

Long-read proteogenomic atlas of human neuronal differentiation reveals isoform diversity informing neurodevelopmental risk mechanisms

Corresponding Author: Dr Shreejoy Tripathy

This file contains all reviewer reports in order by version, followed by all author rebuttals in order by version.

Version 0:

Reviewer comments:

Reviewer #1

(Remarks to the Author)

This manuscript by Xu et al represents an important advance in human neuronal transcriptomics on several fronts. It provides a deep dataset of long read RNA-Sequencing during human iPSC neuronal differentiation, uncovering thousands of novel transcript isoforms. These isoforms are investigated with relation to ASD risk and mutational constraint. Widespread coordination between alternative splicing events and alternative 3'UTR selection is uncovered. Support for some novel isoforms is demonstrated using proteomics. The web application for viewing full length isoforms and abundances is very useful. I have the following questions that should be addressed prior to publication.

Major points:

Question 1: Line 83. The reader might not know about the Kinnex library prep details and thus miss the fact that polyA containing mRNA is sequenced (not total RNA).

Question 2: 689 million long reads (line 95) does seem very deep for PacBio data. But the reader has to take the authors' word that it is one of the most deeply sequenced to date. How do other recent studies compare? There are long read datasets from the Encode project, for example: Reese et al., bioRxiv, "The ENCODE4 long-read RNA-seq collection reveals distinct classes of transcript structure diversity". A comparison of novel isoform detection across a few other recent datasets might help support the annotation of the novel transcripts and even give insights into cell type expression patterns.

Question 3: There are no plots shown with actual reads. Showing representative reads is a common practice in long read transcriptomics. One package I can recommend that has been used extensively for showing long reads comes from the Hagen Tilgner lab- ScisrWiz. But UCSC genome browser or IGV plots could also work.

Question 4 (Related to 3): Figure 5,E,F. By plotting PSI of a given exon (0-1 scale) for each different pA isoform, one does not get an appreciation of the expression levels of each isoform. Some indication of read depth should be provided in E,F.

Question 5: I might be buried somewhere, but can you bring to light the proportion of the novel isoforms and alternative exons that are "neuronal diff increased/decreased"? Line 204, Fig. 4A describes that 37% of isoforms undergoing DTU were novel. This will bring to clarity with regard to whether these novel forms are likely expressed in glutamatergic neurons vs undifferentiated cells. Perhaps this can be integrated with the formatting of "Down (D), Up (U), unchanged (-)" like in fig. 3?

Question 6: A strong trend for coordination between alternative splicing and APA events is uncovered (Fig. 5). Recent papers have also found evidence for coordination of alternative first exons (implying alternative transcription start sites) and 3' end usage (Calvo-Roitberg et al., Science, 2025). This analysis also used PacBio long reads. I'm curious to know if the trend for coordination (PITA) is confirmed in this new and very deep dataset of long reads. This is also an important aspect to consider given that in Drosophila, a study (Alfonso-Gonzalez et al., Cell, 2023) found that the coordination between alternative first exons and APA is more pervasive than between alternative splicing and APA.

Question 7: pPAU is used as a metric to compare APA with gene expression. (Fig. S11). This was the default metric for

QAPA in the original paper. However, others have used dPAU as a metric for QAPA output (Zhang et al., Nat. Comm, 2023). Given that distal polyA sites are preferentially used during neuronal differentiation, dPAU should be tested. This is important because the neuronal APA patterns often involve low level increases in very distal polyA sites. In Fig. S11 there does seem to be a significant very weak correlation (bottom right panel; $R = 0.047$, $P < 0.001$). Perhaps when using dPAU the trend will be stronger?

Question 8. I am concerned about the 4,332 putative gene fusion transcripts identified. It needs to be resolved whether these might be artifacts of the library preparation protocol, or alignment problems, versus Bonafide transcript fusions. I assume these are novel? Comparing to other datasets is important to resolve this. It would not even need to be exclusively long read RNA-seq data, because junction spanning reads from short read RNA-Seq data would also suffice as support. The read depth for the genes in this category is also important to explore. It would be odd that so many thousands of fusion transcripts have gone undetected in other studies. Some exploration of relevant literature is warranted.

Question 9. In Figure S12, the Y axis scale for "novel" and "known" categories are different. Because of this, it is harder to convince the reader that the splice donor/acceptors of the novel category are significant when it comes to MAPS.

Question 10. The novel statistical framework using generalized linear models (GLMs) for AS-APA coordination might represent a methodology advance. Perhaps this could be described in greater detail in the supplement.

Minor points:

Figure legend font and spacing problems in Supp Fig. 11.

(Remarks on code availability)

Reviewer #2

(Remarks to the Author)

The manuscript by Xu, Rynard and colleagues reports an analysis of alternative splicing other RNA isoform variation in differentiating iPSC-derived neurons using long read RNA sequencing (PacBio Kinnex). The study addresses an important knowledge gap in the field: the limitations of short read RNA-seq have prevented comprehensive characterization of isoform specific regulation of RNA expression in differentiating neurons, which could play an important role in neural regulation and neurodevelopmental disease. The data quality and analyses are generally very high quality, although I found some specific analyses to be hard to understand or underdeveloped. Overall I think this study will be an important contribution to the field.

One major caveat/limitation of this study is the uncertainty of generalizing findings from iPSC derived neurons, all generated from the same source donor, to specific neuron or glial cell types. In particular this paradigm does not capture the diversity of brain cell types, which have highly cell type specific expression and isoform regulation.

Specific comments:

- It appears that all data were generated from one cell line, and they used three independent iPSC stocks as replicates. It would be helpful if some validation of the findings in a slightly different context could be provided. This could potentially be done by additional bioinformatic analysis, e.g. comparing with other long-read RNA datasets from post-mortem human brain
- Fig. S2 - the "Duffy neuron" is provided as a comparison but the pattern of expression for neuronal genes looks quite different from the t30 neurons on this heatmap. Some more quantitative analysis of these data and a discussion of the reason for these differences would be helpful.
- Only 30% of the "novel" isoforms overlapped with a study from fetal cortex. What is the fraction of the fetal isoforms that are captured here? This overlap seems low, but the authors' interpretation that the biological context is quite different makes sense. However, I'm confused by the reference to Fig. S5A - I don't understand how Fig S5A supports this statement. Could you clarify?
- Line 131: Only 79% of the novel splice junctions were confirmed by short read sequencing. This is decent, but I wonder what accounts for the 21% which are not validated. Is this due to insufficient depth in the short read data? Or does it reflect a rate of false positives in the long read junction calling?
- Line 152: What was the rate of validation for novel peptides? The text only reports the number of protein isoforms.
- The schematic in Fig. 3A is not very helpful, it would be better to show heatmaps or actual trajectories.
- Fig 3 - Can you statistically test "discordance"? Maybe a likelihood ratio test?
- I found Fig. 4 very hard to interpret. In particular, in Fig 4B - I can't understand what this is showing at all.
- Text refers to wrong figure panels for Figs. 4B,C,D
- Line 213 - Is this shown in Fig. 4C? The relationship between these results and Fig. 4 is not clear.
- The AGO1 finding is interesting. Could this be confirmed/validated functionally? What are the predictions for loss of AgoN domain?
- It would be nice to see some IGV screenshots showing examples of the structure of reads, e.g. for polyA analysis
- Fig 5 - what time point was used for Fig. 5A,B?
- It would be helpful to see some kind of histogram or volcano plot showing the magnitudes of the effects, rather than just reporting the number of events.
- Fig. 5C,D - It's not clear how to interpret these distributions. Can you compare with some kind of shuffled null distribution? Are the distances and # of exons different from what you might expect under some kind of random model?

(Remarks on code availability)

Reviewer #3

(Remarks to the Author)

Tripathy and colleagues have produced a deeply sequenced long-read RNA-seq atlas of cortical neuron development, which they combine with proteomics to examine changes in splicing and polyadenylation.

Overall the paper seems methodologically sound and they perform quality control of their isoform set.

Isoform discovery in long-read RNA-seq is still in its infancy so I would appreciate if the authors compared their isoform catalogue from Pigeon to a few other tools (Stringtie, Bambu, IsoQuant) to check robustness, particularly for their novel isoforms. The assumption would be that the overlap between tools is low - not necessarily a bad thing as this is a hard problem, but important to test as I don't believe Pigeon has been benchmarked against these other tools by LRGASP or other papers.

I would compare the ORF estimates from ORFAnage to a few other tools (GeneMarkS, CPAT) to assess the robustness of those predictions.

The shiny app browser they provide is very helpful but could benefit from a few additional features. Specifically, the bar plots should show the variance between the replicates (boxplots, mean +- SEM etc) and it would be helpful to include some way of showing which isoforms are predicted to be protein-coding and have peptide support.

(Remarks on code availability)

Appears reproducible, includes dependencies as a conda environment.

Version 1:

Reviewer comments:

Reviewer #1

(Remarks to the Author)

I'm satisfied with the response by the authors.

(Remarks on code availability)

Reviewer #2

(Remarks to the Author)

The authors have thoroughly responded to my previous comments and those of the other reviewers. I have some remaining small suggestions (below), but I believe the manuscript is greatly improved.

- The authors added Fig. S3 showing pileups of the reads at the Ago1 locus. However, it would be very helpful to show IGV browser shots that display (a downsampled subset of) individual reads, i.e. showing the actual structure of some example long reads. I believe this was the intent of Reviewer 1 question 3, and I agree with that reviewer that such a representation would be very valuable. I don't understand the comment: "Displaying reads post-filtering does not provide any variability amongst the reads."

- The Shiny browser also does not show individual reads. An IGV browser showing BAM files would be most helpful.

- In the new version of Fig 3A, there is no explanation in the legend for the Venn diagrams - Please add this to the legend.

(Remarks on code availability)

Reviewer #3

(Remarks to the Author)

I thank for the authors for their hard work in making their results more robust. I particularly appreciate the improvements to the Shiny app as it's now much clearer.

I am satisfied that the study is now suitable for publication.

(Remarks on code availability)

README seems extensive with instructions for installation and running.

Open Access This Peer Review File 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 license, and indicate if changes were made.

In cases where reviewers are anonymous, credit should be given to 'Anonymous Referee' and the source.

The images or other third party material in this Peer Review File are included in the article's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons license 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 license, visit <https://creativecommons.org/licenses/by/4.0/>

Executive Summary of Manuscript Changes

We sincerely thank all three reviewers for their detailed and constructive feedback. Guided by their comments, we have performed substantial new analyses, expanded external validation, refined our methodology, and clarified the text and figures throughout. A fully revised manuscript has now been submitted. Below we summarise the principal revisions.

First, Reviewers 1, 2, and 3 all asked us to contextualise our novel isoform catalog against external long-read transcriptomic resources and to better support generalizability beyond our iPSC-derived neuron system.

- We now compare our isoform catalog against the ENCODE4 long-read RNA-seq brain collection (Reese et al., 2023) and a long-read atlas of human fetal cortex (Patowary et al.) at line 128 and Figure S6.

Second, Reviewer 3 asked us to benchmark isoform and ORF discovery against alternative tools.

