

Sequence analysis

Pacybara: accurate long-read sequencing for barcoded mutagenized allelic libraries

Jochen Weile ^{1,2,3,4,*}, Gabrielle Ferra⁵, Gabriel Boyle⁵, Sriram Pendyala⁵, Clara Amorosi⁵, Chiann-Ling Yeh ⁵, Atina G. Cote^{1,2,3,4}, Nishka Kishore^{1,2,3,4}, Daniel Tabet^{1,2,3,4}, Warren van Loggerenberg^{1,2,3,4,8}, Ashyad Rayhan^{1,2,3,4}, Douglas M. Fowler^{5,6,7}, Maitreya J. Dunham ⁵, Frederick P. Roth ^{1,2,3,4,8,*}

¹Lunenfeld-Tanenbaum Research Institute, Sinai Health, Toronto, ON M5G 1X5, Canada

²Donnelly Centre, University of Toronto, Toronto, ON M5S 3E1, Canada

³Department of Molecular Genetics, University of Toronto, Toronto, ON M5S 3E1, Canada

⁴Department of Computer Science, University of Toronto, Toronto, ON M5S 2E4, Canada

⁵Department of Genome Sciences, University of Washington, Seattle, WA 98195, United States

⁶Department of Bioengineering, University of Washington, Seattle, WA 98195, United States

⁷Brotman Baty Institute for Precision Medicine, Seattle, WA 98195, United States

⁸Department of Computational & Systems Biology, School of Medicine, University of Pittsburgh, Pittsburgh, PA 15213, United States

*Corresponding authors. Department of Computational & Systems Biology, School of Medicine, University of Pittsburgh, 800 Murdoch I Bldg 3420 Forbes Avenue, Pittsburgh, PA 15213, USA. E-mail: fritz@pitt.edu (F.P.R.); The Donnelly Centre, University of Toronto, 160 College Street, Toronto, ON M5S 3E1, Canada. E-mail: jochen.weile@mail.utoronto.ca (J.W.)

Associate Editor: Can Alkan

Abstract

Motivation: Long-read sequencing technologies, an attractive solution for many applications, often suffer from higher error rates. Alignment of multiple reads can improve base-calling accuracy, but some applications, e.g. sequencing mutagenized libraries where multiple distinct clones differ by one or few variants, require the use of barcodes or unique molecular identifiers. Unfortunately, sequencing errors can interfere with correct barcode identification, and a given barcode sequence may be linked to multiple independent clones within a given library.

Results: Here we focus on the target application of sequencing mutagenized libraries in the context of multiplexed assays of variant effects (MAVEs). MAVEs are increasingly used to create comprehensive genotype-phenotype maps that can aid clinical variant interpretation. Many MAVE methods use long-read sequencing of barcoded mutant libraries for accurate association of barcode with genotype. Existing long-read sequencing pipelines do not account for inaccurate sequencing or nonunique barcodes. Here, we describe Pacybara, which handles these issues by clustering long reads based on the similarities of (error-prone) barcodes while also detecting barcodes that have been associated with multiple genotypes. Pacybara also detects recombinant (chimeric) clones and reduces false positive indel calls. In three example applications, we show that Pacybara identifies and correctly resolves these issues.

Availability and implementation: Pacybara, freely available at <https://github.com/rothlab/pacybara>, is implemented using R, Python, and bash for Linux. It runs on GNU/Linux HPC clusters via Slurm, PBS, or GridEngine schedulers. A single-machine simplex version is also available.

1 Introduction

Multiplexed assays of variant effect (MAVEs) often involve the use of clone libraries in which each mutagenized allele is associated with a barcode (Tabet *et al.* 2022), requiring determination of full-length amplicon sequences. Such libraries are typically designed with barcode loci closely adjacent to either or both sides of a mutagenized ORFs sequence. Long-read sequencing technologies such as Pacbio Sequel II, Revio, and Oxford Nanopore have become an attractive alternative to short-read amplicon assembly methods (Hiatt *et al.* 2010, Weile *et al.* 2017).

The first pipeline for long-read barcoded library assembly was AssemblyByPacbio (ABP) (Matreyek *et al.* 2018). ABP extracts barcode sequences from each read, groups reads by identical barcodes, and derives the clone genotype from the

highest quality read in each group. Recently, the PacRAT pipeline was developed in which the ABP output is post-processed to derive a multiple sequence alignment-based consensus for each read group sharing a barcode (Yeh *et al.* 2022).

