

RESEARCH ARTICLE

HUMAN GENOMICS

The complete sequence of a human genome

Sergey Nurk^{1†}, Sergey Koren^{1†}, Arang Rhie^{1†}, Mikko Rautiainen^{1†}, Andrey V. Bzikadze², Alla Mikheenko³, Mitchell R. Vollger⁴, Nicolas Altemose⁵, Lev Uralsky^{6,7}, Ariel Gershman⁸, Sergey Aganezov^{9†}, Savannah J. Hoyt¹⁰, Mark Diekhans¹¹, Glennis A. Logsdon⁴, Michael Alonge⁹, Stylianos E. Antonarakis¹², Matthew Borchers¹³, Gerard G. Bouffard¹⁴, Shelise Y. Brooks¹⁴, Gina V. Caldas¹⁵, Nae-Chyun Chen⁹, Haoyu Cheng^{16,17}, Chen-Shan Chin¹⁸, William Chow¹⁹, Leonardo G. de Lima¹³, Philip C. Dishuck⁴, Richard Durbin^{19,20}, Tatiana Dvorkina³, Ian T. Fiddes²¹, Giulio Formenti^{22,23}, Robert S. Fulton²⁴, Arkarachai Functammasan¹⁸, Erik Garrison^{11,25}, Patrick G. S. Grady¹⁰, Tina A. Graves-Lindsay²⁶, Ira M. Hall²⁷, Nancy F. Hansen²⁸, Gabrielle A. Hartley¹⁰, Marina Haukness¹¹, Kerstin Howe¹⁹, Michael W. Hunkapiller²⁹, Chirag Jain^{1,30}, Miten Jain¹¹, Erich D. Jarvis^{22,23}, Peter Kerpedjiev³¹, Melanie Kirsche⁹, Mikhail Kolmogorov³², Jonas Koriach²⁹, Milinn Kremitzki²⁶, Heng Li^{16,17}, Valerie V. Maduro³³, Tobias Marschall³⁴, Ann M. McCartney¹, Jennifer McDaniel³⁵, Danny E. Miller^{4,36}, James C. Mullikin^{14,28}, Eugene W. Myers³⁷, Nathan D. Olson³⁵, Benedict Paten¹¹, Paul Peluso²⁹, Pavel A. Pevzner³², David Porubsky⁴, Tamara Potapova¹³, Evgeny I. Rogae^{6,7,38,39}, Jeffrey A. Rosenfeld⁴⁰, Steven L. Salzberg^{9,41}, Valerie A. Schneider⁴², Fritz J. Sedlaczek⁴³, Kishwar Shafin¹¹, Colin J. Shew⁴⁴, Alaina Shumate⁴¹, Ying Sims¹⁹, Arian F. A. Smit⁴⁵, Daniela C. Soto⁴⁴, Ivan Sovic^{29,46}, Jessica M. Storer⁴⁵, Aaron Streets^{5,47}, Beth A. Sullivan⁴⁸, Françoise Thibaud-Nissen⁴², James Torrance¹⁹, Justin Wagner³⁵, Brian P. Walenz¹, Aaron Wenger²⁹, Jonathan M. D. Wood¹⁹, Chunlin Xiao⁴², Stephanie M. Yan⁴⁹, Alice C. Young¹⁴, Samantha Zarate⁹, Urvashi Surti⁵⁰, Rajiv C. McCoy⁴⁹, Megan Y. Dennis⁴⁴, Ivan A. Alexandrov^{3,7,51}, Jennifer L. Gerton^{13,52}, Rachel J. O'Neill¹⁰, Winston Timp^{8,41}, Justin M. Zook³⁵, Michael C. Schatz^{9,49}, Evan E. Eichler^{4,53*}, Karen H. Miga^{11,54*}, Adam M. Phillippy^{1*}

Since its initial release in 2000, the human reference genome has covered only the euchromatic fraction of the genome, leaving important heterochromatic regions unfinished. Addressing the remaining 8% of the genome, the Telomere-to-Telomere (T2T) Consortium presents a complete 3.055 billion-base pair sequence of a human genome, T2T-CHM13, that includes gapless assemblies for all chromosomes except Y, corrects errors in the prior references, and introduces nearly 200 million base pairs of sequence containing 1956 gene predictions, 99 of which are predicted to be protein coding. The completed regions include all centromeric satellite arrays, recent segmental duplications, and the short arms of all five acrocentric chromosomes, unlocking these complex regions of the genome to variational and functional studies.

The current human reference genome was released by the Genome Reference Consortium (GRC) in 2013 and most recently patched in 2019 (GRCh38.p13) (1). This reference traces its origin to the publicly

funded Human Genome Project (2) and has been continually improved over the past two decades. Unlike the competing Celera effort (3) and most modern sequencing projects based on “shotgun” sequence assembly (4),

the GRC assembly was constructed from sequenced bacterial artificial chromosomes (BACs) that were ordered and oriented along the human genome by means of radiation hybrid, genetic linkage, and fingerprint maps. However, limitations of BAC cloning led to an underrepresentation of repetitive sequences, and the opportunistic assembly of BACs derived from multiple individuals resulted in a mosaic of haplotypes. As a result, several GRC assembly gaps are unsolvable because of incompatible structural polymorphisms on their flanks, and many other repetitive and polymorphic regions were left unfinished or incorrectly assembled (5).

The GRCh38 reference assembly contains 151 mega-base pairs (Mbp) of unknown sequence distributed throughout the genome, including pericentromeric and subtelomeric regions, recent segmental duplications, ampliconic gene arrays, and ribosomal DNA (rDNA) arrays, all of which are necessary for fundamental cellular processes (Fig. 1A). Some of the largest reference gaps include human satellite (HSat) repeat arrays and the short arms of all five acrocentric chromosomes, which are represented in GRCh38 as multimegabase stretches of unknown bases (Fig. 1, B and C). In addition to these apparent gaps, other regions of GRCh38 are artificial or are otherwise incorrect. For example, the centromeric alpha satellite arrays are represented as computationally generated models of alpha satellite monomers to serve as decoys for resequencing analyses (6), and sequence assigned to the short arm of chromosome 21 appears falsely duplicated and poorly assembled (7). When compared with other human genomes, GRCh38 also shows a genome-wide deletion bias that is indicative of incomplete assembly (8). Despite finishing efforts from both the Human Genome Project (9) and GRC (1) that improved the quality of the reference, there was limited

¹Genome Informatics Section, Computational and Statistical Genomics Branch, National Human Genome Research Institute, National Institutes of Health, Bethesda, MD, USA. ²Graduate Program in Bioinformatics and Systems Biology, University of California, San Diego, La Jolla, CA, USA. ³Center for Algorithmic Biotechnology, Institute of Translational Biomedicine, Saint Petersburg State University, Saint Petersburg, Russia. ⁴Department of Genome Sciences, University of Washington School of Medicine, Seattle, WA, USA. ⁵Department of Bioengineering, University of California, Berkeley, Berkeley, CA, USA. ⁶Sirius University of Science and Technology, Sochi, Russia. ⁷Vavilov Institute of General Genetics, Moscow, Russia. ⁸Department of Molecular Biology and Genetics, Johns Hopkins University, Baltimore, MD, USA. ⁹Department of Computer Science, Johns Hopkins University, Baltimore, MD, USA. ¹⁰Institute for Systems Genomics and Department of Molecular and Cell Biology, University of Connecticut, Storrs, CT, USA. ¹¹UC Santa Cruz Genomics Institute, University of California, Santa Cruz, Santa Cruz, CA, USA. ¹²University of Geneva Medical School, Geneva, Switzerland. ¹³Stowers Institute for Medical Research, Kansas City, MO, USA. ¹⁴NIH Intramural Sequencing Center, National Human Genome Research Institute, National Institutes of Health, Bethesda, MD, USA. ¹⁵Department of Molecular and Cell Biology, University of California, Berkeley, Berkeley, CA, USA. ¹⁶Department of Data Sciences, Dana-Farber Cancer Institute, Boston, MA, USA. ¹⁷Department of Biomedical Informatics, Harvard Medical School, Boston, MA, USA. ¹⁸DNA Nexus, Mountain View, CA, USA. ¹⁹Wellcome Sanger Institute, Cambridge, UK. ²⁰Department of Genetics, University of Cambridge, Cambridge, UK. ²¹Inscripta, Boulder, CO, USA. ²²Laboratory of Neurogenetics of Language and The Vertebrate Genome Lab, The Rockefeller University, New York, NY, USA. ²³Howard Hughes Medical Institute, The Rockefeller University, New York, NY, USA. ²⁴Department of Genetics, Washington University School of Medicine, St. Louis, MO, USA. ²⁵University of Tennessee Health Science Center, Memphis, TN, USA. ²⁶McDonnell Genome Institute, Washington University in St. Louis, St. Louis, MO, USA. ²⁷Department of Genetics, Yale University School of Medicine, New Haven, CT, USA. ²⁸Comparative Genomics Analysis Unit, Cancer Genetics and Comparative Genomics Branch, National Human Genome Research Institute, National Institutes of Health, Bethesda, MD, USA. ²⁹Pacific Biosciences, Menlo Park, CA, USA. ³⁰Department of Computational and Data Sciences, Indian Institute of Science, Bangalore KA, India. ³¹Reservoir Genomics LLC, Oakland, CA, USA. ³²Department of Computer Science and Engineering, University of California, San Diego, La Jolla, CA, USA. ³³Undiagnosed Diseases Program, National Human Genome Research Institute, National Institutes of Health, Bethesda, MD, USA. ³⁴Heinrich Heine University Düsseldorf, Medical Faculty, Institute for Medical Biometry and Bioinformatics, Düsseldorf, Germany. ³⁵Biosystems and Biomaterials Division, National Institute of Standards and Technology, Gaithersburg, MD, USA. ³⁶Department of Pediatrics, Division of Genetic Medicine, University of Washington and Seattle Children's Hospital, Seattle, WA, USA. ³⁷Max Planck Institute of Molecular Cell Biology and Genetics, Dresden, Germany. ³⁸Department of Psychiatry, University of Massachusetts Medical School, Worcester, MA, USA. ³⁹Faculty of Biology, Lomonosov Moscow State University, Moscow, Russia. ⁴⁰Cancer Institute of New Jersey, New Brunswick, NJ, USA. ⁴¹Department of Biomedical Engineering, Johns Hopkins University, Baltimore, MD, USA. ⁴²National Center for Biotechnology Information, National Library of Medicine, National Institutes of Health, Bethesda, MD, USA. ⁴³Human Genome Sequencing Center, Baylor College of Medicine, Houston, TX, USA. ⁴⁴Genome Center, MIND Institute, Department of Biochemistry and Molecular Medicine, University of California, Davis, CA, USA. ⁴⁵Institute for Systems Biology, Seattle, WA, USA. ⁴⁶Digital BioLogic d.o.o., Ivanić-Grad, Croatia. ⁴⁷Chan Zuckerberg Biohub, San Francisco, CA, USA. ⁴⁸Department of Molecular Genetics and Microbiology, Duke University School of Medicine, Durham, NC, USA. ⁴⁹Department of Biology, Johns Hopkins University, Baltimore, MD, USA. ⁵⁰Department of Pathology, University of Pittsburgh, Pittsburgh, PA, USA. ⁵¹Research Center of Biotechnology of the Russian Academy of Sciences, Moscow, Russia. ⁵²Department of Biochemistry and Molecular Biology, University of Kansas Medical School, Kansas City, MO, USA. ⁵³Howard Hughes Medical Institute, University of Washington, Seattle, WA, USA. ⁵⁴Department of Biomolecular Engineering, University of California, Santa Cruz, Santa Cruz, CA, USA.