- We now compare our Iso-Seq isoform set against Bambu (v3.0.5) at line 707 and Figure S19A, B. In parallel, we benchmark ORFanage ORF predictions against CPAT and GeneMarkS-T, finding 94% overlap between ORFanage and CPAT (line 747 and Figure S20).

Third, Reviewer 1 asked us to address the authenticity of the 4,332 putative gene fusion transcripts.

- We now clarify that all fusion transcripts arise from adjacent, same-strand genes, making reverse-transcription/PCR artifacts unlikely, at line 132.
- Additionally, 66.0% of fusion transcript splice junctions are independently confirmed by short-read RNA-seq, and 10.4% are also detected in the Patowary et al. long-read embryonic dataset, supporting their authenticity (Figure S7E, F).
- Fourth, Reviewer 1 asked whether positional initiation-termination axis (PITA) coupling, recently described by Calvo-Roitberg et al. (Science, 2025), is recapitulated in our deep long-read dataset.
- We now perform a full PITA analysis, computing per-gene Spearman's ρ and genome-wide Δ AUC metrics between alternative TSSs and PASs, confirming robust PITA coupling across neurodevelopment, at line 300 and Figure S14.
- We additionally extended our GLM-based AS-APA coordination framework to alternative transcription start site and alternative splicing (ATSS-AS) events, identifying 2,467 significant ATSS-AS coordination events across the time course (line 324 and new Figures 6, 7).

Fifth, Reviewers 1 and 2 asked us to improve visualisation and quantification across several figures.

- We have added representative unfiltered long reads for AGO1 across timepoints (Figure S3)

- Added per-timepoint read counts to the pA and TSS usage panels (Figures 6–7E, F).
- Added log2FC/dIF directionality plots broken down by isoform category with accompanying Fisher's exact tests (Figure 4B, Table S5).
- Added effect-size histograms and shuffled null distributions for coordination analyses (genomic distance, exon number; Figures 6–7C, D, Figure S15).
- Figure 3A has been redrawn to show mean $\pm$ SD expression trajectories with Venn diagrams of gene-set overlap.

Sixth, Reviewer 2 asked whether mRNA–protein discordance could be formally tested rather than just described.

- We now report a permutation test ($n = 10,000$; $P = 0$) showing observed concordance of $\sim 45.5\%$ vs. a permuted mean of $\sim 27.1\%$, indicating statistically significant coupling yet also leaving 54.5% of genes (4,630 of 8,498) genuinely discordant (line 208).

Seventh, Reviewer 2 asked for a more quantitative comparison between our t30 neurons and the Duffy et al. reference neurons.

- We now include a PCA and correlation plot in Figure S2 (TGF- β inhibitor supplementation, differing doxycycline induction schedule).

Overall, the revised manuscript now provides a more rigorously benchmarked, externally validated, and quantitatively interpreted long-read transcriptomic atlas of human iPSC-derived cortical neuron differentiation, directly addressing every major and minor point raised by the three reviewers.

Reviewer #1 (Remarks to the Author):

This manuscript by Xu et al represents an important advance in human neuronal transcriptomics on several fronts. It provides a deep dataset of long read RNA-Sequencing during human iPSC neuronal differentiation, uncovering thousands of novel transcript isoforms. These isoforms are investigated with relation to ASD risk and mutational constraint. Widespread coordination between alternative splicing events and alternative 3'UTR selection is uncovered. Support for some novel isoforms is demonstrated using proteomics. The web application for viewing full length isoforms and abundances is very useful.

We thank this reviewer for their very positive reception of our paper and their thoughtful comments below.

I have the following questions that should be addressed prior to publication.

Major points:

Question 1: Line 83. The reader might not know about the Kinnex library prep details and thus miss the fact that polyA containing mRNA is sequenced (not total RNA).

We have now revised the Methods section to clarify that, although total RNA was used as the input material, the Iso-Seq Express 2.0 Kit employs oligo-dT priming for cDNA synthesis, which selectively captures polyadenylated mRNA (line 639). This means our dataset represents the polyA+ transcriptome rather than total RNA.

Question 2:

- 689 million long reads (line 95) does seem very deep for PacBio data. But the reader has to take the authors' word that it is one of the most deeply sequenced to date. How do other recent studies compare? There are long read datasets from the Encode project, for example: Reese et al., bioRxiv, "The ENCODE4 long-read RNA-seq collection reveals distinct classes of transcript structure diversity".

Our paper utilized PacBio Kinnex, which is a cDNA concatenation approach that increases read yield on average by 8-fold relative to previous protocols. Only another two published benchmark papers have utilized PacBio Kinnex for bulk RNA-seq so far (Wissel et al., 2025 and You et al., 2025). The ENCODE4 long-read RNA-seq collection did not use Kinnex.

- A comparison of novel isoform detection across a few other recent datasets might help support the annotation of the novel transcripts and even give insights into cell type expression patterns.

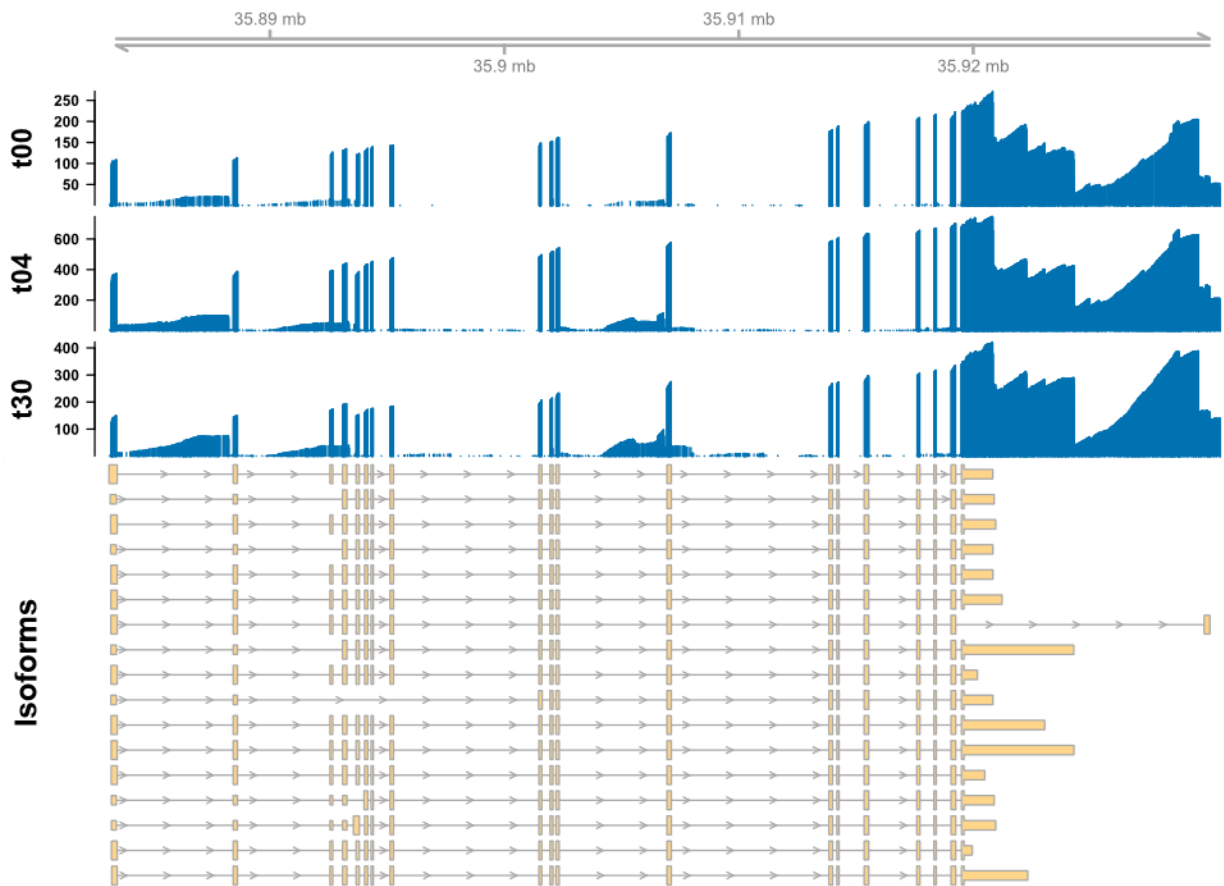
We thank the reviewer for this suggestion and have performed a bioinformatic comparison of our novel isoform catalog against the ENCODE4 long-read RNA-seq brain dataset (Reese et al., 2023). In addition, we also systematically compared our novel isoforms with a fetal cortex long-read dataset, showing considerable overlap between isoforms between datasets (30%; Fig. S6).

We have added this analysis as a supplementary figure and expanded our Discussion to address comparison with additional datasets (line 501).

Question 3: There are no plots shown with actual reads. Showing representative reads is a common practice in long read transcriptomics. One package I can recommend that has been used extensively for showing long reads comes from the Hagen Tilgner lab- ScisorWiz. But UCSC genome browser or IGV plots could also work.

We appreciate the suggestion to display actual reads. We have added total (unfiltered) reads for AGO1, per timepoint (Figure S3, reproduced below). It should be noted that these are all reads from the AGO1 genomic region, pre-filtering. Displaying reads post-filtering does not provide any variability amongst the reads, and as such, we have collapsed these reads into isoforms.

AGO1



Question 4 (Related to 3): Figure 5,E.F. By plotting PSI of a given exon (0-1 scale) for each different pA isoform, one does not get an appreciation of the expression levels of each isoform. Some indication of read depth should be provided in E,F.