At least three problems have remained unsolved: First, sequencing errors in the barcode can cause a single clone to appear as two or more distinct clones. Second, a barcode sequence may be ‘nonunique’, in that the same sequence is associated with multiple clones of different genotypes within a given library. Although an error-aware clustering method for short barcode sequences (Zhao *et al.* 2018) addresses the first issue for short read barcode identification, it does not consider clonality of genotypes beyond the barcode itself. Third, problems arising during library preparation, e.g. PCR crossover events, can yield consensus genotypes that are recombinant

Received: 19 December 2023; Revised: 5 March 2024; Editorial Decision: 27 March 2024; Accepted: 2 April 2024

© The Author(s) 2024. Published by Oxford University Press.

This is an Open Access article distributed under the terms of the Creative Commons Attribution License (<https://creativecommons.org/licenses/by/4.0/>), which permits unrestricted reuse, distribution, and reproduction in any medium, provided the original work is properly cited.

chimeras of wild-type (WT) or other variant genotypes. Others have described double-barcode methods able to detect chimeras, but these approaches are not applicable to singly barcoded libraries in which clones might differ by only a few SNVs, as is the case for mutagenized libraries used in MAVEs (Karst *et al.* 2021, Ramírez-Rojas *et al.* 2024).

To address these issues, we extended the clustering approaches cited above using not only the barcode sequences from each read but also the full set of candidate variants and their quality metrics. Although initial stages of the clustering process, like previous methods, focused on merging pairs of reads with identical candidate barcodes, Pacybara accounts for variant quality and the extent of overlap between candidate variants. Pacybara then further considers merging reads with similar but nonidentical barcodes.

2 Materials and methods

Pacybara, designed to run on high-performance computing clusters, deploys and supervises individual jobs on cluster nodes and collates the results. In the first processing step, long reads are distributed across jobs and aligned to the reference sequence via *bwa* (Li and Durbin 2009). From the alignments, barcode sequences and lists of candidate variants (i.e. basecalls that apparently differ from the reference sequence) are extracted, including their respective quality scores. For libraries designed with multiple barcodes per molecule, the latter can optionally be combined to form a single ‘virtual barcode’ each. Reads in which quality scores within the barcode region fall below a tunable threshold (default Q62) are filtered out, as they often indicate multi-occupancy SMRT wells. Edit distances between all pairs of barcodes are calculated via alignment with *bowtie2* (Langmead and Salzberg 2012).

Clustering first considers sets of reads with identical barcode sequences. Within each identical-barcode set, putative sequencing errors are excluded by filtering out candidate variants seen in only one read and with a below-threshold quality score. Then for each identical-barcode read set, a graph is constructed such that each node corresponds to a read, and an edge is added between two reads either if both are WT or if they harbor a sufficiently similar set of candidate variants (defined by a Jaccard index threshold). A ‘seed cluster’ is then formed for each graph component (including ‘singleton’ nodes without edges).

Cluster definitions, initially defined by the set of seed clusters, are updated in a cluster-merging process. Pairs of clusters, chosen by the minimum edit distance (ED) of their respective member reads, are iteratively considered as candidate clusters up to a maximum edit distance (default 2). After filtering sequencing errors from this candidate cluster as above, the candidate cluster is accepted if all of the following are true: (i) all reads in both clusters are WT or the sets of remaining candidate variants in the respective pair of clusters have an above-threshold Jaccard coefficient, (ii) no pair of reads in the proposed cluster has a barcode edit distance greater than the maximally allowed distance, and (iii) the pair of clusters has sufficiently divergent sizes, as might be expected if one cluster represents a subset of reads derived from the same clone, but with more base-calling errors in the barcode. Here the clusters were considered to have sufficiently divergent sizes if $|\log_2(\text{size}_1/\text{size}_2)| > \text{ED}$, similar to a previous approach for short-read barcode clustering (Zhao *et al.* 2018). We used a Jaccard threshold of 0.2 in each case.

For each cluster in the final set, consensus barcodes are derived via *muscle* (Edgar 2022), variants are filtered as above to remove putative base-calling errors, and (for coding sequences) remaining variants are translated to protein consequences.

3 Results and discussion

We evaluated performance of both Pacybara and the previously established method PacRAT. Both tools were applied to three different barcoded mutagenized open reading frame libraries, designed for use in a MAVE: (i) a mutagenized human *CYP2C9* library (Amorosi *et al.* 2021), (ii) a mutagenized human *CYP2C19* library (Boyle *et al.* 2023), and (iii) a previously unpublished mutagenized human *LDLR* library, which we suspected of harboring nonunique barcodes and PCR chimeras.

All of the above datasets were generated via PacBio HiFi sequencing. While we have not evaluated Pacybara on Nanopore sequencing data, we expect that the latter would be compatible as well, although perhaps requiring tuning of quality thresholds.