*Corresponding author. Email: eee@gs.washington.edu (E.E.E.); khmiga@ucsc.edu (K.H.M.); adam.phillippy@nih.gov (A.M.P.)

†These authors contributed equally to this work. ‡Present address: Oxford Nanopore Technologies Inc., Lexington, MA, USA.

progress toward closing the remaining gaps in the years that followed (Fig. 1D).

Long-read shotgun sequencing overcomes the limitations of BAC-based assembly and bypasses the challenges of structural polymorphism between genomes. PacBio's multi-kilobase, single-molecule reads (10) proved capable of resolving complex structural variation and gaps in GRCh38 (8, 11), whereas Oxford Nanopore's >100-kbp "ultralong" reads (12) enabled complete assemblies of a human centromere (chromosome Y) (13) and, later, an entire chromosome (chromosome X) (14). However, the high error rate (>5%) of these technologies posed challenges for the assembly of long, near-identical repeat arrays. PacBio's most recent "HiFi" circular consensus sequencing offers a compromise of 20-kbp read lengths with an error rate of 0.1% (15). Whereas ultralong reads are useful for spanning repeats, HiFi reads excel at differentiating subtly diverged repeat copies or haplotypes (16).

To finish the last remaining regions of the genome, we leveraged the complementary aspects of PacBio HiFi and Oxford Nanopore ultralong-read sequencing to assemble the uniformly homozygous CHM13hTERT cell line (hereafter, CHM13) (17). The resulting T2T-CHM13 reference assembly removes a 20-year-old barrier that has hidden 8% of the genome from sequence-based analysis, including all centromeric regions and the entire short arms of five human chromosomes. Here, we describe the construction, validation, and initial analysis of a truly complete human reference genome and discuss its potential impact on the field.

Cell line and sequencing

As with many prior reference genome improvement efforts (1, 8, 17–20), including the T2T assemblies of human chromosomes X (14) and 8 (21), we targeted a complete hydatidiform mole (CHM) for sequencing. Most CHM genomes arise from the loss of the maternal complement and duplication of the paternal complement postfertilization and are, therefore, homozygous with a 46,XX karyotype (22). Sequencing of CHM13 confirmed nearly uniform homozygosity, with the exception of a few thousand heterozygous variants and a megabase-scale heterozygous deletion within the rDNA array on chromosome 15 (23) (figs. S1 and S2). Local ancestry analysis shows that most of the CHM13 genome is of European origin, including regions of Neanderthal introgression, with some predicted admixture (23) (Fig. 1A). Compared with diverse samples from the 1000 Genomes Project (1KGP) (24), CHM13 possesses no apparent excess of singleton alleles or loss-of-function variants (25).

We extensively sequenced CHM13 with multiple technologies (23), including 30× PacBio

Table 1. Comparison of GRCh38 and T2T-CHM13v1.1 human genome assemblies. GRCh38 summary statistics exclude "alts" (110 Mbp), patches (63 Mbp), and chromosome Y (58 Mbp). Assembled bases include all non-N bases. Unplaced bases are those not assigned or positioned within a chromosome. GRCh38 scaffolds were split at three consecutive Ns to obtain the number of contigs. Contig NG50 is the largest value such that contigs of at least this size total more than half of the 3.05-Gbp genome size. The number of exclusive genes or transcripts is as follows: for GRCh38, GENCODE genes and transcripts not found in CHM13; and for CHM13, extra putative paralogs that are not in GENCODE. Segmental duplication analysis is from (42). RepeatMasker analysis is from (49). Blank spaces indicate not applicable.

STATISTICS	GRCH38	T2T-CHM13	DIFFERENCE (±%)
Summary			
Assembled bases (Gbp)	2.92	3.05	+4.5
Unplaced bases (Mbp)	11.42	0	–100.0
Gap bases (Mbp)	120.31	0	–100.0
Number of contigs	949	24	–97.5
Contig NG50 (Mbp)	56.41	154.26	+173.5
Number of issues	230	46	–80.0
Issues (Mbp)	230.43	8.18	–96.5
Gene annotation			
Number of genes	60,090	63,494	+5.7
Protein coding	19,890	19,969	+0.4
Number of exclusive genes	263	3,604	
Protein coding	63	140	
Number of transcripts	228,597	233,615	+2.2
Protein coding	84,277	86,245	+2.3
Number of exclusive transcripts	1,708	6,693	
Protein coding	829	2,780	
Segmental duplications			
Percentage of segmental duplications (%)	5.00	6.61	
Segmental duplication bases (Mbp)	151.71	201.93	+33.1
Number of segmental duplications	24097	41528	+72.3
RepeatMasker			
Percentage of repeats (%)	51.89	53.94	
Repeat bases (Mbp)	1,516.37	1,647.81	+8.7
Long interspersed nuclear elements	626.33	631.64	+0.8
Short interspersed nuclear elements	386.48	390.27	+1.0
Long terminal repeats	267.52	269.91	+0.9
Satellite	76.51	150.42	+96.6
DNA	108.53	109.35	+0.8
Simple repeat	36.5	77.69	+112.9
Low complexity	6.16	6.44	+4.6
Retroposon	4.51	4.65	+3.3
rRNA	0.21	1.71	+730.4

circular consensus sequencing (HiFi) (16, 20), 120× Oxford Nanopore ultralong-read sequencing (ONT) (14, 21), 100× Illumina PCR-Free sequencing (ILMN) (1), 70× Illumina Arima Genomics Hi-C (Hi-C) (14), BioNano

optical maps (14), and single-cell DNA template strand sequencing (Strand-seq) (20) (table S1). To enable assembly of the highly repetitive centromeric satellite arrays and closely related segmental duplications, we developed methods

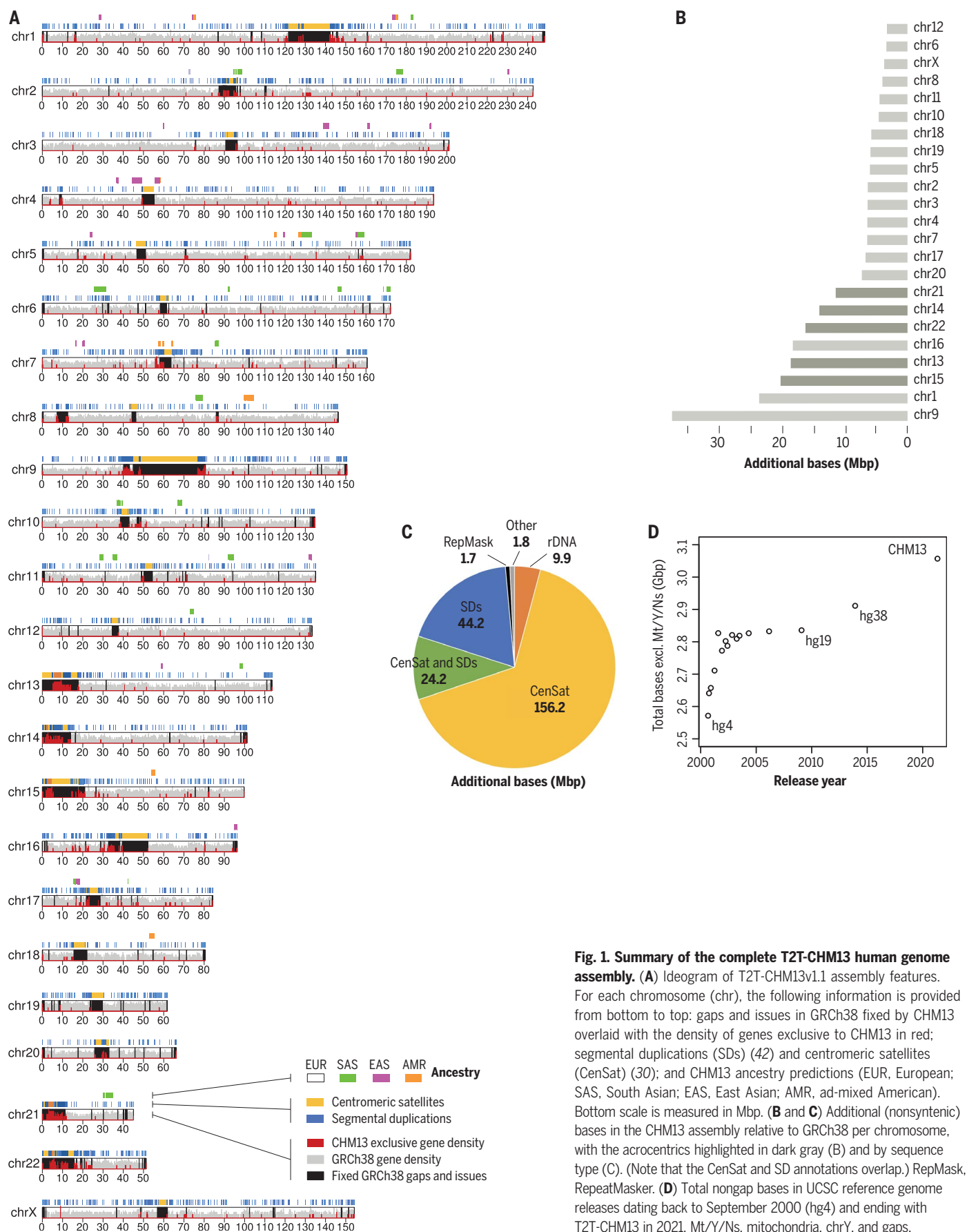


Fig. 1. Summary of the complete T2T-CHM13 human genome assembly. (A) Ideogram of T2T-CHM13v1.1 assembly features.