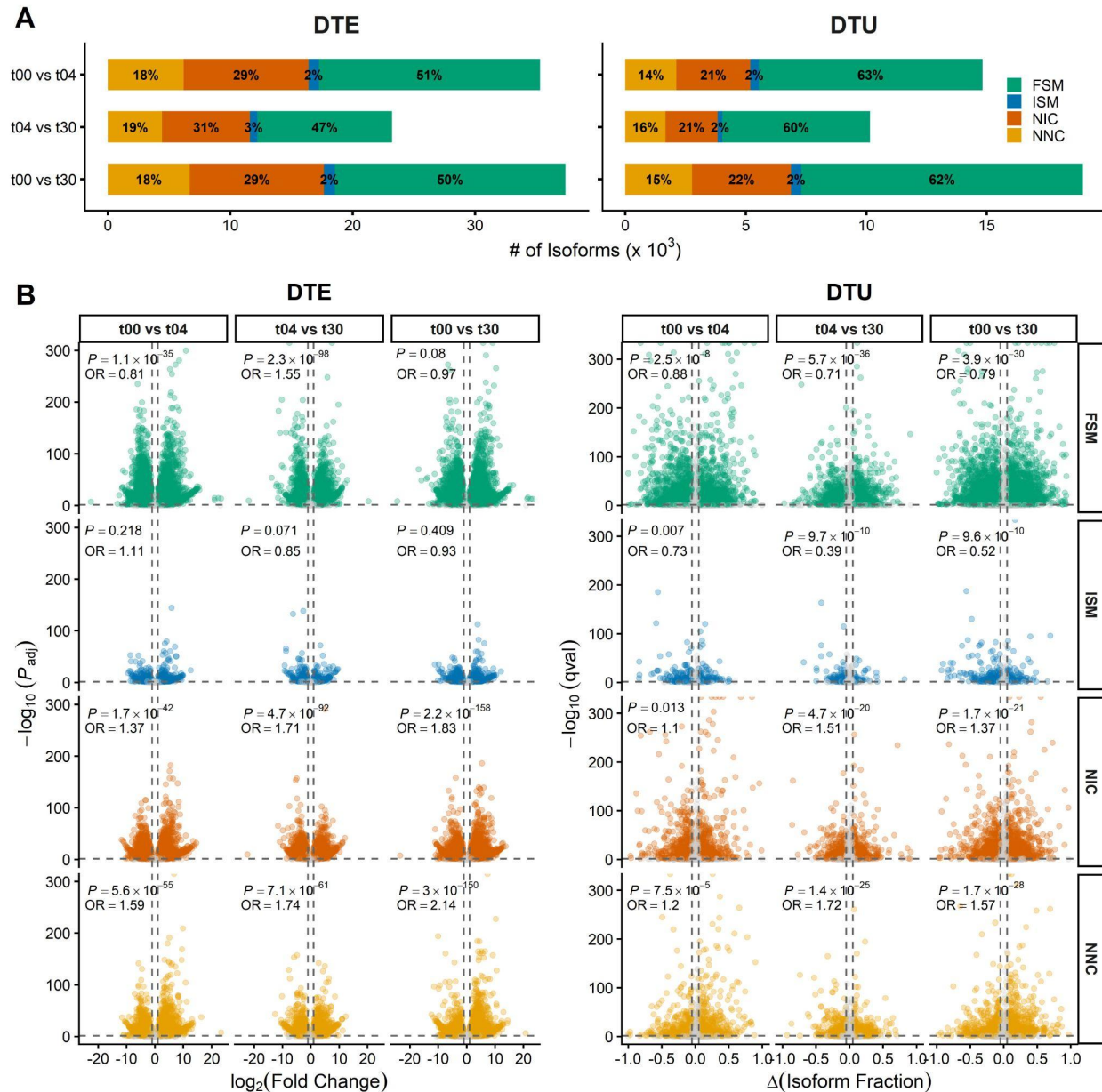
We appreciate the suggestion to display expression levels of each pA site. To address this, we have added read counts for each pA site (and each TSS in the new ATSS-AS analysis), at each timepoint in Figure 6-7E & F.

Question 5: It might be buried somewhere, but can you bring to light the proportion of the novel isoforms and alternative exons that are “neuronal diff increased/decreased”? Line 204, Fig. 4A describes that 37% of isoforms undergoing DTU were novel. This will bring to clarity with regard to whether these novel forms are likely expressed in glutamatergic neurons vs undifferentiated cells. Perhaps this can be integrated with the formatting of “Down (D), Up (U), unchanged (-)” like in fig. 3?

We appreciate the feedback to emphasize the proportion of novel isoforms increasing/decreasing across timepoint comparisons. To address this, we have added

plots displaying log2FC or dIF, per isoform category, in Figure 4 (reproduced below; lines 223-232). Additionally, we performed a Fisher's exact test to compare the number of isoforms increasing/decreasing per category to provide an indication of the magnitude and numbers of these changes. We found that FSM isoforms undergoing differential expression showed significantly more isoforms decreasing in expression from t00 to t04, with the opposite from t04 to t30 (Fig. 4B). ISM isoforms undergoing differential expression were evenly distributed between increasing and decreasing across the timecourse (Fig. 4B). NIC and NNC transcripts undergoing differential expression, at each timepoint comparison, had significantly more isoforms increasing across the timecourse, rather than decreasing (Fig. 4B). FSM and ISM isoforms undergoing differential usage, showed significantly more isoforms decreasing in usage at all timepoint comparisons (Fig. 4B). NIC and NNC transcripts undergoing differential usage showed significantly more isoforms increasing in usage at each timepoint comparison (Fig. 4B).

The data used for these plots can additionally be found in Table S5.

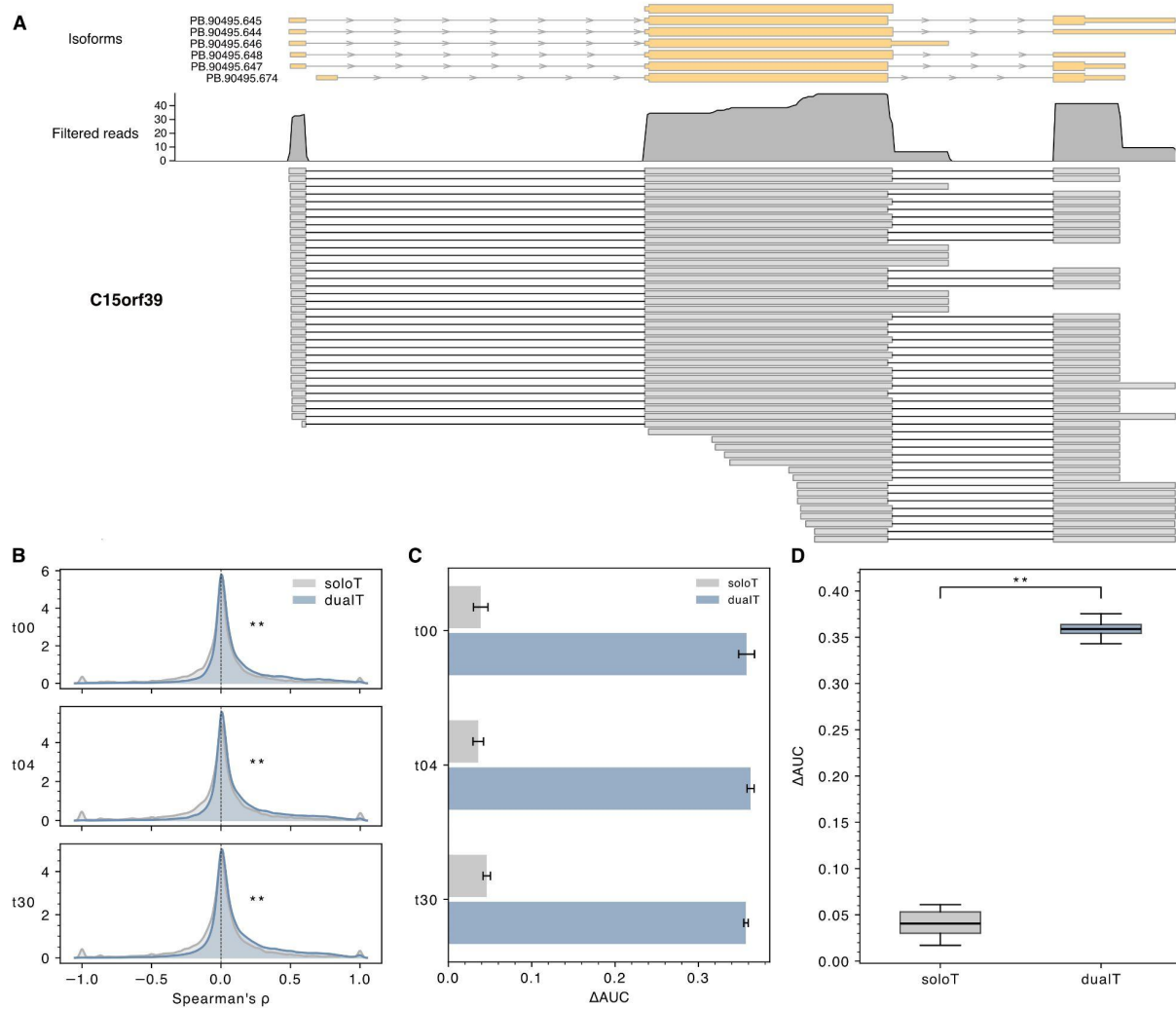


Question 6: A strong trend for coordination between alternative splicing and APA events is uncovered (Fig. 5). Recent papers have also found evidence for coordination of alternative first exons (implying alternative transcription start sites) and 3'end usage (Calvo-Roitberg et al., Science, 2025). This analysis also used PacBio long reads. I'm curious to know if the trend for coordination (PITA) is confirmed in this new and very deep dataset of long reads. This is also an important aspect to consider given that in *Drosophila*, a study (Alfonso-Gonzalez et al., Cell, 2023) found that the coordination between alternative first exons and APA is more pervasive than between alternative splicing and APA.

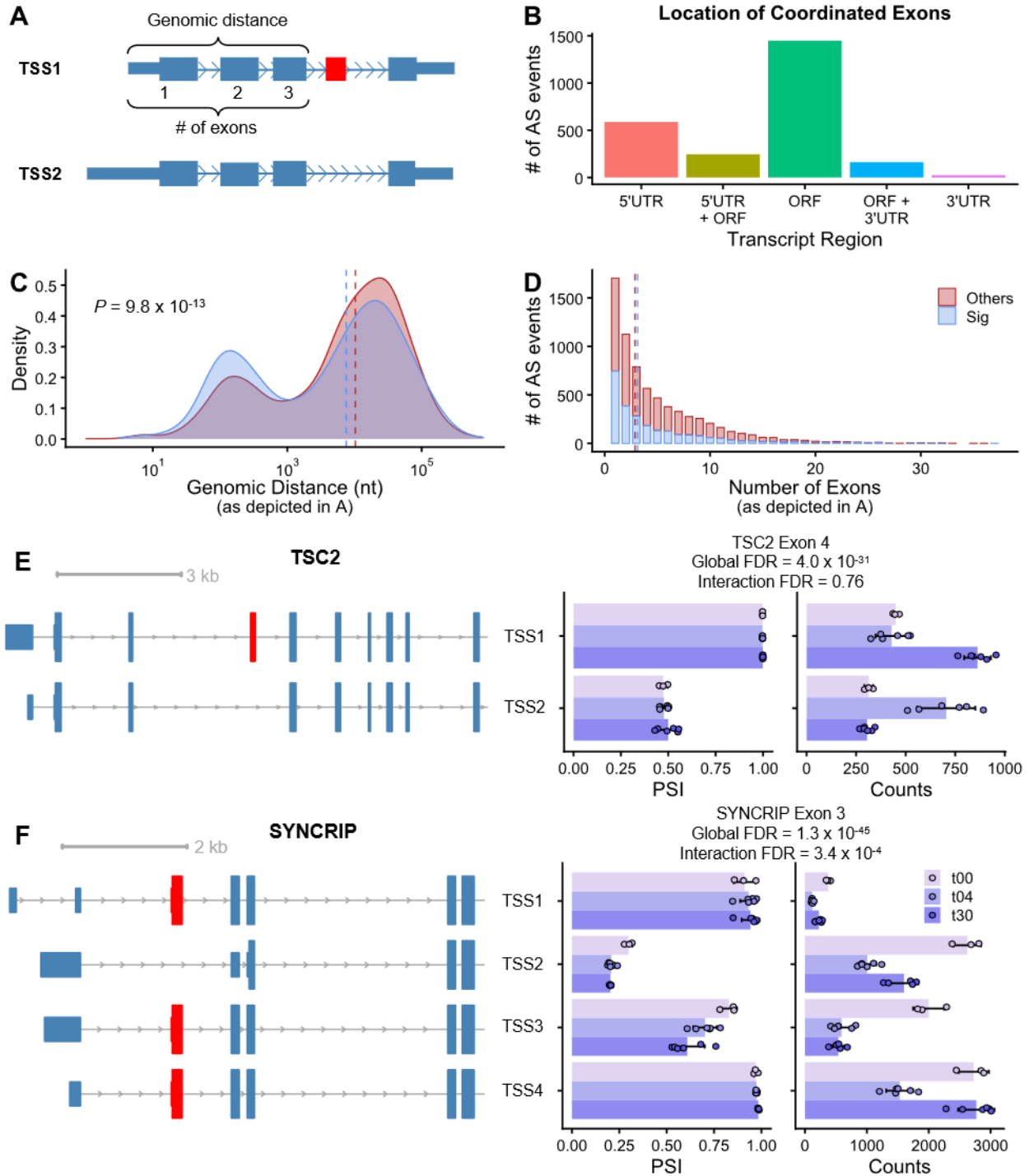
We thank the reviewer for suggesting mRNA initiation and termination coordination. To address this comment, we performed a positional initiation-termination axis (PITA) analysis following the framework recently described by Calvo-Roitberg et al., which quantifies the ordinal coupling between alternative transcription start sites (TSSs) and polyadenylation sites (PASs) within individual genes. Briefly, for each gene we calculated a Spearman's ρ between the rank coordinates of TSSs and PASs across all reads within a sample, where positive ρ values indicate positional coupling (PITA) — i.e., upstream TSSs tend to pair with upstream PASs and downstream TSSs with downstream PASs. We also calculated a Δ AUC metric (area under the Spearman's ρ distribution shifted toward positive values minus that shifted toward negative values) to quantify genome-wide enrichment of PITA-coupled genes. We then compared these metrics between dual alternative termini (dualT) genes — genes expressing at least two alternative first exons and at least two alternative polyadenylation sites — and single termini (singleT) genes — genes using only one first exon or one PAS per sample — across all time points in our dataset.