Differences between Pacybara and PacRAT were immediately apparent in the distribution of cluster sizes (Supplementary Fig. S1). In all three datasets, Pacybara produced a larger number of small clusters, which we attribute to (i) more aggressive filtering of low-quality reads and (ii) fewer erroneous merges between clones with nonunique barcodes by considering the coherence of clone genotypes before merging.

For all three datasets, both pipelines reported similar numbers of clusters and similar distributions of the number of mutations. A notable exception to this is frameshift variants in *CYP2C9*, nearly half of which were filtered out by Pacybara (Supplementary Fig. S2). This is consistent with false-positive indels being a common type of sequencing error in PacBio sequencing, especially at homopolymer sites. Indeed, we found that frameshift calls at the 5' end of homopolymer runs of length 4 or greater were 2.3 times as likely to be filtered out by Pacybara compared to other positions (Supplementary Fig. S3).

We next wished to examine how well variant calls agreed between PacRAT and Pacybara for each barcoded clone. Of those clones that were not filtered out by either tool, 88%, 79%, and 96% agreed completely (Supplementary Fig. S4). We then examined the raw read data of five arbitrarily chosen clusters with disagreeing genotypes. In three out of five cases, the pipelines primarily disagreed on the exact lengths of long indels. The remaining two cases disagreed on whether to include or exclude low-quality variants that only occur in some but not all reads within a cluster (see Supplementary Text S1).

To validate whether the barcode-genotype associations produced by either pipeline are more accurate, we tested them in terms of their impact on protein function. As nonsense and frameshift variants are likely to damage protein function while synonymous variants are not, we compared, for each pipeline, the distributions of functionality scores for clones assigned a nonsense and frameshift variant against those bearing only synonymous variants. Functional scores were calculated based on enrichment or depletion in a fluorescence-based assay (measuring protein abundance for *CYP2C19*, enzyme activity for *CYP2C9*, and the ability of the cell to uptake LDL particles for *LDLR*). For all three datasets, score distributions based on both pipelines were very similar, with Pacybara's frameshift scores being slightly more likely to be deleterious (Supplementary Fig. S5).

Next, we examined whether Pacybara was able to detect clones with nonunique barcodes. For the *CYP2C9* and *CYP2C19* libraries, only a small fraction of barcodes were nonunique (Supplementary Fig. S6). For the *LDLR* library 22.4% of barcodes were nonunique. However, 17.3% were classified as ‘remediable’, in that a single clone represented more than two thirds of reads for that barcode and would thus likely dominate readouts for that barcode in a screen. We manually examined eight arbitrarily chosen nonunique barcodes from the *LDLR* library which had conflicting genotypes assigned by the two pipelines. In one case, PacRAT called the genotype correctly, while Pacybara did not, as it included a frequently occurring sequencing error occurring in multiple reads. However, in six out of the eight cases Pacybara called the correct genotype by splitting the cluster. One final case could not be conclusively decided due to ambiguous read data (see Supplementary Text S2).

Pacybara’s use of virtual barcodes, concatenating barcodes from both flanks of the same construct, allows it to identify putative PCR crossover events. The *LDLR* library made use of such flanking barcodes, allowing us to test this feature. Indeed, after enabling virtual barcodes we identified 74 485 sets of clusters in which upstream barcodes were identical and either complete or partial overlap in genotype existed, but which had entirely different downstream barcodes (Supplementary Table S1). We consider these cases likely PCR chimeras. To investigate the relationship between the putative chimeras and barcodes that are found to be nonunique without the use of virtual barcodes, we counted how often single upstream barcodes were flagged as a chimera or nonunique or both. We found that for 99.99% of cases in which a barcode was flagged as nonunique, it was also found to be involved in a putative PCR chimera. Similarly, if a barcode was not involved in a chimera, it had a 99.99% chance of being unique. However, out of all barcodes putatively involved in chimeras, it had only a 21% chance of being detected as nonunique when virtual barcodes were disabled. This observation confirms the value of dual flanking barcodes, which were unfortunately not a feature of the *CYP2C9* and *CYP2C19* libraries, while showing that detecting nonunique barcodes can at least serve as a way of alleviating problems caused by chimeras. No other pipeline has been described that detects chimeras in barcoded clone libraries where clones differ by only one or few nucleotides.

The *CYP2C19*, *CYP2C9*, and *LDLR* datasets represent a diverse array of library sizes and sequencing depths. They contain ~170 000, ~60 000, and ~80 000 barcoded clones with average cluster sizes of 7.9, 3.6, and 6.8, respectively. Even at a modest depth of 3.6 reads per cluster, the *CYP2C9* dataset still produced reliable results.