For each chromosome (chr), the following information is provided from bottom to top: gaps and issues in GRCh38 fixed by CHM13 overlaid with the density of genes exclusive to CHM13 in red; segmental duplications (SDs) (42) and centromeric satellites (CenSat) (30); and CHM13 ancestry predictions (EUR, European; SAS, South Asian; EAS, East Asian; AMR, admixed American). Bottom scale is measured in Mbp. (B and C) Additional (nonsynthetic) bases in the CHM13 assembly relative to GRCh38 per chromosome, with the acrocentrics highlighted in dark gray (B) and by sequence type (C). (Note that the CenSat and SD annotations overlap.) RepMask, RepeatMasker. (D) Total nongap bases in UCSC reference genome releases dating back to September 2000 (hg4) and ending with T2T-CHM13 in 2021. Mt/Y/Ns, mitochondria, chrY, and gaps.

for assembly, polishing, and validation that better utilize these available datasets.

Genome assembly

The basis of the T2T-CHM13 assembly is a high-resolution assembly string graph (26) built directly from HiFi reads. In a bidirected string graph, nodes represent unambiguously assembled sequences, and edges correspond to the overlaps between them, owing to either repeats or true adjacencies in the underlying genome. The CHM13 graph was constructed using a purpose-built method that combines components from existing assemblers (16, 27) along with specialized graph processing (23). Most HiFi errors are small insertions or deletions within homopolymer runs and simple sequence repeats (16), so homopolymer runs were first “compressed” to a single nucleotide (e.g., $A_1 \dots A_n$ becomes A_1 for $n > 1$). All compressed reads were then aligned to one another to identify and correct small errors, and differences within simple sequence repeats were masked. After compression, correction, and masking, only exact read overlaps were considered during graph construction, followed by iterative graph simplification (23).

In the resulting graph, most components originate from a single chromosome and have an almost linear structure (Fig. 2A), which suggests that few perfect repeats greater than roughly 10 kbp exist between different chromosomes or distant loci. Two notable exceptions are the five acrocentric chromosomes, which form a single connected component in the graph, and a recent multimegabase HSat3 duplication on chromosome 9, consistent with the 9qh+ karyotype of CHM13 (fig. S3). Minor fragmentation of the chromosomes into multiple components resulted from a lack of HiFi sequencing coverage across GA-rich sequences (16). These gaps were later filled with a prior ONT-based assembly (CHM13v0.7) (14).

Ideally, the complete sequence for each chromosome should exist as a walk through the string graph where some nodes may be traversed multiple times (repeats) and some not at all (errors and heterozygous variants). To help identify the correct walks, we estimated coverage depth and multiplicity of the nodes (23), which allowed most tangles to be manually resolved as unique walks visiting each node the appropriate number of times (Fig. 2B and fig. S4). In the remaining cases, the correct path was ambiguous and required integration of ONT reads (Fig. 2, C and D). Where possible, ONT reads were aligned to candidate traversals or directly to the HiFi graph (28) to guide the correct walk (fig. S5), but more elaborate strategies were required for recent satellite array duplications on chromosomes 6 and 9 (23). Only the five rDNA arrays, constituting about 10 Mbp of sequence, could not be re-

solved with the string graph and required a specialized approach (described later). An accurate consensus sequence for the selected graph walks was computed from the uncompressed HiFi reads (23), resulting in the CHM13v0.9 draft assembly.

For comparative genomics of the centromere (29, 30), we repeated this process on an additional X chromosome from the Coriell GM24385 cell line [National Institute of Standards and Technology (NIST) ID: HG002]. The resulting T2T-HG002-ChrX assembly shows comparable accuracy to T2T-CHM13 (23) (figs. S6 to S8).

rDNA assembly

The most complex region of the CHM13 string graph involves the human rDNA arrays and their surrounding sequence (Fig. 2D). Human rDNAs are 45-kbp near-identical repeats that encode the 45S rRNA and are arranged in large, tandem repeat arrays embedded within the short arms of the acrocentric chromosomes. The length of these arrays varies between individuals (31) and even somatically, especially with aging and certain cancers (32). A typical diploid human genome has an average of 315 rDNA copies, with a standard deviation of 104 copies (31). We estimate that the diploid CHM13 genome contains about 400 rDNA copies based on ILMN depth of coverage (23) (fig. S9) or 409 ± 9 (mean $\pm$ SD) rDNA copies by droplet digital polymerase chain reaction (ddPCR) (fig. S10).

To assemble these highly dynamic regions of the genome and overcome limitations of the string graph construction (23) (fig. S11), we constructed sparse de Bruijn graphs for each of the five rDNA arrays (33) (fig. S12). ONT reads were aligned to the graphs to identify a set of walks, which were converted to sequence, segmented into individual rDNA units, and clustered into “morphs” according to their sequence similarity. The copy number of each morph was estimated from the number of supporting ONT reads, and consensus sequences were polished with mapped HiFi reads. ONT reads spanning two or more rDNA units were used to build a morph graph representing the structure of each array (fig. S12).

The shorter arrays on chromosomes 14 and 22 consist of a single primary morph arranged in a head-to-tail array, whereas the longer arrays on chromosomes 13, 15, and 21 exhibit a more mosaic structure involving multiple, interspersed morphs. In these cases, the ONT reads were not long enough to fully resolve the ordering, and the primary morphs were artificially arranged in consecutive blocks reflecting their estimated copy number. These three arrays capture the chromosome-specific morphs but should be treated as model sequences. The final T2T-CHM13 assembly contains 219

complete rDNA copies, totaling 9.9 Mbp of sequence.

Assembly validation and polishing

To evaluate concordance between the reads and the assembly, we mapped all available primary data—including HiFi, ONT, ILMN, Strand-seq, and Hi-C—to the CHM13v0.9 draft assembly to identify both small and structural variants [see (34) for a complete description]. Manual curation corrected 4 large and 993 small errors, resulting in the CHM13v1.0 assembly, and identified 44 large and 3901 small heterozygous variants (34). Further telomere polishing and addition of the rDNA arrays (23) resulted in a complete, telomere-to-telomere assembly of a human genome, T2T-CHM13v1.1.

The T2T-CHM13 assembly is consistent with previously validated assemblies of chromosomes X (14) and 8 (21), and the sizes of assembled satellite arrays match ddPCR copy-number estimates for those tested (fig. S10 and tables S2 and S3). Mapped Strand-seq (figs. S13 and S14) and Hi-C (fig. S15) data show no signs of misorientations or other large-scale structural errors. The assembly correctly resolves 644 of 647 previously sequenced CHM13 BACs at $>99.99\%$ identity, with the three others reflecting errors in the BACs themselves (figs. S16 to S19).

Mapped sequencing read depth shows uniform coverage across all chromosomes (Fig. 3A), with 99.86% of the assembly within three standard deviations of the mean coverage for either HiFi or ONT (HiFi coverage 34.70 ± 7.03 and ONT coverage 116.16 ± 16.96 , excluding the mitochondrial genome). Ignoring the 10 Mbp of rDNA sequence, where most of the coverage deviation resides, 99.99% of the assembly is within three standard deviations (23). Alignment-free analysis of ILMN and HiFi copy-number data also shows concordance with the assembly (figs. S20 and S21). This is consistent with uniform coverage of the genome and confirms both the accuracy of the assembly and the absence of aneuploidy in the sequenced CHM13 cells.

Coverage increases or decreases were observed across multiple satellite arrays (Fig. 3, B to D). However, given the uniformity of coverage across these arrays, association with specific satellite classes, and the sometimes opposite effect observed for HiFi and ONT, we hypothesize that these anomalies are related to biases introduced during sample preparation, sequencing, or base calling, rather than assembly error (23) (figs. S22 to S26 and table S4). Although the specific mechanisms require further investigation, prior studies have noted similar biases within certain satellite arrays and sequence contexts for both ONT and HiFi (35, 36).

Because they are the most difficult regions of the genome to assemble, we performed targeted validation of long tandem repeats to

identify any errors missed by the genome-wide approach. The assembled rDNA morphs, being only 45 kbp each, were manually validated by inspection of the read alignments used for polishing. Alpha satellite higher-order repeats (HORs) were validated using a purpose-built method (37) (fig. S27 and table S5) and compared with independent ILMN-based HOR copy-number estimates (fig. S28). All centromeric satellite arrays, including beta satellite

(BSat) and HSat repeats, were further validated by measuring the ratio of primary to secondary variants identified by HiFi reads (38) (fig. S29).

The consensus accuracy of the T2T-CHM13 assembly is estimated to be about one error per 10 Mbp (23, 34), which exceeds the historical standard of “finished” sequence by orders of magnitude. However, regions of low HiFi coverage were found to be associated

with an enrichment of potential errors, as estimated from both HiFi and ILMN data (34). To guide future use of the assembly, we have cataloged all low-coverage, low-confidence, and known heterozygous sites identified by the above validation procedures (34). The total number of bases covered by potential issues in the T2T-CHM13 assembly is just 0.3% of the total assembly length compared with 8% for GRCh38 (Fig. 3A).