This analysis revealed that, at all time points examined in our study, both Spearman's ρ values and Δ AUC values were significantly higher for dualT genes compared to singleT genes (K-S test, $p < 10^{-16}$ at all time points; Figure S14). This finding is consistent with the results reported by Calvo-Roitberg et al. across multiple human tissues, and demonstrates that PITA coupling — the tendency for transcripts to co-select TSSs and PASs of similar ordinal position — is a robust feature of our dataset across the neurodevelopmental time course. Importantly, singleT genes, which serve as a control for potential biases in terminal coordinates from long-read RNA-seq, showed near-zero Δ AUC values at all time points, confirming that the enrichment of positive Spearman's ρ values in dualT genes reflects genuine positional coupling rather than a technical artifact.

We have added a description of this analysis and its results to the manuscript (see the new paragraph in the Results section below, and Figure S14).



Additionally, to address the possibility of coordination between ATSSs and AS, we performed the same GLM-based AS-APA analysis using ATSSs. We identified 2,467 unique significant ATSS-AS coordination events across all three timepoints (Fig. 6).

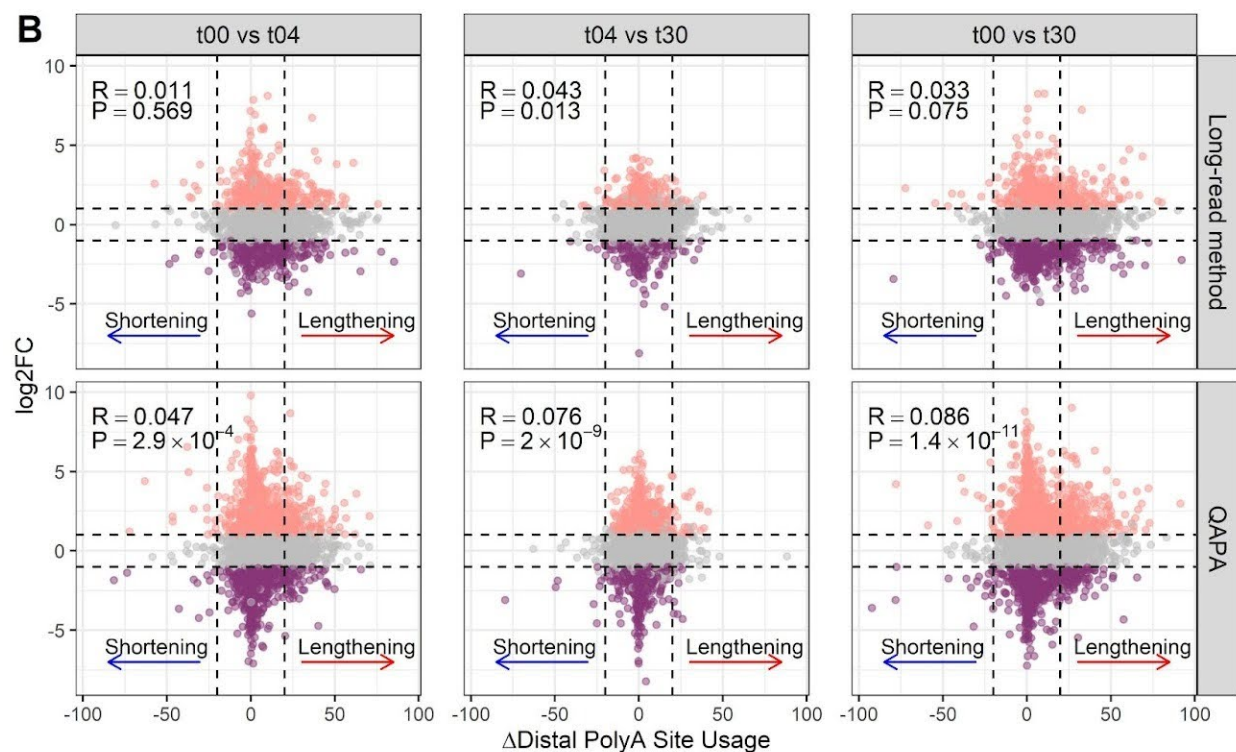


Each of these pairwise analyses (i.e. ATSS-APA, ATSS-AS and AS-APA) were performed with different transcripts, as well as different reads. Thus, we did not directly compare overlapping events/genes.

Question 7: pPAU is used as a metric to compare APA with gene expression. (Fig. S11). This was the default metric for QAPA in the original paper. However, others have used dPAU as a

metric for QAPA output (Zhang et al., Nat. Comm, 2023). Given that distal polyA sites are preferentially used during neuronal differentiation, dPAU should be tested. This is important because the neuronal APA patterns often involve low level increases in very distal polyA sites. In Fig. S11 there does seem to be a significant very weak correlation (bottom right panel; $R = 0.047$, $P < 0.001$). Perhaps when using dPAU the trend will be stronger?

We appreciate the suggestion to use distal PAU, rather than proximal PAU, given that distal polyA sites are preferentially used during neuronal differentiation. To address this, we have re-analyzed PAU to consider distal polyA site usage (Figure S16B; lines 381-382). We find a very marginal increase in the correlation, suggesting that our prior inference is unlikely to be meaningfully impacted by distal vs proximal PAU sites.

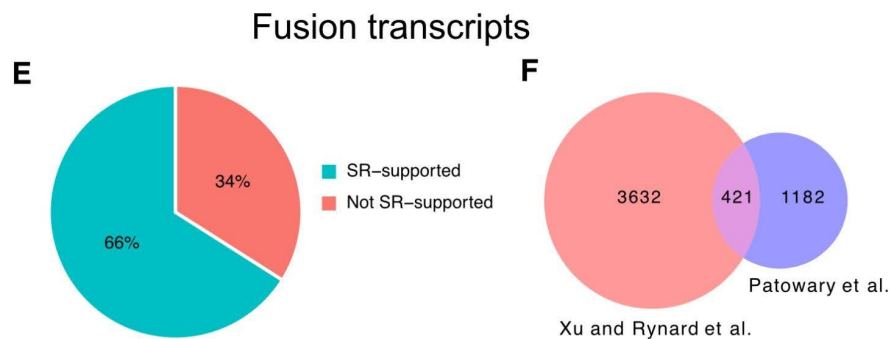


Question 8. I am concerned about the 4,332 putative gene fusion transcripts identified. It needs to be resolved whether these might be artifacts of the library preparation protocol, or alignment problems, versus Bonafide transcript fusions. I assume these are novel? Comparing to other datasets is important to resolve this. It would not even need to be exclusively long read RNA-seq data, because junction spanning reads from short read RNA-Seq data would also suffice as support. The read depth for the genes in this category is also important to explore. It would be odd that so many thousands of fusion transcripts have gone undetected in other studies. Some exploration of relevant literature is warranted.

We thank the reviewer for this suggestion. We have now added additional explanations for why we feel that these fusion transcripts are not likely to be entirely artifacts.

First, we have additionally clarified that all fusion transcripts come from genes that are adjacent to one another in the genome and are encoded on the same strand, making it unlikely that they represent artifacts of reverse transcription or PCR (line 132).

Second, we have additionally shown that 66.0% of all fusion transcripts have their splice junctions independently confirmed by short-read RNA-seq data, and 10.4% are also detected in an external long-read RNA-seq dataset from human embryonic tissues, Patowary et al., further supporting for the authenticity of these fusion transcripts (Fig. S7E, F, reproduced below).



Question 9. In Figure S12, the Y axis scale for “novel” and “known” categories are different. Because of this, it is harder to convince the reader that the splice donor/acceptors of the novel category are significant when it comes to MAPS.

We have adjusted the Y axis scales for “novel” and “known” categories in Figure. S12 to be the same.

Question 10. The novel statistical framework using generalized linear models (GLMs) for AS-APA coordination might represent a methodology advance. Perhaps this could be described in greater detail in the supplement.

We have added more detail to the Methods section on our quasi-binomial models (see Methods - ‘Identification and Characterization of Alternative Transcript Processing Event Coordination’; lines 933-1001). Additionally, we have added a vignette illustrating this method to our Github repository.

Minor points:

Figure legend font and spacing problems in Supp Fig. 11.

We thank the reviewer for raising this issue. Figure legend font and spacing problems have been adjusted accordingly (new Figure S15).

Reviewer #2 (Remarks to the Author):

The manuscript by Xu, Rynard and colleagues reports an analysis of alternative splicing other RNA isoform variation in differentiating iPSC-derived neurons using long read RNA sequencing (PacBio Kinnex). The study addresses an important knowledge gap in the field: the limitations of short read RNA-seq have prevented comprehensive characterization of isoform specific regulation of RNA expression in differentiating neurons, which could play an important role in neural regulation and neurodevelopmental disease. The data quality and analyses are generally very high quality, although I found some specific analyses to be hard to understand or underdeveloped. Overall I think this study will be an important contribution to the field.

We thank this reviewer for their overall very positive reception of this work.