In summary, we developed Pacybara to process long-read data from libraries of barcoded mutagenized clones, and found that it successfully detects nonunique barcodes and PCR crossover events, combines reads with barcodes that differ only due to sequencing error, reduces false-positive indel calls, and thus improves the quality of a downstream MAVE experiment.

Acknowledgements

We would like to acknowledge Kelley Harris for her advice.

Supplementary data

Supplementary data are available at *Bioinformatics* online.

© The Author(s) 2024. Published by Oxford University Press.

Bioinformatics, 2024, 40, 1–3

<https://doi.org/10.1093/bioinformatics/btae182>

Applications Note

Conflict of interest

F.P.R. is a shareholder and advisor for SeqWell, Constantiam, BioSymetrics, and a shareholder of Ranomics.

Funding

This work was supported by the National Human Genome Research Institute (NHGRI) of the National Institutes of Health (NIH) Center of Excellence in Genomic Science Initiative (CEGS) [RM1 HG010461]; the NHGRI Impact of Genomic Variation on Function Initiative’s Center for Actionable Variant Analysis [UM1 HG011969]; the NIH [R01 HL52066]; National Institute of General Medical Sciences of the NIH [R01 GM132162 to M.J.D. and D.M.F.]; the NIH/NIGMS [R35GM133428 to Kelley Harris]; the NIH [UM1 HG011989 and R01 HL164675 to F.P.R.]; and a Canadian Institutes of Health Research Foundation Grant [to F.P.R.]

Data availability

The data underlying this article are available in NCBI GEO at <https://www.ncbi.nlm.nih.gov/gds/> and NCBI SRA at <https://www.ncbi.nlm.nih.gov/sra/>, and can be accessed with GSE165412, GSE244489, and SRP477344.

References

- Amorosi CJ, Chiasson MA, McDonald MG *et al.* Massively parallel characterization of CYP2C9 variant enzyme activity and abundance. *Am J Hum Genet* 2021;108:1735–51.
- Boyle GE, Sitko K, Galloway JG *et al.* Deep mutational scanning of CYP2C19 reveals a substrate specificity-abundance tradeoff. bioRxiv, <https://doi.org/10.1101/2023.10.06.561250>, 2023, preprint: not peer reviewed.
- Edgar RC. Muscle5: high-accuracy alignment ensembles enable unbiased assessments of sequence homology and phylogeny. *Nat Commun* 2022;13:6968. <https://doi.org/10.1038/s41467-022-34630-w>
- Hiatt JB, Patwardhan RP, Turner EH *et al.* Parallel, tag-directed assembly of locally derived short sequence reads. *Nat Methods* 2010;7: 119–22. <https://doi.org/10.1038/nmeth.1416>
- Karst SM, Ziels RM, Kirkegaard RH *et al.* High-accuracy long-read amplicon sequences using unique molecular identifiers with nanopore or PacBio sequencing. *Nat Methods* 2021;18:165–9.
- Langmead B, Salzberg S. Fast gapped-read alignment with Bowtie 2. *Nat Methods* 2012;9:357–9. <https://doi.org/10.1038/nmeth.1923>
- Li H, Durbin R. Fast and accurate short read alignment with Burrows-Wheeler transform. *Bioinformatics* 2009;25:1754–60. <https://doi.org/10.1093/bioinformatics/btp324>
- Matreyek KA, Starita LM, Stephany JJ *et al.* Multiplex assessment of protein variant abundance by massively parallel sequencing. *Nat Genet* 2018;50:874–82.
- Ramírez-Rojas AA *et al.* DuBA.flow—a low-cost, long-read amplicon sequencing workflow for the validation of synthetic DNA constructs. *ACS Synth Biol* 2024;13:457–65. <https://doi.org/10.1021/acssynbio.3c00522>
- Tabet D, Parikh V, Mali P *et al.* Scalable functional assays for the interpretation of human genetic variation. *Annu Rev Genet* 2022;56: 441–65. <https://doi.org/10.1146/annurev-genet-072920-032107>
- Weile J, Sun S, Cote AG *et al.* A framework for exhaustively mapping functional missense variants. *Mol Syst Biol* 2017;13:957.
- Yeh C-LC, Amorosi CJ, Showman S *et al.* PacRAT: a program to improve barcode-variant mapping from PacBio long reads using multiple sequence alignment. *Bioinformatics* 2022;38:2927–9.
- Zhao L, Liu Z, Levy SF *et al.* Bartender: a fast and accurate clustering algorithm to count barcode reads. *Bioinformatics* 2018;34:739–47.