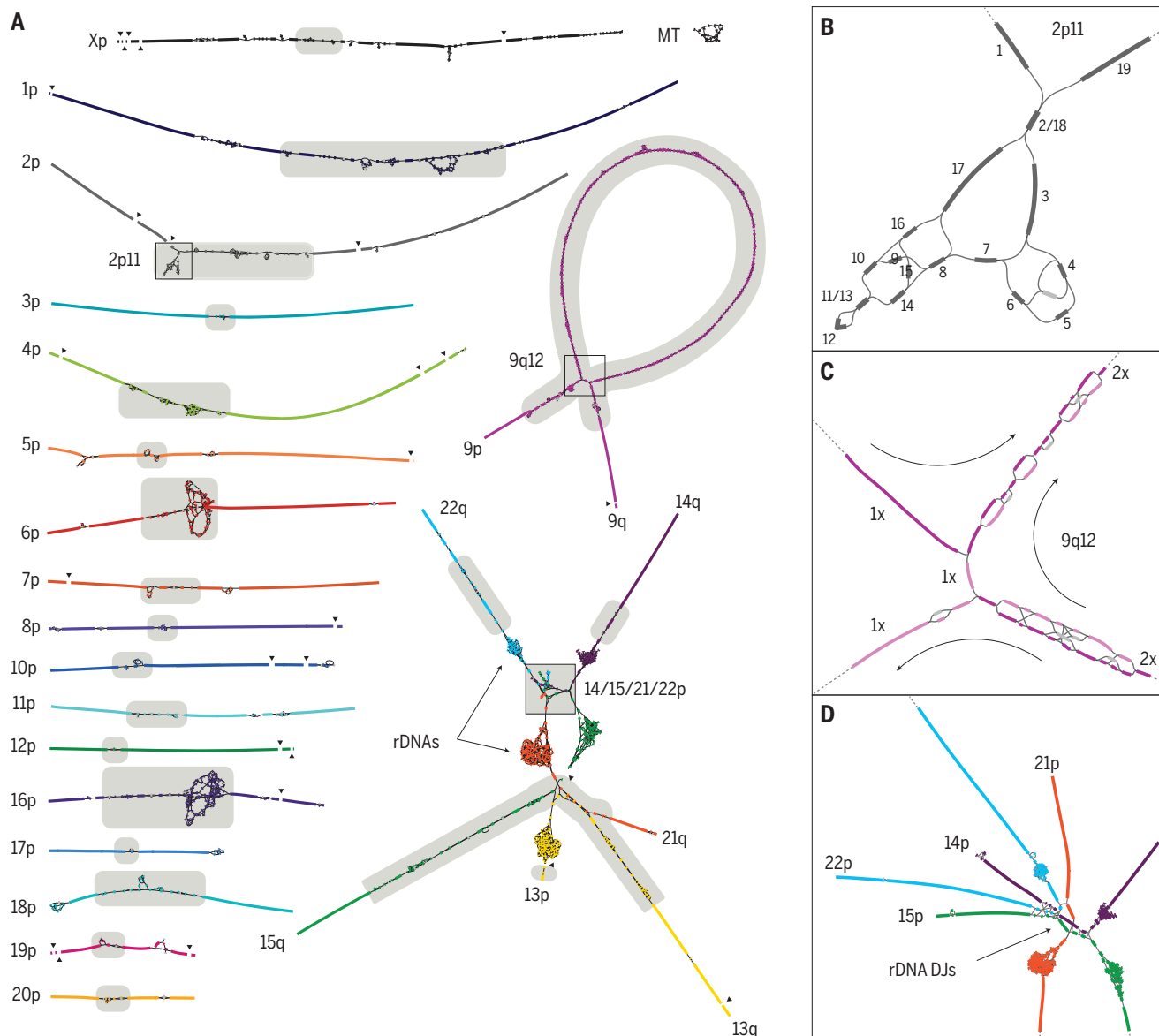


Fig. 2. High-resolution assembly string graph of the CHM13 genome.

(A) Bandage (60) visualization, where nodes represent unambiguously assembled sequences scaled by length and edges correspond to the overlaps between node sequences. Each chromosome is both colored and numbered on the short (p) arm. Long (q) arms are labeled where unclear. The five acrocentric chromosomes (bottom right) are connected owing to similarity between their short arms, and the rDNA arrays form five dense tangles because of their high copy number. The graph is partially fragmented because of HiFi coverage dropout surrounding GA-rich sequence (black triangles). Centromeric satellites (30) are the source of most ambiguity in the graph (gray highlights). MT, mitochondria.

(B) The ONT-assisted graph traversal for the 2p11 locus is given by numerical order. Based on low depth of coverage, the unlabeled light gray node represents an artifact or heterozygous variant and was not used. (C) The multimegabase tandem HSat3 duplication (9qh+) at 9q12 requires two traversals of the large loop structure. (The size of the loop is exaggerated because graph edges are of constant size.) Nodes used by the first traversal are in dark purple, and nodes used by the second traversal are in light purple. Nodes used by both traversals typically have twice the sequencing coverage. (D) Enlargement of the distal short arms of the acrocentrics, showing the colored graph walks and edges between highly similar sequences in the distal junctions (DJs) adjacent to the rDNA arrays.

A truly complete genome

T2T-CHM13 includes gapless telomere-to-telomere assemblies for all 22 human autosomes and chromosome X, comprising 3,054,815,472 bp of nuclear DNA, plus a 16,569-bp mitochondrial genome. This complete assembly adds or corrects 238 Mbp of sequence that does not colinearly align to GRCh38 over a 1-Mbp interval (i.e., is non-synthetic), primarily comprising centromeric satellites (76%), nonsatellite segmental duplications (19%), and rDNAs (4%) (Fig. 1C). Of this, 182 Mbp of sequence has no primary alignments to GRCh38 and is exclusive to T2T-CHM13. As a result, T2T-CHM13 increases the number of known genes and repeats in the human genome (Table 1).

To provide an initial annotation, we used both the Comparative Annotation Toolkit (CAT) (39) and Liftoff (40) to project the GENCODE v35 (41) reference annotation onto the T2T-CHM13 assembly. Additionally, CHM13 full-length isoform sequencing (Iso-seq) transcriptome reads were assembled into transcripts and provided as complemen-

tary input to CAT. A comprehensive annotation was built by combining the CAT annotation with genes identified only by Liftoff (23).

The draft T2T-CHM13 annotation totals 63,494 genes and 233,615 transcripts, of which 19,969 genes (86,245 transcripts) are predicted to be protein coding, with 683 predicted frameshifts in 385 genes (469 transcripts) (Table 1, fig. S30, and tables S6 to S8). Only 263 GENCODE genes (448 transcripts) are exclusive to GRCh38 and have no assigned ortholog in the CHM13 annotation (tables S9 and S10). Of these, 194 are due to a lower copy number in the CHM13 annotation (fig. S31), 46 do not align well to CHM13, and 23 correspond to known false duplications in GRCh38 (25) (fig. S32). Most of these genes are noncoding and associated with repetitive elements. Only four are annotated as being medically relevant (*CFHR1*, *CFHR3*, *OR51A2*, *UGT2B28*), all of which are absent owing to a copy number difference, and the only protein-coding genes that align poorly are immunoglobulin and T cell receptor genes, which are known to be highly diverse.

In comparison, a total of 3604 genes (6693 transcripts) are exclusive to CHM13 (tables S11 and S12). Most of these genes represent putative paralogs and localize to pericentromeric regions and the short arms of the acrocentrics, including 876 rRNA transcripts. Only 48 of the CHM13-exclusive genes (56 transcripts) were predicted solely from de novo assembled transcripts. Of all genes exclusive to CHM13, 140 are predicted to be protein coding based on their GENCODE paralogs and have a mean of 99.5% nucleotide and 98.7% amino acid identity to their most similar GRCh38 copy (table S13). Although some of these additional paralogs may be present (but unannotated) in GRCh38 (23), 1956 of the genes exclusive to CHM13 (99 protein coding) are in regions with no primary alignment to GRCh38 (table S11). A broader set of 182 multi-exon protein-coding genes fall within nonsynthetic regions, 36% of which were confirmed to be expressed in CHM13 (42).

Compared with GRCh38, T2T-CHM13 is a more complete, accurate, and representative reference for both short- and long-read variant calling across human samples of all ancestries