One major caveat/limitation of this study is the uncertainty of generalizing findings from iPSC derived neurons, all generated from the same source donor, to specific neuron or glial cell types. In particular this paradigm does not capture the diversity of brain cell types, which have highly cell type specific expression and isoform regulation.

We appreciate this caveat and have taken steps to better generalize our findings to those collected from a greater diversity of cells and cell types (detailed in specific comments below).

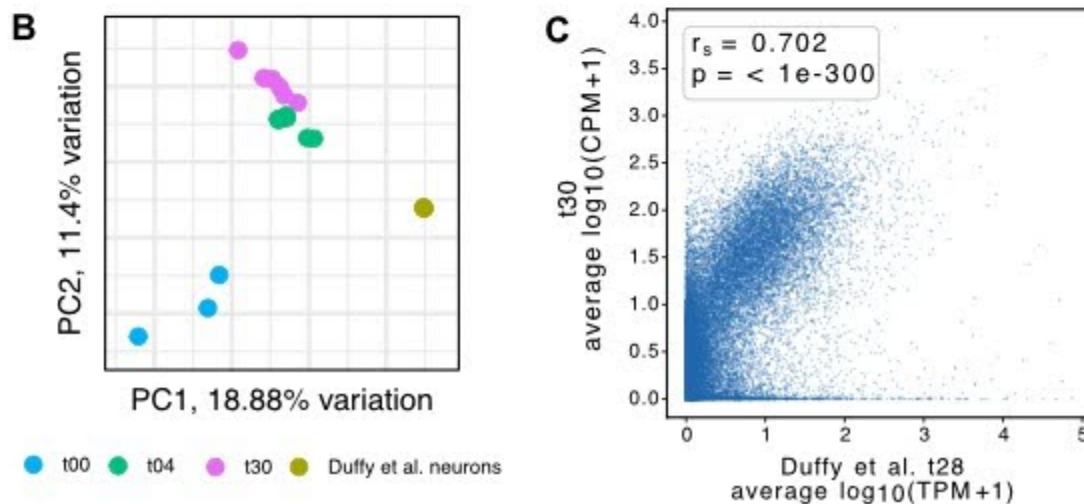
In addition, we have more carefully highlighted this specific limitation in the Discussion (line 498), detailed below:

We acknowledge several limitations. First, our atlas was generated using cells from a single male donor differentiated through an NGN2-induced excitatory lineage, and this iPSC-derived model does not capture the full cellular diversity, 3D architecture, or signaling environment of the in vivo human brain. Reassuringly, 30.8% of the novel isoforms present in our dataset are also detected in a recent long-read atlas of human fetal cortex (Patowary et al.), and the evidence for their functional relevance draws on population-scale analyses spanning hundreds of thousands of individuals. Despite this, larger long-read studies across diverse donors and ancestries remain needed to establish the generalizability of specific isoforms, particularly those expressed at low levels.

Specific comments:

- Fig. S2 - the "Duffy neuron" is provided as a comparison but the pattern of expression for neuronal genes looks quite different from the t30 neurons on this heatmap. Some more quantitative analysis of these data and a discussion of the reason for these differences would be helpful.

We appreciate this comment and have now more directly quantified the degree of similarity between the timepoints represented in our dataset and those from the “Duffy neuron” dataset (Duffy et al.), where mRNA from cells were harvested at day 28. Specifically, we have now added a PCA plot and a correlation plot to Fig. S2B and C (reproduced below) to quantify the degree of similarity between our dataset and the Duffy dataset.



We reason that some of the differences between the datasets could be explained by subtle differences in the differentiation protocols used. These differences include 1) Duffy et al. added additional small molecules to supplement their media, such as SB431542, a TGF-beta inhibitor, and 2) Duffy et al. included doxycycline in their media on days 1-4, 8, 15 and 22, and we included doxycycline on days 1-7. Both including a TGF-beta inhibitor, and changing doxycycline induction length have been shown to alter the final neuronal state^{18,61}.

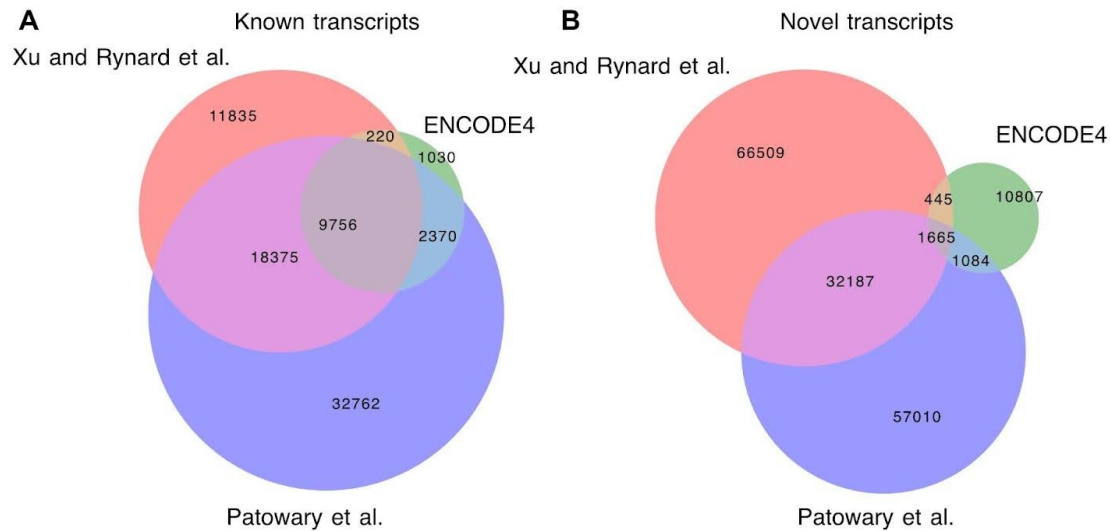
- It appears that all data were generated from one cell line, and they used three independent iPSC stocks as replicates. It would be helpful if some validation of the findings in a slightly different context could be provided. [Further comparison] could potentially be done by additional bioinformatic analysis, e.g. comparing with other long-read RNA datasets from post-mortem human brain
- Only 30% of the "novel" isoforms overlapped with a study from the fetal cortex. What is the fraction of the fetal isoforms that are captured here? This overlap seems low, but the authors' interpretation that the biological context is quite different makes sense.

To clarify, 17.1% (34,871) of the isoforms detected in the study from fetal human neocortex (Patowary et al.) overlap with “novel isoforms” in our study.

In addition, we now also perform a comparison of isoforms detected in our atlas to those from the ENCODE4 long-read RNA-seq adult post-mortem brain dataset (Reese et al.,

2023). We find that 2.1% of our novel isoforms are detected in ENCODE4. We reason that the lower detection rate of novel isoforms in this dataset compared to the fetal dataset is likely due to a large degree of developmental stage-related differences between our fetal-like NGN2 derived neurons at t30 and adult post-mortem brain tissues.

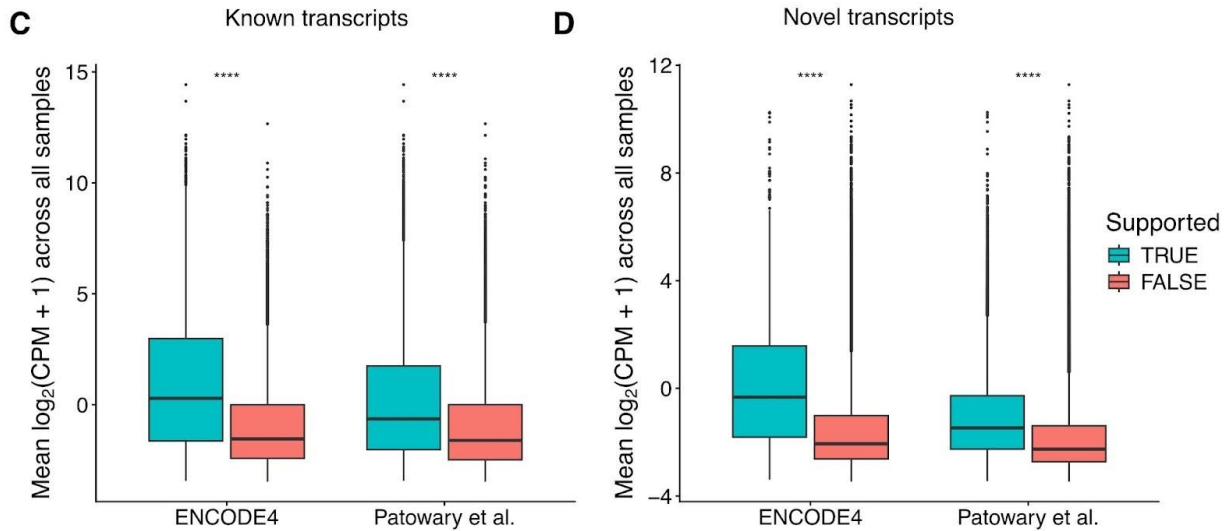
We have now added two venn diagrams to Fig S6 (A and B) that show the degree of overlap between known and novel isoforms from our study (Xu and Rynard et al.) with those from the fetal brain dataset (Patowary et al) and adult post-mortem brain dataset (ENCODE 4).



However, I'm confused by the reference to Fig. S5A - I don't understand how Fig S5A supports this statement. Could you clarify?

We have updated the prior Fig. S5A (now Fig. S6C and D, see image below) to more clearly illustrate that known and novel isoforms present in our study that are replicated in external studies (Patowary and ENCODE 4) tend to be more highly expressed than isoforms that are not detected in these additional studies.

This further helps support the idea that low overlap between our studies with prior studies is in part caused by the higher read depth in our dataset, which enables detection of low-abundance isoforms that fall below the detection threshold of external, more shallowly sequenced datasets.



- Line 131: Only 79% of the novel splice junctions were confirmed by short read sequencing. This is decent, but I wonder what accounts for the 21% which are not validated. Is this due to insufficient depth in the short read data? Or does it reflect a rate of false positives in the long read junction calling?

This is an astute observation and the reality is that we are not completely sure what accounts for why there are splice junctions that are observed in long-read but not short-read RNAseq, given that our short-read RNAseq is should be sequenced sufficiently deeply to identify such splice junctions (100M reads per sample).

One (potentially overly simplistic) explanation might be due to technical differences, for example, that rRNA depletion was used for short-read RNA-seq and oligo(dT) enrichment for long-read RNA-seq. To incorporate this explanation, we now state at line 144, “The remaining 21.1% of junctions may not have been detected due to technical differences between rRNA depletion (for short-read RNA-seq) and oligo(dT) enrichment (long-read RNA-seq).”

Beyond this somewhat simplistic explanation, a more detailed investigation into this discrepancy is likely necessary but is beyond the scope of this study.

- Line 152: What was the rate of validation for novel peptides? The text only reports the number of protein isoforms.

