A, Kondratova L, Salguero P, Mestre-Tomás J, Amorín R, et al. SQANTI3: curation of long-read transcriptomes for accurate identification of known and novel isoforms. *Nat Methods.* 2024;21(5):793–7.
- Al'Khafaji AM, Smith JT, Garimella KV, Babadi M, Popic V, Sade-Feldman M, et al. High-throughput RNA isoform sequencing using programmed cDNA concatenation. *Nat Biotechnol.* 2024;42(4):582–6.
- Pardo-Palacios FJ, Wang D, Reese F, Diekhans M, Carbonell-Sala S, Williams B, et al. Systematic assessment of long-read RNA-seq methods for transcript identification and quantification. *Nat Methods.* 2024;21(7):1349–63.
- Chen Y, Davidson NM, Wan YK, Yao F, Su Y, Gamaarachchi H, et al. A systematic benchmark of Nanopore long-read RNA sequencing for transcript-level analysis in human cell lines. *Nat Methods.* 2025;1–12.
- Dong X, Du MR, Gouil Q, Tian L, Jabbari JS, Bowden R, et al. Benchmarking long-read RNA-sequencing analysis tools using in silico mixtures. *Nat Methods.* 2023;20(11):1810–21.
- Su Y, Yu Z, Jin S, Ai Z, Yuan R, Chen X, et al. Comprehensive assessment of mRNA isoform detection methods for long-read sequencing data. *Nat Commun.* 2024;15(1):3972.
- Nelson EA, Qiu J, Chavkin NW, Hirschi KK. Directed Differentiation of Hemogenic Endothelial Cells from Human Pluripotent Stem Cells. *J Vis Exp.* 2021;(169).
- Soneson C, Yao Y, Bratus-Neuenschwander A, Patrignani A, Robinson MD, Hussain S. A comprehensive examination of Nanopore native RNA sequencing for characterization of complex transcriptomes. *Nat Commun.* 2019;10(1):3359.
- Baldoni PL, Chen Y, Hediye-Zadeh S, Liao Y, Dong X, Ritchie ME, et al. Dividing out quantification uncertainty allows efficient assessment of differential transcript expression with edgeR. *Nucleic Acids Res.* 2024;52(3):e13.
- Zhang C, Li H, Wang S. Common gene signatures and molecular mechanisms of diabetic nephropathy and metabolic syndrome. *Front Public Health.* 2023;11:1150122.
- Romero MF, Chen AP, Parker MD, Boron WF. The SLC4 family of bicarbonate HCO transporters. *Mol Aspects Med.* 2013;34(2–3):159–82.
- Liao Z, Cantor JM. Endothelial cells require CD98 for efficient angiogenesis-brief report. *Arterioscler Thromb Vasc Biol.* 2016;36(11):2163–6.
- Nichol D, Stuhlmann H. EGFL7: a unique angiogenic signaling factor in vascular development and disease. *Blood.* 2012;119(6):1345–52.
- Yeh YT, Hur SS, Chang J, Wang KC, Chiu JJ, Li YS, et al. Matrix stiffness regulates endothelial cell proliferation through septin 9. *PLoS ONE.* 2012;7(10):e46889.
- Chen Y, Sim A, Wan YK, Yeo K, Lee JX, Ling MH, et al. Context-aware transcript quantification from long-read RNA-seq data with Bambu. *Nat Methods.* 2023;20(8):1187–95.
- Varabyou A, Erdogdu B, Salzberg SL, Pertea M. Investigating open reading frames in known and novel transcripts using ORFAnage. *Nat Comput Sci.* 2023;3(8):700–8.
- Amarasinghe SL, Ritchie ME, Gouil Q. long-read-tools.org: an interactive catalogue of analysis methods for long-read sequencing data. *GigaScience.* 2021;10(2):giab003.
- Patro R, Duggal G, Love MI, Irizarry RA, Kingsford C. Salmon provides fast and bias-aware quantification of transcript expression. *Nat Methods.* 2017;14(4):417–9.