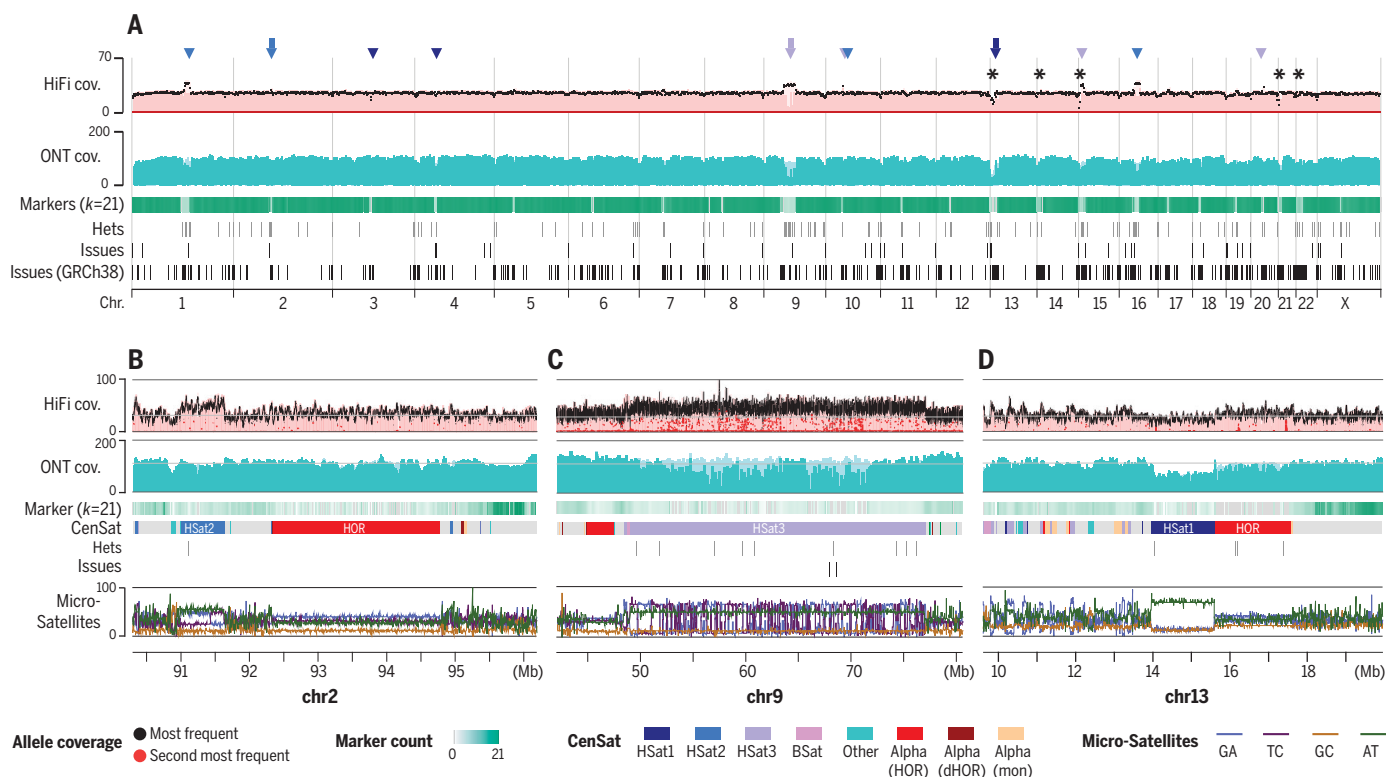


Fig. 3. Sequencing coverage and assembly validation. (A) Uniform whole-genome coverage of mapped HiFi and ONT reads is shown with primary alignments in light shades and marker-assisted alignments overlaid in dark shades. Large HSat arrays (30) are noted by triangles, with inset regions marked by arrowheads and the location of the rDNA arrays marked with asterisks. Regions with low unique marker frequency (light green) correspond to drops in unique marker density but are recovered by the lower-confidence primary alignments. Annotated assembly issues are compared for T2T-CHM13 and

GRCh38. Hets, heterozygous variants; k, marker size. (B to D) Enlargements corresponding to regions of the genome featured in Fig. 2, B to D, respectively. Uniform coverage changes within certain satellites are reproducible and likely caused by sequencing bias. Identified heterozygous variants and assembly issues are marked below and typically correspond with low coverage of the primary allele (black) and increased coverage of the secondary allele (red). The percentage of microsatellite repeats for every 128-bp window is shown at the bottom. dHOR, divergent HOR; mon, monomeric.

(25). Reanalysis of 3202 short-read datasets from the 1KGP showed that T2T-CHM13 simultaneously reduces both false-negative and false-positive variant calls because of the addition of 182 Mbp of missing sequence and the exclusion of 1.2 Mbp of falsely duplicated sequence in GRCh38. These improvements, combined with a lower frequency of rare variants and errors in T2T-CHM13, eliminate tens of thousands of spurious variants per 1KGP sample (25). In addition, the T2T-CHM13 reference was found to be more representative of human copy-number variation than GRCh38 when compared against 268 human genomes from the Simons Genome Diversity Project (SGDP) (42, 43). Specifically, within non-syntenic segmentally duplicated regions of

the genome, T2T-CHM13 is nine times more predictive of SGDP copy number than GRCh38 (42). These results underscore both the quality of the assembly and the genomic stability of the cell line from which it was derived.

Acrocentric chromosomes

T2T-CHM13 uncovers the genomic structure of the short arms of the five acrocentric chromosomes, which, despite their importance for cellular function (44), have remained largely unsequenced to date. This omission has been due to their enrichment for satellite repeats and segmental duplications, which has prohibited sequence assembly and limited their characterization to cytogenetics, restriction mapping, and BAC sequencing

(45–47). All five of CHM13's short arms follow a similar structure consisting of an rDNA array embedded within distal and proximal repeat arrays (Fig. 4). From telomere to centromere, the short arms vary in size from 10.1 Mbp (chromosome 14) to 16.7 Mbp (chromosome 15), with a combined length of 66.1 Mbp.

Compared with other human chromosomes, the short arms of the acrocentrics are unusually similar to one another. Specifically, we find that 5-kbp windows align with a median identity of 98.7% between the short arms, creating many opportunities for interchromosomal exchange (Fig. 4). This high degree of similarity is presumably due to recent non-allelic or ectopic recombination stemming from their colocalization in the nucleolus

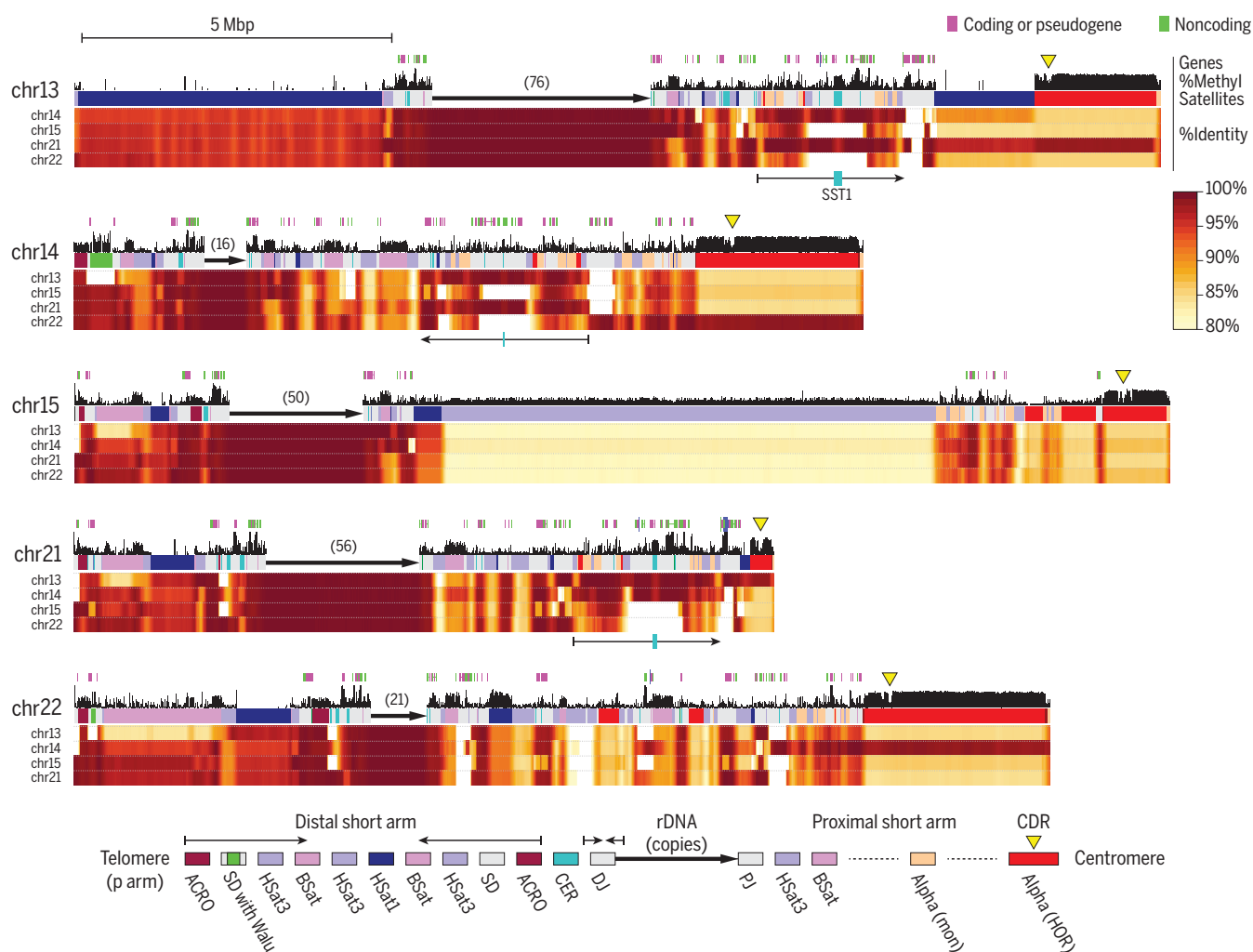


Fig. 4. Short arms of the acrocentric chromosomes. Each short arm is shown along with annotated genes, percentage of methylated CpGs (29), and a color-coded satellite repeat annotation (30). The rDNA arrays are represented by a directional arrow and copy number because of their high self-similarity, which prohibits ONT mapping. Percent identity heatmaps versus the other four arms were computed in 10-kbp windows and smoothed over 100-kbp intervals. Each position shows the maximum identity of that window to any window in the other chromosome. The distal short arms include conserved satellite structure and

inverted repeats (thin arrows), whereas the proximal short arms show a diversity of structures. The proximal short arms of chromosomes 13, 14, and 21 share a segmentally duplicated core, including small alpha satellite HOR arrays and a central, highly methylated SST1 array (thin arrows with teal block). Yellow triangles indicate hypomethylated centromeric dip regions (CDRs), marking the sites of kinetochore assembly (29). Numbers in parentheses indicate rDNA copy number. ACRO, acrocentric repeat; CER, centromeric repeat; DJ, distal junction; PJ, proximal junction; SD, segmental duplication.