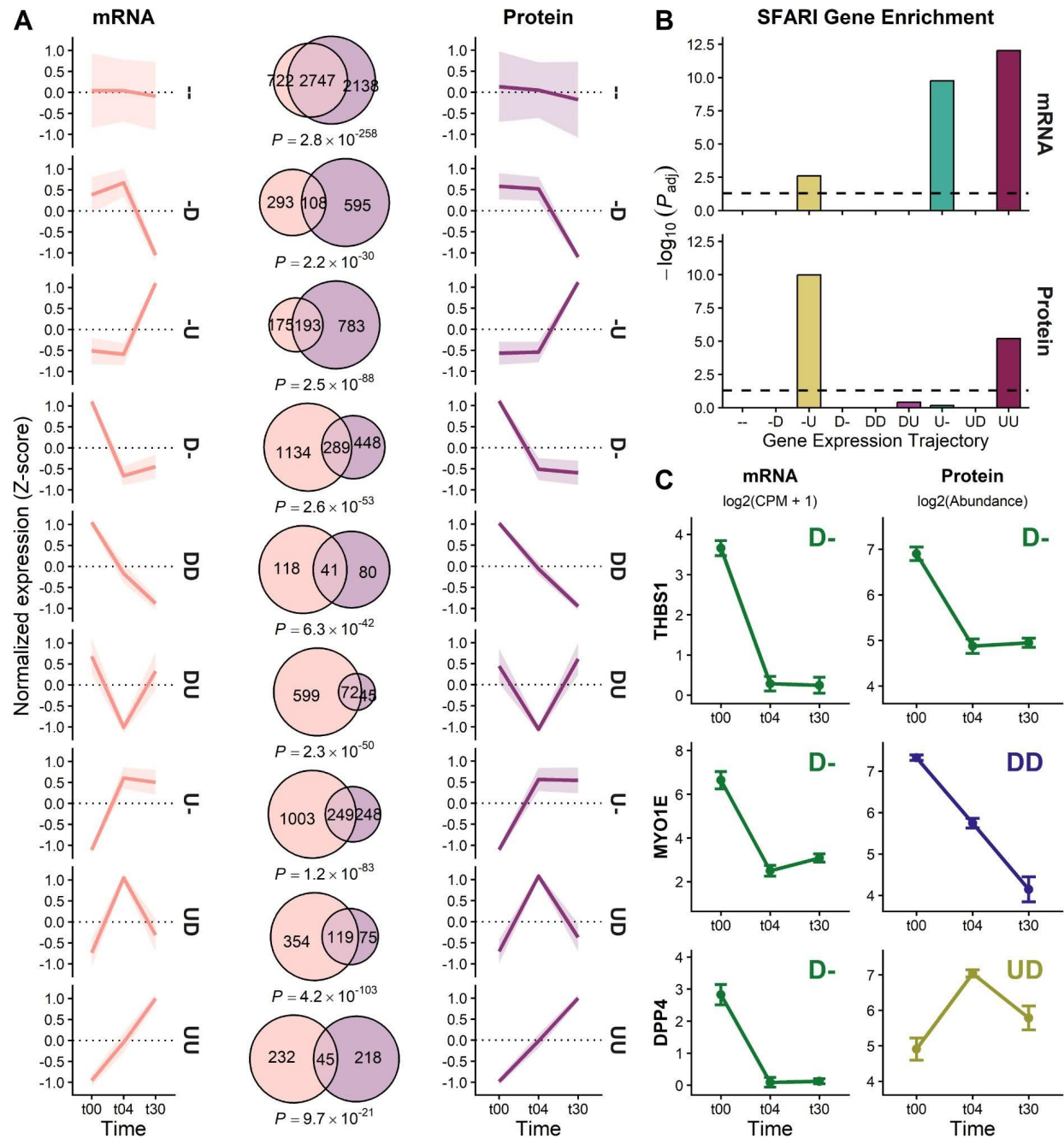
Calculating a validation rate for novel peptides requires knowing the total number of novel peptides theoretically detectable by our pipeline — i.e., all peptides derived from in silico digestion of our enhanced reference proteome that fall within the detectable range of our instrument. While we used an enhanced reference proteome containing novel protein sequences as our search database in Comet (the bioinformatics tool commonly

used in such proteomics analyses), Comet does not provide a list of the candidate peptides generated during in silico digestion.

Without this list, we cannot define a denominator and therefore cannot compute a meaningful validation rate. We instead report the absolute number of novel peptides identified and validated and note that this is how such data are typically reported by others previously (Humphrey et al., 2025) at line 168.

- The schematic in Fig. 3A is not very helpful, it would be better to show heatmaps or actual trajectories.

We appreciate the reviewer's thoughtful feedback. The updated Figure 3 has been comprehensively modified, and now shows mean and standard deviation expression values for each expression trajectory. Additionally, to clarify that each mRNA/protein expression trajectory may contain different genes, we have also added venn diagrams displaying the significant overlap (Fisher's exact test) between the sets of genes in each trajectory. As expected, this suggests that mRNA and protein expression are highly concordant (lines 208-210).

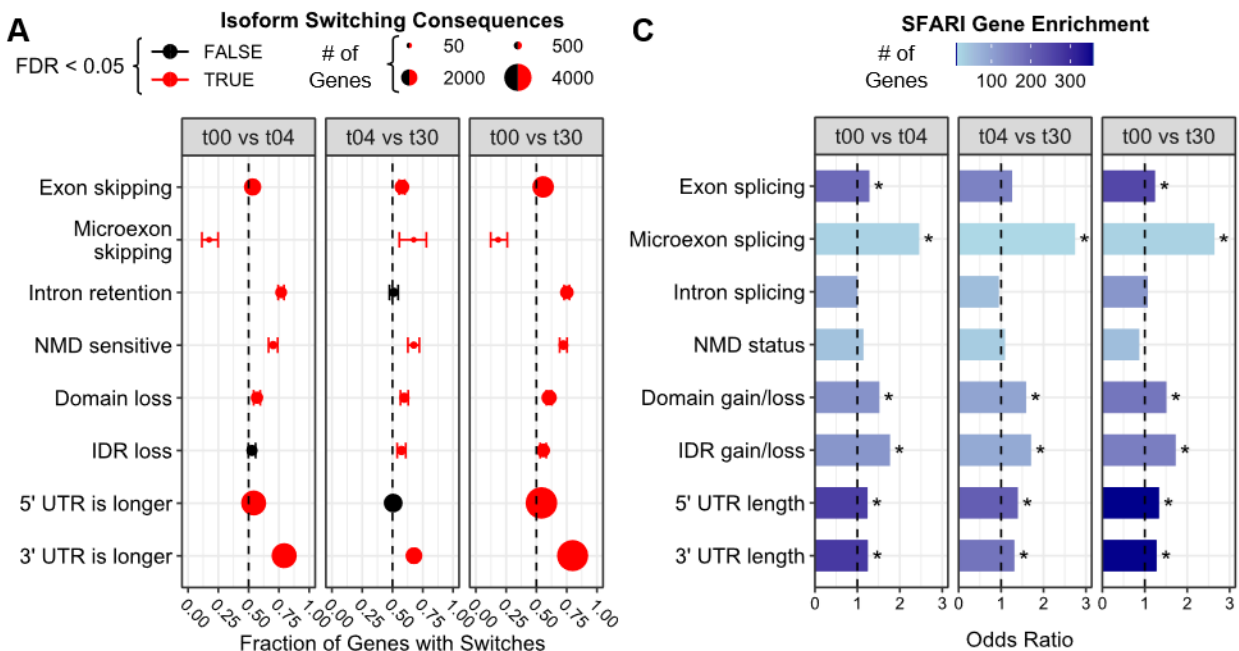


- Fig 3 - Can you statistically test “discordance”? Maybe a likelihood ratio test?

We thank the reviewer for this helpful suggestion. We performed a permutation test ($n = 10,000$ permutations; $P = 0$), where mRNA and protein expression are dependent. Our concordance was ~45.5%, and the mean concordance across permutation was ~27.1%. We would not expect this statistical test to assess mRNA and protein expression as entirely independent. Thus, despite this statistically significant concordance, we still find a large number of genes (4630 out of 8493; 54.5%) of genes displaying discordance (lines 208-210).

- I found Fig. 4 very hard to interpret. In particular, in Fig 4B - I can't understand what this is showing at all.
- Text refers to wrong figure panels for Figs. 4B,C,D
- Line 213 - Is this shown in Fig. 4C? The relationship between these results and Fig. 4 is not clear.

We thank the reviewer for this valuable input. We have added additional text in the results to describe Figure 4B (now Figure 5A; see 'Widespread isoform switching reveals dynamic regulation of splicing during neuronal differentiation'; lines 233-243). In addition, we have now increased the size of this figure (in part by breaking the figure up into two figures, reproduced below), which should help better illustrate what this figure is supposed to show. We have also now fixed figure references in the text.



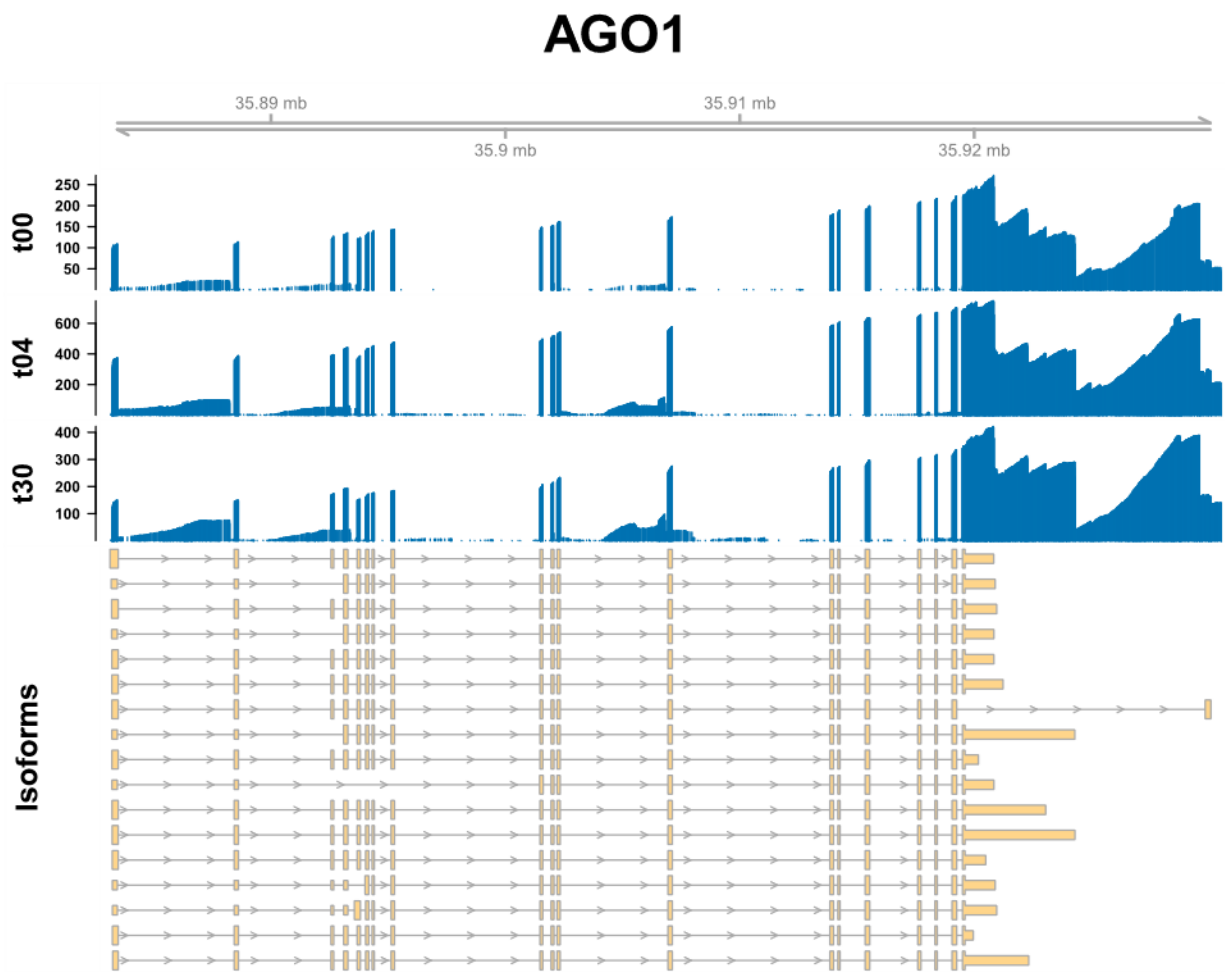
- The AGO1 finding is interesting. Could this be confirmed/validated functionally? What are the predictions for loss of AgoN domain?