20. Pribelski AD, Mikheenko A, Joglekar A, Smetanin A, Jarroux J, Lapidus AL, et al. Accurate isoform discovery with IsoQuant using long reads. *Nat Biotechnol.* 2023;41(7):915–8.
21. Loving RK, Sullivan DK, Reese F, Rebboah E, Sakr J, Rezaie N, et al. Long-read sequencing transcriptome quantification with Ir-kallisto. *PLoS Comput Biol.* 2025;21:e1013692.
22. Zare Jousheghani Z, Singh NP, Patro R. Oarfish: enhanced probabilistic modeling leads to improved accuracy in long read transcriptome quantification. *Bioinformatics.* 2025;41 Supplement_1:i304–13.
23. Zhu A, Srivastava A, Ibrahim JG, Patro R, Love MI. Nonparametric expression analysis using inferential replicate counts. *Nucleic Acids Res.* 2019;47(18):e105.
24. Qiu J, Nordling S, Vasavada HH, Butcher EC, Hirschi KK. Retinoic Acid Promotes Endothelial Cell Cycle Early G1 State to Enable Human Hemogenic Endothelial Cell Specification. *Cell Rep.* 2020;33(9):108465.
25. Dobin A, Davis CA, Schlesinger F, Drenkow J, Zaleski C, Jha S, et al. STAR: ultrafast universal RNA-seq aligner. *Bioinformatics.* 2013;29(1):15–21.
26. Li H, Handsaker B, Wysoker A, Fennell T, Ruan J, Homer N, et al. The sequence alignment/map format and SAMtools. *Bioinformatics.* 2009;25(16):2078–9.
27. Li H. Minimap2: pairwise alignment for nucleotide sequences. *Bioinformatics.* 2018;34(18):3094–100.
28. Frankish A, Carbonell-Sala S, Diekhans M, Jungreis I, Loveland JE, Mudge JM, et al. GENCODE: reference annotation for the human and mouse genomes in 2023. *Nucleic Acids Res.* 2023;51(D1):D942–9.
29. Robinson MD, McCarthy DJ, Smyth GK. edgeR: a Bioconductor package for differential expression analysis of digital gene expression data. *Bioinformatics.* 2010;26(1):139–40.
30. Wang L, Wang S, Li W. RSeQC: quality control of RNA-seq experiments. *Bioinformatics.* 2012;28(16):2184–5.
31. Pertea M, Pertea GM, Antonescu CM, Chang TC, Mendell JT, Salzberg SL. StringTie enables improved reconstruction of a transcriptome from RNA-seq reads. *Nat Biotechnol.* 2015;33(3):290–5.
32. Pertea G, Pertea M. GFF utilities: GffRead and GffCompare. *F1000Research.* 2020;9:ISCB–Comm.
33. Ritchie ME, Phipson B, Wu D, Hu Y, Law CW, Shi W, et al. limma powers differential expression analyses for RNA-sequencing and microarray studies. *Nucleic Acids Res.* 2015;43(7):e47.
34. Köster J, Rahmann S. Snakemake—a scalable bioinformatics workflow engine. *Bioinformatics.* 2012;28(19):2520–2.
35. R Core Team. R: A Language and Environment for Statistical Computing. Vienna, Austria: R Foundation for Statistical Computing; 2025. <https://www.R-project.org/>.
36. Wickham H. ggplot2: Elegant Graphics for Data Analysis. Springer-Verlag New York; 2016. <https://ggplot2.tidyverse.org>.
37. Steenwyk JL, Rokas A. ggpubfigs: colorblind-friendly color palettes and ggplot2 graphic system extensions for publication-quality scientific figures. *Microbiol Resour Announc.* 2021;10(44):10–1128.
38. Wissel D, Mehlferber M, Nguyen K, Pavelko V, Tseng E, Robinson M, et al. A Systematic Benchmark of High-Accuracy PacBio Long-Read RNA Sequencing for Transcript-Level Quantification. Zenodo; 2026. Version v4. <https://doi.org/10.5281/zenodo.15763514>.
39. Wissel D, Mehlferber M, Nguyen K, Pavelko V, Tseng E, Robinson M, et al. A Systematic Benchmark of High-Accuracy PacBio Long-Read RNA Sequencing for Transcript-Level Quantification. Zenodo. 2026. <https://doi.org/10.5281/zenodo.15763514>.
40. Wissel D, Mehlferber M, Nguyen K, Pavelko V, Tseng E, Robinson M, et al. A Systematic Benchmark of High-Accuracy PacBio Long-Read RNA Sequencing for Transcript-Level Quantification. GitHub; 2026. <https://github.com/dnwissel/kinnex-data>.
41. Coriell Institute for Medical Research. WTC11 Induced Pluripotent Stem Cell Line (GM25256). NIGMS Human Genetic Cell Repository; 2016. Human-induced pluripotent stem cells (hiPSCs) used for endothelial differentiation. https://catalog.coriell.org/0/Sections/Search/Sample_Detail.aspx?Ref=GM25256.
42. Lexogen GmbH. SIRV Set 4 (Spike-In RNA Variant Control Mixes). Lexogen; 2021. Spike-in controls used for transcript discovery and quantification benchmarking. <https://www.lexogen.com/sirvs/>.
43. GENCODE Project. GENCODE Reference Annotation for the Human Genome (V45). GENCODE; 2023. Primary assembly used for transcript-level quantification. <https://www.gencodegenes.org/human/>.
44. Genome Reference Consortium. Human Genome Assembly GRCh38 (GCA_000001405.15). National Center for Biotechnology Information (NCBI); 2013. Analysis set without alt contigs used for alignment. https://www.ncbi.nlm.nih.gov/datasets/genome/GCA_000001405.15/.

Publisher's Note

Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.