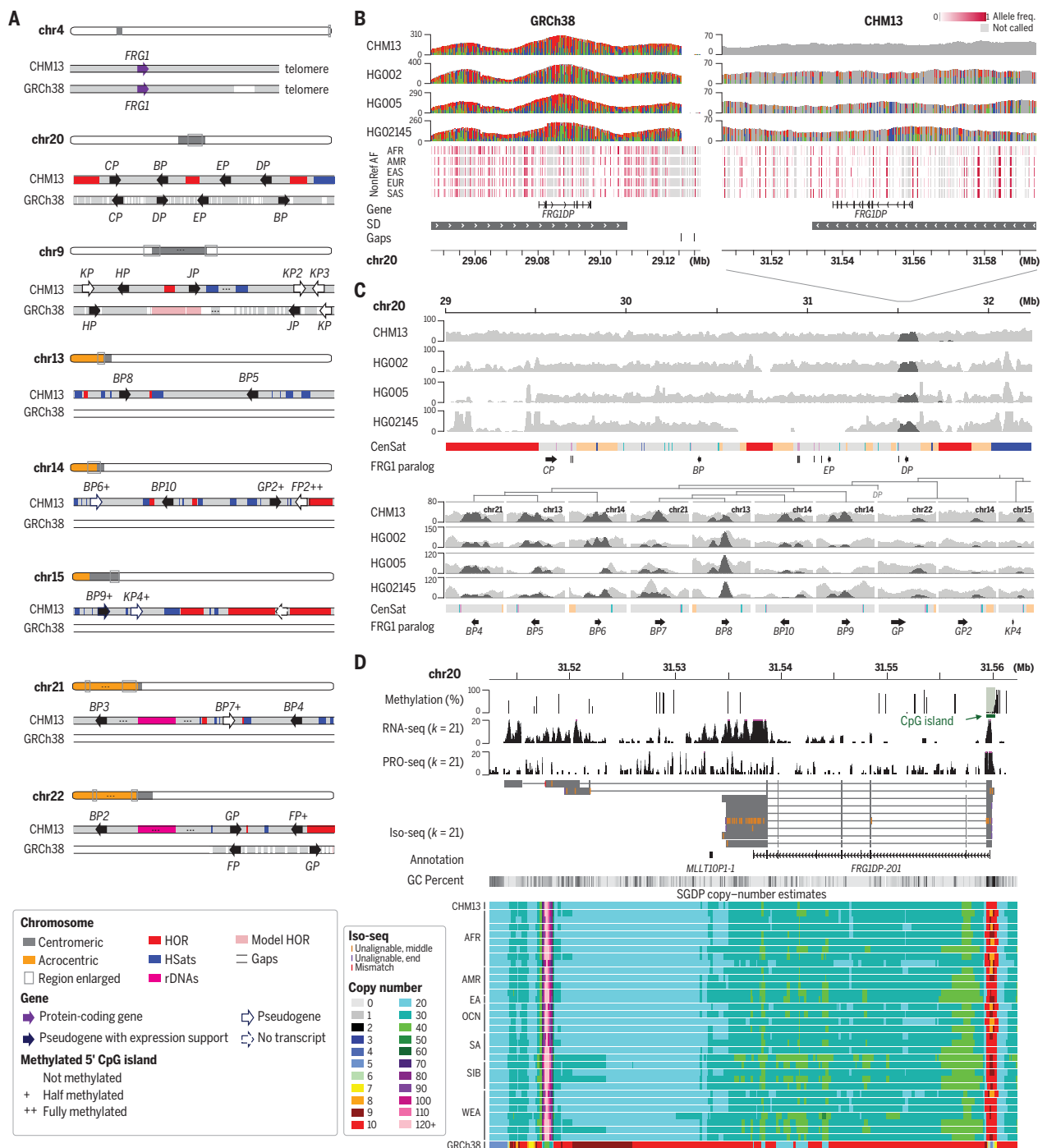


Fig. 5. Resolved *FRG1* paralogs. (A) Protein-coding gene *FRG1* and its 23 paralogs in CHM13. Only nine are found in GRCh38. Genes are drawn larger than their actual size, and the “*FRG1*” prefix is omitted for brevity. All paralogs are found near satellite arrays. Most copies exhibit evidence of expression, including CpG islands present at the 5' start site with varying degrees of methylation. (B) Reference (gray) and variant (colored) allele coverage is shown for four human HiFi samples mapped to the paralog *FRG1DP*. When mapped to GRCh38, the region shows excessive HiFi coverage and variants, indicating that reads from the missing paralogs are mismapped to *FRG1DP* (variants >80% frequency shown). When mapped to CHM13, HiFi reads show the expected coverage and a typical heterozygous variation pattern for the three non-CHM13 samples (variants >20% frequency shown). These nonreference alleles are also found in other populations from 1KGP ILMN data. NonRef AF, nonreference allele frequency; AFR,

African; AMR, ad-mixed American; EAS, East Asian; EUR, European; SAS, South Asian. (C) Mapped HiFi read coverage for other *FRG1* paralogs, with an extended context shown for chromosome 20. Coverage of HiFi reads that mapped to *FRG1DP* in GRCh38 is highlighted (dark gray), showing the paralogous copies they originate from (*FRG1BP4* to *FRG1BP10*, *FRG1GP*, *FRG1GP2*, and *FRG1KP4*). Background coverage is variable for some paralogs, suggesting the presence of copy-number polymorphism in the population. (D) Methylation and expression profiles suggest transcription of *FRG1DP* in CHM13. In the copy-number display (bottom), 100-bp windows from the CHM13 assembly are highlighted with a color representing the estimated copy number of that sequence in an SGDP sample. The CHM13 and GRCh38 tracks show the estimated copy number of these same sequences in the respective assemblies. CHM13 copy number resembles all samples from the SGDP, whereas GRCh38 underrepresents the true copy number.

(46). Additionally, considering an 80% identity threshold, no 5-kbp window on the short arms is unique, and 96% of the non-rDNA sequence can be found elsewhere in the genome, suggesting that the acrocentrics are dynamic sources of segmental duplication.

The rDNA arrays of CHM13 vary in size from 0.7 Mbp (chromosome 14) to 3.6 Mbp (chromosome 13) and are in the expected arrangement, organized as head-to-tail tandem arrays with all 45S transcriptional units pointing toward the centromere. No inversions were noted within the arrays, and nearly all rDNA units are full length, in contrast to some prior studies that reported embedded inversions and other noncanonical structures (47, 48). Each array appears highly homogenized, and there is more variation between rDNA units on different chromosomes than within chromosomes (fig. S33), suggesting that intrachromosomal exchange of rDNA units through nonallelic homologous recombination is more common than interchromosomal exchange.

Many 45S gene copies on the same chromosome are identical to one another, whereas the identity of the most frequent 45S morphs between chromosomes ranges from 99.4 to 99.7%. A chromosome 15 rDNA morph shows the highest identity (98.9%) to the current KY962518.1 rDNA reference sequence, originally derived from a human chromosome 21 BAC clone (47). As expected, the 13-kbp 45S is more conserved than the intergenic spacer, with all major 45S morphs aligning between 99.4 and 99.6% identity to KY962518.1. Certain rDNA variants appear to be chromosome specific, including single-nucleotide variants within the 45S and its upstream promoter region (fig. S34). The most evident variants are repeat expansions and contractions within the tandem “R” repeat that immediately follows the 45S and the CT-rich “long” repeat located in the middle of the intergenic spacer. The most frequent morph in each array can be specifically distinguished by these two features (fig. S35).

From the telomere to the rDNA array, the structure of all five distal short arms follows a similar pattern that involves a symmetric arrangement of inverted segmental duplications and acrocentric, HSat3, BSat, and HSat1 repeats (Fig. 4); however, the sizes of these repeat arrays vary among chromosomes. Chromosome 13 is missing the distal half of the inverted duplication and has an expanded HSat1 array relative to the others. Despite their variability in size, all satellite arrays share a high degree of similarity (typically >90% identity) both within and between acrocentric chromosomes. Chromosomes 14 and 22 also feature the expansion of a 64-bp Alu-associated satellite repeat (“Walu”) within the distal inverted duplication (49), the loca-

tion of which was confirmed by fluorescence in situ hybridization (FISH) (fig. S36). The distal junction immediately before the rDNA array includes centromeric repeats and a highly conserved and actively transcribed 200-kbp palindromic repeat, which agrees with previous characterizations of the rDNA flanking sequences (46, 50).

Extending from the rDNA array to the centromere, the proximal short arms are larger in size and show a higher diversity of structures, including shuffled segmental duplications (42), composite transposable element arrays (49), satellite arrays (including HSat3, BSat, HSat1, and HSat5), and alpha satellite arrays (both monomeric and HORs) (30). Some proximal BSat arrays show a mosaic inversion structure that was also observed in HSat arrays elsewhere in the genome (30) (fig. S37). The proximal short arms of chromosomes 13, 14, and 21 appear to share the highest degree of similarity with a large region of segmental duplication, including similar HOR subsets and a central and highly methylated SST1 array (Fig. 4). This coincides with these three chromosomes being most frequently involved in Robertsonian translocations (51). Alpha satellite HORs on chromosomes 13 and 21 and chromosomes 14 and 22 also share high similarity within each pair, but not between them (52, 53). Nonsatellite sequences within these segmental duplications often exceed 99% identity and show evidence of transcription (29, 42, 49). Using the T2T-CHM13 reference as a basis, further study of additional genomes is now needed to understand which of these features are conserved across the human population.

Analyses and resources

A number of companion studies were carried out to characterize the complete sequence of a human genome, including comprehensive analyses of centromeric satellites (30), segmental duplications (42), transcriptional (49) and epigenetic profiles (29), mobile elements (49), and variant calls (25). Up to 99% of the complete CHM13 genome can be confidently mapped with long-read sequencing, opening these regions of the genome to functional and variational analysis (23) (fig. S38 and table S14). We have produced a rich collection of annotations and omics datasets for CHM13—including RNA sequencing (RNA-seq) (30), Iso-seq (21), precision run-on sequencing (PRO-seq) (49), cleavage under targets and release using nuclease (CUT&RUN) (30), and ONT methylation (29) experiments—and have made these datasets available via a centralized University of California, Santa Cruz (UCSC), Assembly Hub genome browser (54).

To highlight the utility of these genetic and epigenetic resources mapped to a complete human genome, we provide the example of

a segmentally duplicated region of the chromosome 4q subtelomere that is associated with facioscapulohumeral muscular dystrophy (FSHD) (55). This region includes FSHD region gene 1 (*FRG1*), FSHD region gene 2 (*FRG2*), and an intervening D4Z4 macrosatellite repeat containing the double homeobox 4 (*DUX4*) gene that has been implicated in the etiology of FSHD (56). Numerous duplications of this region throughout the genome have complicated past genetic analyses of FSHD.

The T2T-CHM13 assembly reveals 23 paralogs of *FRG1* spread across all acrocentric chromosomes as well as chromosomes 9 and 20 (Fig. 5A). This gene appears to have undergone recent amplification in the great apes (57), and approximate locations of *FRG1* paralogs were previously identified by FISH (58). However, only nine *FRG1* paralogs are found in GRCh38, hampering sequence-based analysis.