We appreciate the suggestions to confirm/validate the novel AGO1 isoform and ORF, however this is outside of the scope of this work.

Thus, we have further elaborated on possible interpretations of the loss of AGO1 N-terminal domain on AGO1 function. We state, "The implications of this isoform switch for AGO1 protein function are unclear. An AGO1 protein lacking the N-terminal domain may be less efficient in generating AGO1-miRNA complexes and/or could impact miRNA maturation. Direct experimental validation of the novel start codon usage, and presence and function of the truncated protein is required" (lines 284-291).

- It would be nice to see some IGV screenshots showing examples of the structure of reads, e.g. for polyA analysis

We appreciate the suggestion to display actual reads. We have added total (unfiltered) reads for AGO1, an ASD risk gene (Figure S3). However, it should be noted that these are all reads from the AGO1 genomic region, pre-filtering. Displaying reads post-filtering does not provide any variability amongst the reads, as they have been collapsed into isoforms. Thus, we display these reads as a supplemental figure associated with Figure 1, rather than with the polyA analysis.

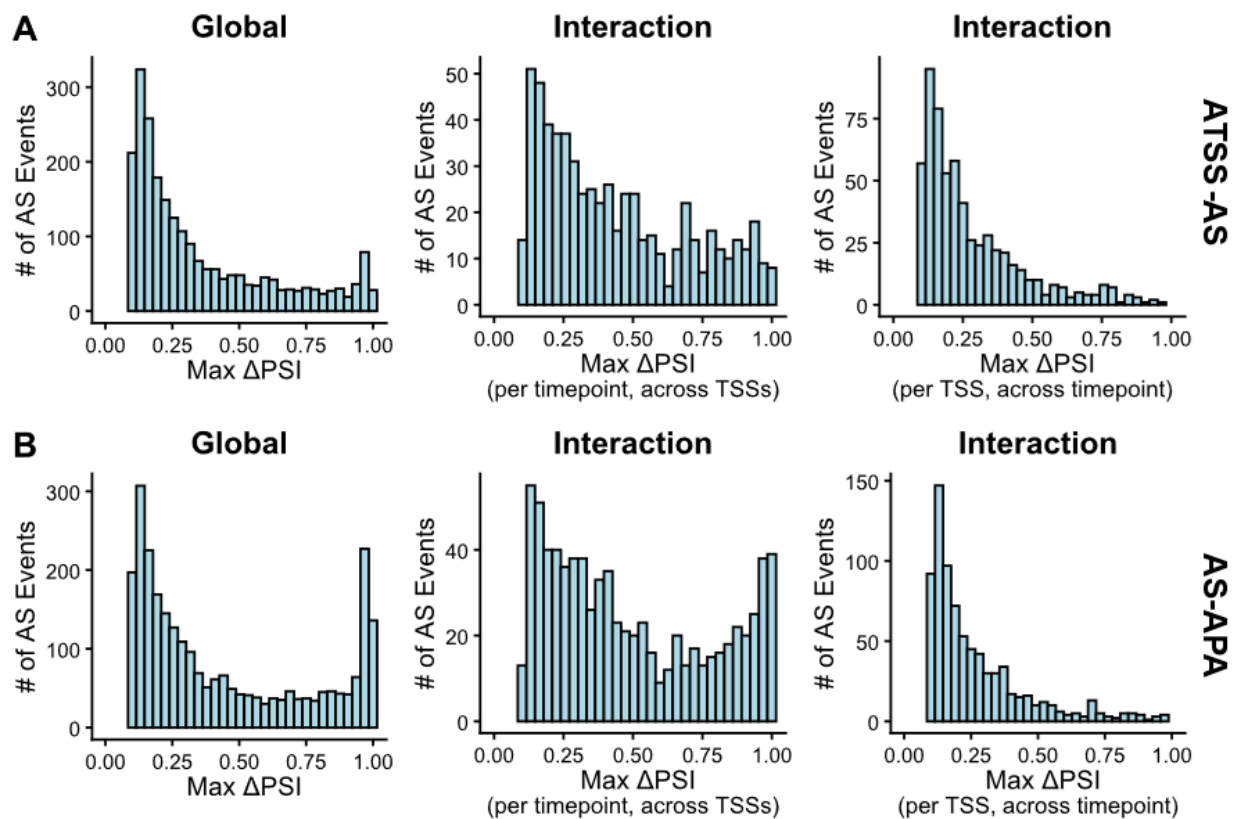


- Fig 5 - what time point was used for Fig. 5A,B?

We thank the reviewer for raising this point. To address this, we have clarified that Figure 7B-D shows all unique coordinated events across all three timepoints in the text and in Figure 7 legend (lines 1395 and 1398).

- It would be helpful to see some kind of histogram or volcano plot showing the magnitudes of the effects, rather than just reporting the number of events.

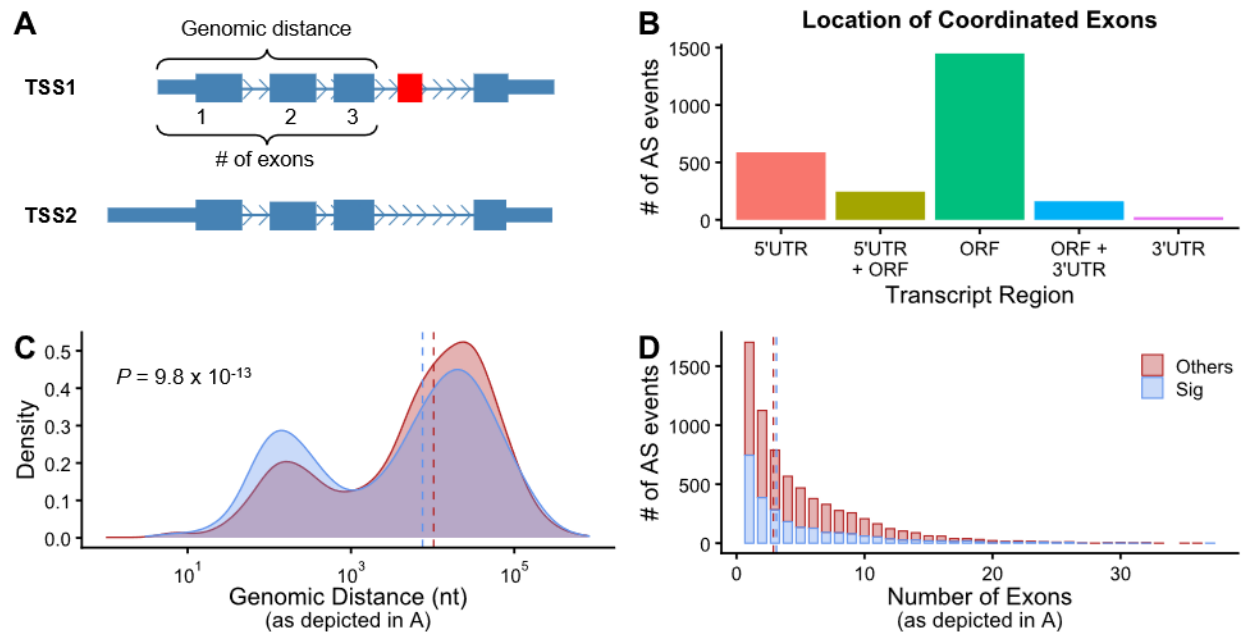
We thank the reviewer for this suggestion. To address this, we have added the maximum Δ PSI (maximum PSI - minimum PSI) between TSSs/pA sites, averaged across time, for all significant global coordination events. As well as the maximum change in PSI between TSSs/pA sites at a single timepoint, and the maximum change in PSI between timepoints, for a single TSS/pA site, for all significant interaction coordination events (Figure S15, reproduced below).



- Fig. 5C,D - It's not clear how to interpret these distributions. Can you compare with some kind of shuffled null distribution? Are the distances and # of exons different from what you might expect under some kind of random model?

We thank the reviewer for this thoughtful suggestion. To address this, we have added the distribution of genomic distances and number of exons for all other (i.e. not statistically coordinated events) (Figure 7-8C & D, Figure 7-AD reproduced below). For the ATSS-AS events, we find the genomic distance of coordinated events to be shorter than all others. For the AS-APA events, we find the genomic distance of coordinated events to be longer than all others. In both cases, the median number of exons is the same. Together, this

suggests that ATSS-AS and APA-AS coordination occurs at statistically shorter and longer, respectively, distances than expected (lines 330-334, 359-362 and 364-366).



Reviewer #3 (Remarks to the Author):

Tripathy and colleagues have produced a deeply sequenced long-read RNA-seq atlas of cortical neuron development, which they combine with proteomics to examine changes in splicing and polyadenylation.

Overall the paper seems methodologically sound and they perform quality control of their isoform set.

We thank this reviewer for their overall very positive reception of our work.