One of the few *FRG1* paralogs included in GRCh38, *FRG1DP*, is located in the centromeric region of chromosome 20 and shares high identity (97%) with several paralogs (*FRG1BP4* to *FRG1BP10*) (23) (fig. S39 and tables S15 and S16). When mapping HiFi reads, the absence of the additional *FRG1* paralogs in GRCh38 causes their reads to incorrectly align to *FRG1DP*, resulting in many false-positive variants (Fig. 5B). Most *FRG1* paralogs appear present in other human genomes (Fig. 5C), and all except *FRG1KP2* and *FRG1KP3* have upstream CpG islands and some degree of expression evidence in CHM13 (Fig. 5D and table S17). Any variants within these paralogs, and others like them, will be overlooked when using GRCh38 as a reference.

Future of the human reference genome

The T2T-CHM13 assembly adds five full chromosome arms and more additional sequence than any genome reference release in the past 20 years (Fig. 1D). This 8% of the genome has not been overlooked because of a lack of importance but rather because of technological limitations. High-accuracy long-read sequencing has finally removed this technological barrier, enabling comprehensive studies of genomic variation across the entire human genome, which we expect to drive future discovery in human genomic health and disease. Such studies will necessarily require a complete and accurate human reference genome.

CHM13 lacks a Y chromosome, and homozygous Y-bearing CHMs are nonviable, so a different sample type will be required to complete this last remaining chromosome. However, given its haploid nature, it should be possible to assemble the Y chromosome from a male sample using the same methods described here and supplement the T2T-CHM13 reference assembly with a Y chromosome as needed.

Extending beyond the human reference genome, large-scale resequencing projects have

revealed genomic variation across human populations. Our reanalyses of the 1KGP (25) and SGDP (42) datasets have already shown the advantages of T2T-CHM13, even for short-read analyses. However, these studies give only a glimpse of the extensive structural variation that lies within the most repetitive regions of the genome assembled here. Long-read resequencing studies are now needed to comprehensively survey polymorphic variation and reveal any phenotypic associations within these regions.

Although CHM13 represents a complete human haplotype, it does not capture the full diversity of human genetic variation. To address this bias, the Human Pangenome Reference Consortium (59) has joined with the T2T Consortium to build a collection of high-quality reference haplotypes from a diverse set of samples. Ideally, all genomes could be assembled at the quality achieved here, but automated T2T assembly of diploid genomes presents a difficult challenge that will require continued development. Until this goal is realized, and any human genome can be completely sequenced without error, the T2T-CHM13 assembly represents a more complete, representative, and accurate reference than GRCh38.

REFERENCES AND NOTES

1. V. A. Schneider *et al.*, *Genome Res.* **27**, 849–864 (2017).
2. International Human Genome Sequencing Consortium, *Nature* **409**, 860–921 (2001).
3. J. C. Venter *et al.*, *Science* **291**, 1304–1351 (2001).
4. E. W. Myers *et al.*, *Science* **287**, 2196–2204 (2000).
5. E. E. Eichler, R. A. Clark, X. She, *Nat. Rev. Genet.* **5**, 345–354 (2004).
6. K. H. Miga *et al.*, *Genome Res.* **24**, 697–707 (2014).
7. M. Gupta, A. R. Dhanasekaran, K. J. Gardiner, *Mamm. Genome* **27**, 538–555 (2016).
8. M. J. P. Chaisson *et al.*, *Nature* **517**, 608–611 (2015).
9. International Human Genome Sequencing Consortium, *Nature* **431**, 931–945 (2004).
10. J. Eid *et al.*, *Science* **323**, 133–138 (2009).
11. K. Berlin *et al.*, *Nat. Biotechnol.* **33**, 623–630 (2015).
12. M. Jain *et al.*, *Nat. Biotechnol.* **36**, 338–345 (2018).
13. M. Jain *et al.*, *Nat. Biotechnol.* **36**, 321–323 (2018).
14. K. H. Miga *et al.*, *Nature* **585**, 79–84 (2020).
15. A. M. Wenger *et al.*, *Nat. Biotechnol.* **37**, 1155–1162 (2019).
16. S. Nurk *et al.*, *Genome Res.* **30**, 1291–1305 (2020).
17. J. Huddleston *et al.*, *Genome Res.* **27**, 677–685 (2017).
18. E. E. Eichler, U. Surti, R. Ophoff, Proposal for construction a human haploid BAC library from hydatidiform mole source material (2002); www.genome.gov/Pages/Research/Sequencing/BACLibrary/HydatidiformMoleBAC021203.pdf.
19. K. M. Steinberg *et al.*, *Genome Res.* **24**, 2066–2076 (2014).
20. M. R. Vollger *et al.*, *Ann. Hum. Genet.* **84**, 125–140 (2020).
21. G. A. Logsdon *et al.*, *Nature* **593**, 101–107 (2021).
22. J.-B. Fan *et al.*, *Genomics* **79**, 58–62 (2002).
23. See supplementary materials.
24. 1000 Genomes Project Consortium, *Nature* **526**, 68–74 (2015).
25. S. Aganezov *et al.*, *Science* **376**, eabl3533 (2022).
26. E. W. Myers, *Bioinformatics* **21**, ii79–ii85 (2005).
27. H. Li, *Bioinformatics* **32**, 2103–2110 (2016).
28. M. Rautiainen, T. Marschall, *Genome Biol.* **21**, 253 (2020).
29. A. Gershman *et al.*, *Science* **376**, eabj5089 (2022).
30. N. Altemose *et al.*, *Science* **376**, eabl4178 (2022).
31. M. M. Parks *et al.*, *Sci. Adv.* **4**, eaao0665 (2018).
32. J. O. Nelson, G. J. Watase, N. Warsinger-Pepe, Y. M. Yamashita, *Trends Genet.* **35**, 734–742 (2019).
33. M. Rautiainen, T. Marschall, *Bioinformatics* **37**, 2476–2478 (2021).
34. A. M. McCartney *et al.*, *Nat. Methods* <https://doi.org/10.1038/s41592-022-01440-3> (2022).
35. J. M. Flynn, M. Long, R. A. Wing, A. G. Clark, *Mol. Biol. Evol.* **37**, 1362–1375 (2020).
36. W. M. Guiblet *et al.*, *Genome Res.* **28**, 1767–1778 (2018).
37. A. Mikheenko, A. V. Bzikadze, A. Gurevich, K. H. Miga, P. A. Pevzner, *Bioinformatics* **36**, i75–i83 (2020).
38. M. R. Vollger *et al.*, *Nat. Methods* **16**, 88–94 (2019).
39. I. T. Fiddes *et al.*, *Genome Res.* **28**, 1029–1038 (2018).
40. A. Shumate, S. L. Salzberg, *Bioinformatics* **37**, 1639–1643 (2021).
41. A. Frankish *et al.*, *Nucleic Acids Res.* **47**, D766–D773 (2019).
42. M. R. Vollger *et al.*, *Science* **376**, eabj6965 (2022).
43. S. Mallick *et al.*, *Nature* **538**, 201–206 (2016).
44. M. S. Lindström *et al.*, *Oncogene* **37**, 2351–2366 (2018).
45. R. Lyle *et al.*, *Genome Res.* **17**, 1690–1696 (2007).
46. I. Floutsakou *et al.*, *Genome Res.* **23**, 2003–2012 (2013).
47. J.-H. Kim *et al.*, *Nucleic Acids Res.* **46**, 6712–6725 (2018).
48. S. Caburet *et al.*, *Genome Res.* **15**, 1079–1085 (2005).
49. S. J. Hoyt *et al.*, *Science* **376**, eabk3112 (2022).
50. M. van Sluis *et al.*, *Genes Dev.* **33**, 1688–1701 (2019).
51. B. A. Sullivan, L. S. Jenkins, E. M. Karson, J. Leana-Cox, S. Schwartz, *Am. J. Hum. Genet.* **59**, 167–175 (1996).
52. G. M. Greig, P. E. Warburton, H. F. Willard, *J. Mol. Evol.* **37**, 464–475 (1993).
53. A. L. Jørgensen, S. Kølvrå, C. Jones, A. L. Bak, *Genomics* **3**, 100–109 (1988).
54. W. J. Kent *et al.*, *Genome Res.* **12**, 996–1006 (2002).
55. C. Wijmenga *et al.*, *Nat. Genet.* **2**, 26–30 (1992).
56. R. J. L. F. Lemmers *et al.*, *Science* **329**, 1650–1653 (2010).
57. P. K. Grewal, M. van Geel, R. R. Frants, P. de Jong, J. E. Hewitt, *Gene* **227**, 79–88 (1999).
58. J. C. van Deutekom *et al.*, *Hum. Mol. Genet.* **5**, 581–590 (1996).
59. K. H. Miga, T. Wang, *Annu. Rev. Genomics Hum. Genet.* **22**, 81–102 (2021).
60. R. R. Wick, M. B. Schultz, J. Zobel, K. E. Holt, *Bioinformatics* **31**, 3350–3352 (2015).
61. S. Nurk, S. Koren, A. Rhie, M. Rautiainen, Zenodo (2021); <https://doi.org/10.5281/zenodo.5598253>.

ACKNOWLEDGMENTS

We thank M. Akeson, A. Carroll, P.-C. Chang, A. Delcher, M. Nattestad, and M. Pop for discussions on sequencing, assembly, and analysis; AnVIL, Amazon Web Services, DNAnexus, the UW Genome Sciences IT Group, and the UConn Computational Biology Core for computational support; and the National Institutes of Health (NIH) Intramural Sequencing Center, the UConn Center for Genome Innovation, and the Stowers Imaging Facility for experimental support. This work used the computational resources of the NIH HPC Biowulf cluster (<https://hpc.nih.gov>). Certain commercial equipment, instruments, or materials are identified to adequately specify experimental conditions or reported results. Such identification does not imply recommendation or endorsement by the National Institute of Standards and Technology, nor does it imply that the equipment, instruments, or materials identified are necessarily the best available for the purpose. **Funding:** Intramural Research Program of the National Human Genome Research Institute, NIH (A.M.P., A.C.Y., A.M.M., M.R., A.R., B.P.W., G.G.B., C.J., J.C.M., N.F.H., S.K., S.N., and S.Y.B.); NIH U01HG010971 (E.E.E., H.L., K.H.M., M.Kr., R.S.F., and T.A.G.-L.); NIH R01HG002385 and R01HG010169 (E.E.E.); NIH R01HG009190 (A.G. and W.T.); NIH R01HG010485 and U01HG010961 (B.P., E.G., and K.S.); NIH U41HG010972 (I.M.H., B.P., E.G., and K.S.); National Science Foundation (NSF) 1627442, 1732253, and 1758800 and NIH U24HG006620, U01CA253481, and R24DK106766 (M.C.S.); NIH U24HG010263 (S.Z. and M.C.S.); Mark Foundation for Cancer Research 19-033-ASP (S.A. and M.C.S.); NIH R01HG006677 (A.Sh. and S.L.S.); NIH U24HG009081 (R.S.F. and T.A.G.-L.); intramural funding at the National Institute of Standards and Technology (J.M., J.M.Z., J.W., and N.D.O.); St. Petersburg State University grant 73023573 (A.M., I.A.A., and T.D.); NIH R01HG002939 (A.F.A.S., I.T.F., and J.M.S.); NIH R01GM124041, R01GM129263, and R21CA238758 (B.A.S.); Intramural Research Program of the National Library of Medicine, NIH (C.X., F.T.-N., and V.A.S.); NIH F31HG011205 (C.J.S.); Damon Runyon Postdoctoral Fellowship and PEW Latin American Fellowship (G.V.C.); Fulbright Fellowship (D.C.S.); Howard Hughes Medical Institute (HHMI) (E.D.J. and G.F.); NIH R01AG054712 (E.I.R.); NIH UMIHG008898 (F.J.S.); NIH

R01GM123312 and R21CA240199 and NSF 1613806 and 1643825 (G.A.H., P.G.S.G., S.J.H., and R.J.O.); NIH R21CA240199, NSF 643825, and Connecticut Innovations 20190200 (R.J.O.); NIH F32GM134558 (G.A.L.); NIH R01HG010040 (H.L.); St. Petersburg State University grant 73023573 (I.A.A.); Wellcome WT206194 (J.T., J.M.D.W., K.H., W.C., and Y.S.); Wellcome WT207492 (R.D.); Stowers Institute for Medical Research (J.L.G.); NIH R01HG011274 (K.H.M.); Ministry of Science and Higher Education of the RF 075-10-2020-116/13.1902.21.0023 (L.U.); supported by the Sirius University (L.U.); Russian Science Foundation (RSF) 19-75-30039 Analysis of genomic repeats (I.A.A.); NIH U41HG007234 (M.D.); NIH DP2MH119424 (M.Y.D.); HHMI Hanna H. Gray Fellowship (N.A.); NIH R35GM133747 (R.C.M.); Childcare Foundation, Swiss National Science Foundation, and European Research Council (ERC) 249968 (S.E.A.); German Federal Ministry for Research and Education 031L0184A (T.M.); Chan Zuckerberg Biohub Investigator Award (A.St.); Common Fund, Office of the Director, NIH (V.V.M.); and the Max Planck Society (E.W.M.). E.E.E. and E.D.J. are investigators of the HHMI. **Author contributions:** Analysis texts are listed, with leads indicated by an asterisk. Assembly: S.N.* S.K.* M.R.* M.A., H.C., C.-S.C., R.D., E.G., M.Ki., M.Ko., H.L., T.M., E.W.M., I.S., B.P.W., A.W., A.M.P. Acrocentrics: A.M.P.* J.L.G.* M.R., S.E.A., M.B., R.D., L.G.d.L., T.P. Validation: A.R.* A.V.B.* A.M.* M.A.* A.M.M.* K.S.* W.C., L.G.d.L., T.D., G.F., A.F., K.H., C.J., E.D.J., D.P., V.A.S., Y.S., B.A.S., F.T.-N., J.T., J.M.D.W., A.M.P. Segmental duplications: M.R.V.* E.E.E.* S.N., S.K., M.D., P.C.D., A.G., G.A.L., D.P., C.J.S., D.C.S., M.Y.D., W.T., K.H.M., A.M.P. Satellite annotation: N.A.* I.A.A.* K.H.M.* A.V.B., L.U., T.D., L.G.d.L., P.A.P., E.I.R., A.St., B.A.S., A.M.P. Epigenetics: A.G.* W.T.* S.K., A.R., M.R.V., N.A., S.J.H., G.A.L., G.V.C., M.C.S., R.J.O., E.E.E., K.H.M., A.M.P. Variants: S.A.* D.C.S.* S.M.Y.* S.Z.* R.C.M.* M.Y.D.* J.M.Z.* M.C.S.* N.F.H., M.Ki., J.M., D.E.M., N.D.O., J.A.R., F.J.S., K.S., A.Sh., J.W., C.X., A.M.P. Repeat annotation: S.J.H.* R.J.O.* A.G., P.G.S.G., G.A.H., L.G.d.L., A.F.A.S., J.M.S. Gene annotation: M.D.* M.H.* A.Sh.* S.N., S.K., P.C.D., I.T.F., S.L.S., F.T.-N., A.M.P. Browsers: M.D.* N.-C.C., P.K. Data generation: S.J.H., G.G.B., S.Y.B., G.V.C., R.S.F., T.A.G.-L., I.M.H., M.W.H., M.J., J.K., M.Kr., V.V.M., J.C.M., B.P., P.A.C.Y., U.S., M.Y.D., J.L.G., R.J.O., T.D., E.E.E., K.H.M., A.M.P. Computational resources: C.-S.C., A.F., R.J.O., M.C.S., K.H.M., A.M.P. Manuscript draft: A.M.P. Figures: S.K., S.N., A.M.P., A.R. Editing: A.M.P., S.N., S.K., A.R., E.E.E., and K.H.M., with the assistance of all authors. Supplement: S.N. and S.K., with the assistance of the working groups. Supervision: R.C.M., M.Y.D., I.A.A., J.L.G., R.J.O., W.T., J.M.Z., M.C.S., E.E.E., K.H.M., A.M.P. Conceptualization: E.E.E., K.H.M., A.M.P. **Competing interests:** A.F. and C.-S.C. are employees of DNAnexus; I.S., J.K., M.W.H., P.P., and A.W. are employees of Pacific Biosciences; S.A. is an employee and stockholder of Oxford Nanopore Technologies; E.E.E. is a science advisory board member of Variant Bio; K.H.M. is a science advisory board member of Centaura; P.K. owns and receives income from Reservoir Genomics LLC; F.J.S. has received travel funds to speak at events hosted by Pacific Biosciences; S.K., D.E.M., F.J.S., and K.H.M. have received travel funds to speak at events hosted by Oxford Nanopore Technologies; and W.T. has licensed two patents to Oxford Nanopore Technologies (US 8748091 and 8394584). **Data and materials availability:** The T2T-CHM13 and T2T-HG002-ChrX assemblies generated by this study are archived under National Center for Biotechnology Information (NCBI) GenBank accession numbers GCA_009914755 and CP086568, respectively. CHM13hTERT cells were obtained for research use through a material transfer agreement with U. Surti and the University of Pittsburgh. The raw sequencing data were described in prior studies and are summarized in table S1. For convenience, links to the sequence data and genome browsers are also available from <https://github.com/marbl/CHM13>. Supplementary data for fig. S39 and the string graph construction code are archived at Zenodo (61) and also https://github.com/snurk/sg_sandbox.

SUPPLEMENTARY MATERIALS

science.org/doi/10.1126/science.abj6987
Material and Methods
Figs. S1 to S39
Tables S1 to S17
References (62–128)
MDAR Reproducibility Checklist

[View/request a protocol for this paper from Bio-protocol.](#)

26 May 2021; accepted 28 January 2022
10.1126/science.abj6987

The complete sequence of a human genome

Sergey NurkSergey KorenArang RhieMikko RautiainenAndrey V. BzikadzeAlla MikheenkoMitchell R. VollgerNicolas AltemoseLev UralskyAriel GershmanSergey AganezovSavannah J. HoytMark DiekhansGlennis A. LogsdonMichael AlongeStylianos E. AntonarakisMatthew BorchersGerard G. BouffardShelise Y. BrooksGina V. CaldasNae-Chyun ChenHaoyu ChengChen-Shan ChinWilliam ChowLeonardo G. de LimaPhilip C. DishuckRichard DurbinTatiana Dvorkinalan T. FiddesGiulio FormentiRobert S. FultonArkarachai FungtammasanErik GarrisonPatrick G. S. GradyTina A. Graves-LindsayIra M. HallNancy F. HansenGabrielle A. HartleyMarina HauknessKerstin HoweMichael W. HunkapillerChirag JainMiten JainErich D. JarvisPeter KerpedjievMelanie KirscheMikhail KolmogorovJonas KorlachMilinn KremitzkiHeng LiValerie V. MaduroTobias MarshallAnn M. McCartneyJennifer McDanielDanny E. MillerJames C. MullikinEugene W. MyersNathan D. OlsonBenedict PatenPaul PelusoPavel A. PevznerDavid PorubskyTamara PotapovaEvgeny I. RogaevJeffrey A. RosenfeldSteven L. SalzbergValerie A. SchneiderFritz J. SedlazeckKishwar ShafinColin J. ShewAlaina ShumateYing SimsArian F. A. SmitDaniela C. Sotolvan Sovi#Jessica M. StorerAaron StreetsBeth A. SullivanFrançoise Thibaud-NissenJames TorranceJustin WagnerBrian P. WalenzAaron WengerJonathan M. D. WoodChunlin XiaoStephanie M. YanAlice C. YoungSamantha ZarateUrvashi SurtiRajiv C. McCoyMegan Y. DennisIvan A. AlexandrovJennifer L. GertonRachel J. O'NeillWinston TimpJustin M. ZookMichael C. SchatzEvan E. EichlerKaren H. MigaAdam M. Phillippy

Science, 376 (6588), • DOI: 10.1126/science.abj6987

View the article online

<https://www.science.org/doi/10.1126/science.abj6987>

Permissions

<https://www.science.org/help/reprints-and-permissions>

Use of this article is subject to the [Terms of service](#)

Science (ISSN) is published by the American Association for the Advancement of Science. 1200 New York Avenue NW, Washington, DC 20005. The title *Science* is a registered trademark of AAAS.

Copyright © 2022 The Authors, some rights reserved; exclusive licensee American Association for the Advancement of Science. No claim to original U.S. Government Works