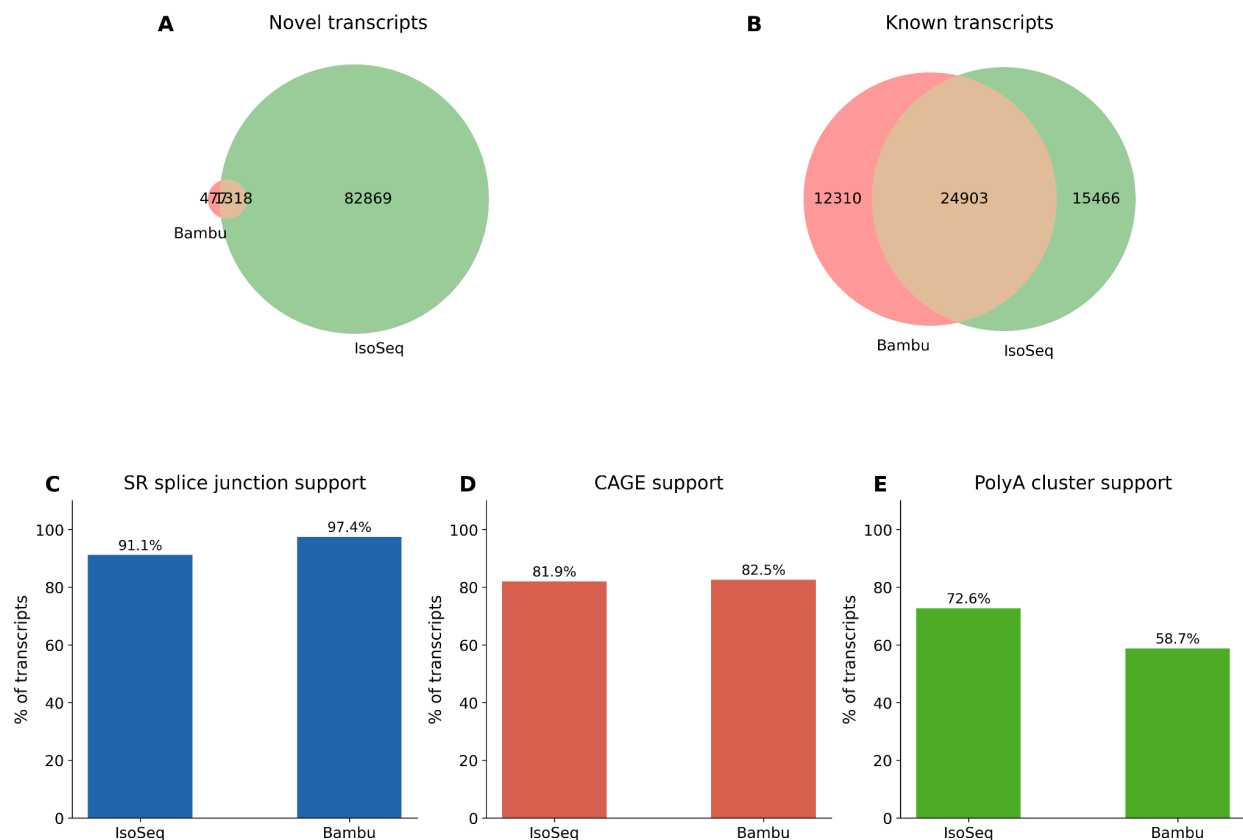
Isoform discovery in long-read RNA-seq is still in its infancy so I would appreciate if the authors compared their isoform catalogue from Pigeon to a few other tools (Stringtie, Bamby, IsoQuant) to check robustness, particularly for their novel isoforms. The assumption would be that the overlap between tools is low - not necessarily a bad thing as this is a hard problem, but important to test as I don't believe Pigeon has been benchmarked against these other tools by LRGASP or other papers.

We thank the reviewer for this suggestion. As a comparison, we ran Bamby (v.3.0.5), an alternative tool for isoform discovery in long-read sequencing data, using default settings and identical downstream filtering (Fig S19., reproduced below).

Transcript sets from both methods were compared by matching shared intron chains. Of the 52,679 known isoforms identified across the two methods, 24,903 (47.3%) were shared, demonstrating strong concordance in the annotated transcriptome. For novel isoforms, agreement was expected to be lower: of the 84,664 novel transcripts identified across both methods, only 1,318 (1.6%) were shared, with the majority of novel calls exclusive to Iso-Seq (97.9%) and a smaller fraction exclusive to Bambu (0.6%). This is consistent with the observation that pairwise overlap of transcript discovery was low, which reflects the different sensitivity thresholds and modeling assumptions of each method.

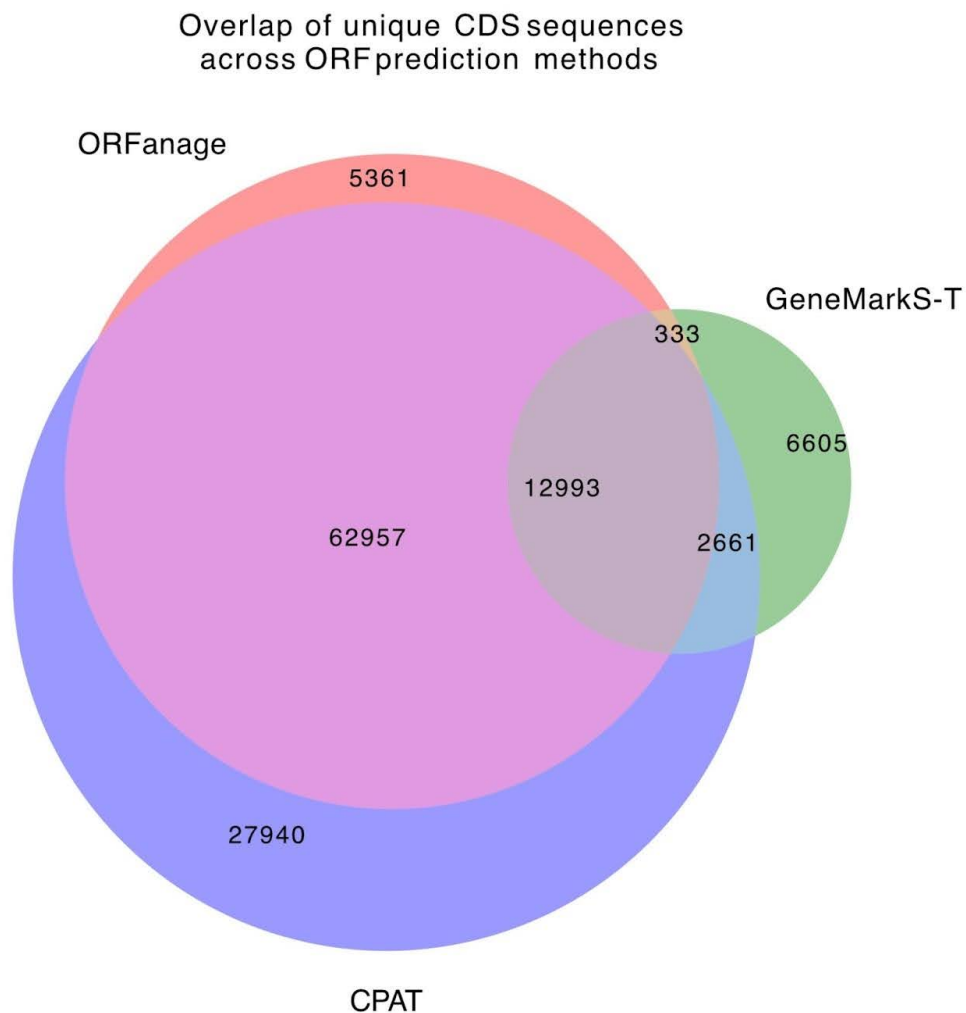
To assess the reliability of the isoforms uniquely identified by Iso-Seq, we performed two independent orthogonal validations. First, 91.1% of these isoforms had all of their splice junctions supported by short-read RNA-seq data, confirming that the splicing events underlying these isoforms are reproducible across sequencing platforms. Second, 81.9% of these isoforms had their 5' end located within 100 bp of a CAGE-seq peak, providing strong evidence for bona fide transcription initiation at these sites. Together, these results suggest that the vast majority of Iso-Seq isoforms represent genuine transcripts, and we are therefore confident in our choice of Iso-Seq for this analysis.

We have added a description of the results from benchmarking transcript discovery methods to the Methods section, at line 707.



I would compare the ORF estimates from ORFanage to a few other tools (GeneMarkS, CPAT) to assess the robustness of those predictions.

We have added a supplementary figure (Fig. S20, reproduced below) that shows the overlap of unique CDS sequences predicted by three ORF calling methods (ORFanage, CPAT and GeneMarkS-T). ORFanage highly overlaps with CPAT, with 94% (75,950/81,644) of ORFanage CDS sequences also predicted by CPAT. GeneMarkS-T has the least overlap with either method, with 30% (6,605/22,592) of its predictions unique to it. The high overlap between ORFanage and CPAT likely reflects their shared dependence on existing genome annotations — ORFanage directly uses reference annotations to guide ORF assignment, while CPAT's coding potential features are trained on known coding sequences. GeneMarkS-T, by contrast, is a self-training probabilistic model that learns codon statistics de novo from the input data, making it algorithmically independent from existing annotations and more likely to produce divergent predictions.



The shiny app browser they provide is very helpful but could benefit from a few additional features. Specifically, the bar plots should show the variance between the replicates (boxplots, mean $\pm$ SEM etc) and it would be helpful to include some way of showing which isoforms are predicted to be protein-coding and have peptide support.

We thank the reviewer for this suggestion. We have updated the shinyApp that allows users to filter for protein-coding isoforms and isoforms with peptide support.

Reviewer #3 (Remarks on code availability):

Appears reproducible, includes dependencies as a conda environment.

REVIEWERS' COMMENTS

Reviewer #1 (Remarks to the Author)

I'm satisfied with the response by the authors. (Remarks on code availability)

We thank this reviewer for their thoughtful feedback.

Reviewer #2 (Remarks to the Author)

The authors have thoroughly responded to my previous comments and those of the other reviewers. I have some remaining small suggestions (below), but I believe the manuscript is greatly improved.

We thank this reviewer for their insightful feedback.

- The authors added Fig. S3 showing pileups of the reads at the Ago1 locus. However, it would be very helpful to show IGV browser shots that display (a downsampled subset of) individual reads, i.e. showing the actual structure of some example long reads. I believe this was the intent of Reviewer 1 question 3, and I agree with that reviewer that such a representation would be very valuable.

We have modified Figure S3 to display pre-filtering read alignments (downsampled to 10%), illustrating the variability among individual reads. However, these alignments do include reads that were subsequently removed during filtering.

- I don't understand the comment: "Displaying reads post-filtering does not provide any variability amongst the reads."

We appreciate the opportunity to clarify this point.

Our analysis pipeline uses isoseq collapse and Pigeon. isoseq collapse is used to merge redundant reads into unique transcript isoforms (based on exonic structures), while allowing for minor variation in splice junction and transcript ends. Pigeon is used to filter out artifacts such as internal priming and RT switching. As a result, post-filtering read alignments show minimal variability per isoform.

As described above, we have modified Figure S3 to display pre-filtering read alignments illustrating the variability among individual reads. However, these alignments do include reads that were subsequently removed during filtering.

- The Shiny browser also does not show individual reads. An IGV browser showing BAM files would be most helpful.

Thank you for this helpful suggestion. We agree that visualization of individual reads can be informative. The BAM files underlying our transcripts are available for download, allowing readers to inspect the aligned reads directly in IGV or another genome browser of their choice.

- In the new version of Fig 3A, there is no explanation in the legend for the Venn diagrams - Please add this to the legend.

We thank the reviewer for highlighting this. We have added an explanation in the legend for Figure 3A to describe the Venn diagrams.

(Remarks on code availability)

Reviewer #3 (Remarks to the Author)

I thank for the authors for their hard work in making their results more robust. I particularly appreciate the improvements to the Shiny app as it's now much clearer.

We thank this reviewer for their thoughtful feedback.

I am satisfied that the study is now suitable for publication.

(Remarks on code availability)

README seems extensive with instructions for installation and running.