

Reviewed Preprint

v1 • April 15, 2025

Not revised

Reviewed Preprint

v2 • September 18, 2026

Revised by authors

✉ For correspondence:

[zhihao.ding@boehringer-
ingelheim.com](mailto:zhihao.ding@boehringer-
ingelheim.com)

^{*}, ⁺ **Author notes:** See [page 21](#)

Reviewing editor: Jun Kim,
Chungnam National University,
Republic of Korea

© 2025, Noyvert et al. This article is
distributed under the terms of the
[Creative Commons Attribution
License](#), which permits unrestricted
use and redistribution provided that
the original author and source are
credited.

Imputation of structural variants using a multi-ancestry long-read sequencing panel enables identification of disease associations

Boris Noyvert^{1,6,*}, A Mesut Erzurumluoglu^{1,*}, Dmitry Drichel^{2,*}, Steffen Omland^{2,*}, Till FM Andlauer^{1,3,*}, Stefanie Mueller¹, Lau Sennels², Christian Becker², Aleksandr Kantorovich², Boris A Bartholdy¹, Ingrid Brænne¹, Julio Cesar Bolivar-Lopez¹, Costas Mistrellides¹, Gillian M Belbin⁴, Jeremiah H Li⁴, Joseph K Pickrell⁴, Jatin Arora¹, Yao Hu¹, Boehringer Ingelheim – Global Computational Biology and Digital Sciences, Clive R Wood⁵, Jan M Kriegel¹, Nikhil Podduturi², Jan N Jensen¹, Jan Stutzki^{2,+}, Zhihao Ding^{1,+} ✉

¹Global Computational Biology and Digital Sciences (gCBDS), Boehringer Ingelheim Pharma GmbH & Co. KG, Biberach an der Riss, Germany • ²BI X GmbH, Ingelheim am Rhein, Germany • ³Department of Neurology, School of Medicine, Technical University of Munich, Munich, Germany • ⁴Gencove, New York, United States • ⁵Discovery Research, Boehringer Ingelheim Pharma GmbH & Co. KG, Ingelheim am Rhein, Germany • ⁶Institute of Cancer and Genomic Sciences, University of Birmingham, Birmingham, United Kingdom

eLife Assessment

This **fundamental** study delivers a population reference panel for long-read sequencing-based structural variants and demonstrates its utility for disease association analyses in the UK Biobank, expanding genetic discovery beyond conventional SNV-based approaches. The authors provide **convincing** evidence through extensive benchmarking and systematic analyses that the panel improves structural variant detection and can support fine-mapping of trait associations. The broader impact will depend on timely access to the imputed UK Biobank resource, or on provision of a practical, reproducible pipeline and associated summary statistics.

<https://doi.org/10.7554/eLife.106115.2.sa2>

Abstract

Advancements in long-read sequencing technology have accelerated the study of large structural variants (SVs). We created a curated, publicly available, multi-ancestry SV imputation panel by long-read sequencing 888 samples from the 1000 Genomes Project. This high-quality panel was used to impute SVs in approximately 500,000 UK Biobank participants. We demonstrated the feasibility of conducting genome-wide SV association studies at biobank scale using 32 disease-relevant phenotypes related to respiratory, cardiometabolic and liver diseases, in addition to 1,463 protein levels. This analysis identified thousands of genome-wide significant SV associations, including hundreds of conditionally independent signals, thereby enabling novel biological insights. Focusing on genetic association studies of lung function as an example, we demonstrate the added value of SVs for prioritising causal genes at gene-rich loci compared to traditional GWAS using only short variants. We envision that future post-GWAS gene-prioritisation workflows will incorporate SV analyses using this SV imputation panel and framework.

Introduction

Human disease associations of single nucleotide variants (SNVs) and short insertions and deletions are routinely identified in genome-wide association studies (GWASs)^{1,2}. By contrast, large structural variants (SVs) of >50 base pairs (bp) are typically neglected, despite functional roles in the context of disease. Each human carries 23,000–31,000 SVs¹, often overlapping protein-coding genes or regulatory regions, thus enabling fine-mapping causal genetic variants^{2,3}. Studies on single populations have already demonstrated the value of sequencing SVs for identifying causal variants underlying disease associations^{4,5}.

Robust SV calling from traditional short-read sequencing is challenging because SVs are often longer than the average short-read length (Figure 1a [↗](#))^{3,6}. Long-read sequencing captures SVs reliably, but the high cost impairs its application to large-scale datasets. Reference panels constructed for imputation of SVs from genotyped samples enable biobank-scale genome-wide analyses of SVs. For example, imputing SVs in UK Biobank (UKB) can accelerate research on the genetic underpinnings of diverse diseases and facilitate the identification of novel therapeutic targets. To this end, we generated a publicly available multi-ancestry long-read sequencing-based SV imputation panel (Figure 1b [↗](#)). Simultaneously, an effort is ongoing to develop new methods for improved SV calling that utilises this dataset⁷. Another Oxford nanopore sequencing-based dataset of 1000 Genomes Project samples is being generated, SV calling of the first 100 samples was reported recently⁸.

Results

Structural variant calling and benchmarking

We performed long-read whole-genome sequencing of 906 individuals sampled from the 1000 Genomes Project, with a median read length of ~6.2 kbp (Supplementary Table 1 [↗](#), Supplementary Figure 1 [↗](#)) and 15x median coverage (alignment with minimap2; Supplementary Table 2 [↗](#), Supplementary Figure 2 [↗](#)). The 888 samples passing quality control (QC) included 164 Europeans, 144 (admixed) Americans, 168 East Asians, 171 South Asians, and 241 Africans (Figures 2a [↗](#) and 2b [↗](#); Supplementary Table 3 [↗](#)).

Joint (multi-sample) calling was performed using Sniffles2⁹. To assess the quality of our SV calling, we benchmarked data generated from the 1000 Genomes project participant NA12878 (also known as HG001). This individual was analysed by us, the Genome in a Bottle (GIAB) initiative using high-coverage PacBio HiFi long-read sequencing¹⁰, and by Byrsk-Bishop *et al.*¹¹ via high-coverage short-read Illumina sequencing (referred to as the NYGC dataset).

We compared our raw SV calls for NA12878 to the GIAB calls of the same individual using the Truvari bench tool¹². In a whole-genome comparison, we observed a precision of 71.8%, a recall of 76.3%, and an F1 score of 74.0%. In a second benchmark excluding tandem repeat regions longer than 200 bases, we achieved a high precision of 90.4%, a recall of 91.5%, and an F1 score of 91.0% (see Methods for further details). Comparing our NA12878 SV calls to the NYGC dataset, we observed a precision of 15.9% and a recall of 85.4%. This indicates that our long-read approach detected most of the SVs identified by NYGC short-read sequencing but also discovered many additional SVs. The majority (67.5%) of SVs not included in the NYGC data were present in the GIAB call set, strongly supporting the validity of our SV calls. Similarly high recall and low precision rates compared to NYGC were observed for other individuals (Methods, Supplementary Figure 3 [↗](#)).

SV characterisation

We used the following criteria for selecting 107,445 SVs for inclusion in the imputation panel: length between 50 bp and 30 Mbp, missingness rate <20%, and being present in at least 2 out of the 888 individuals (Methods, Supplementary Table 4 [↗](#), Supplementary Figures 4 and 5 [↗](#)). Based on the benchmarking of our SV calls against the NYGC call set, we annotated SVs in the panel with

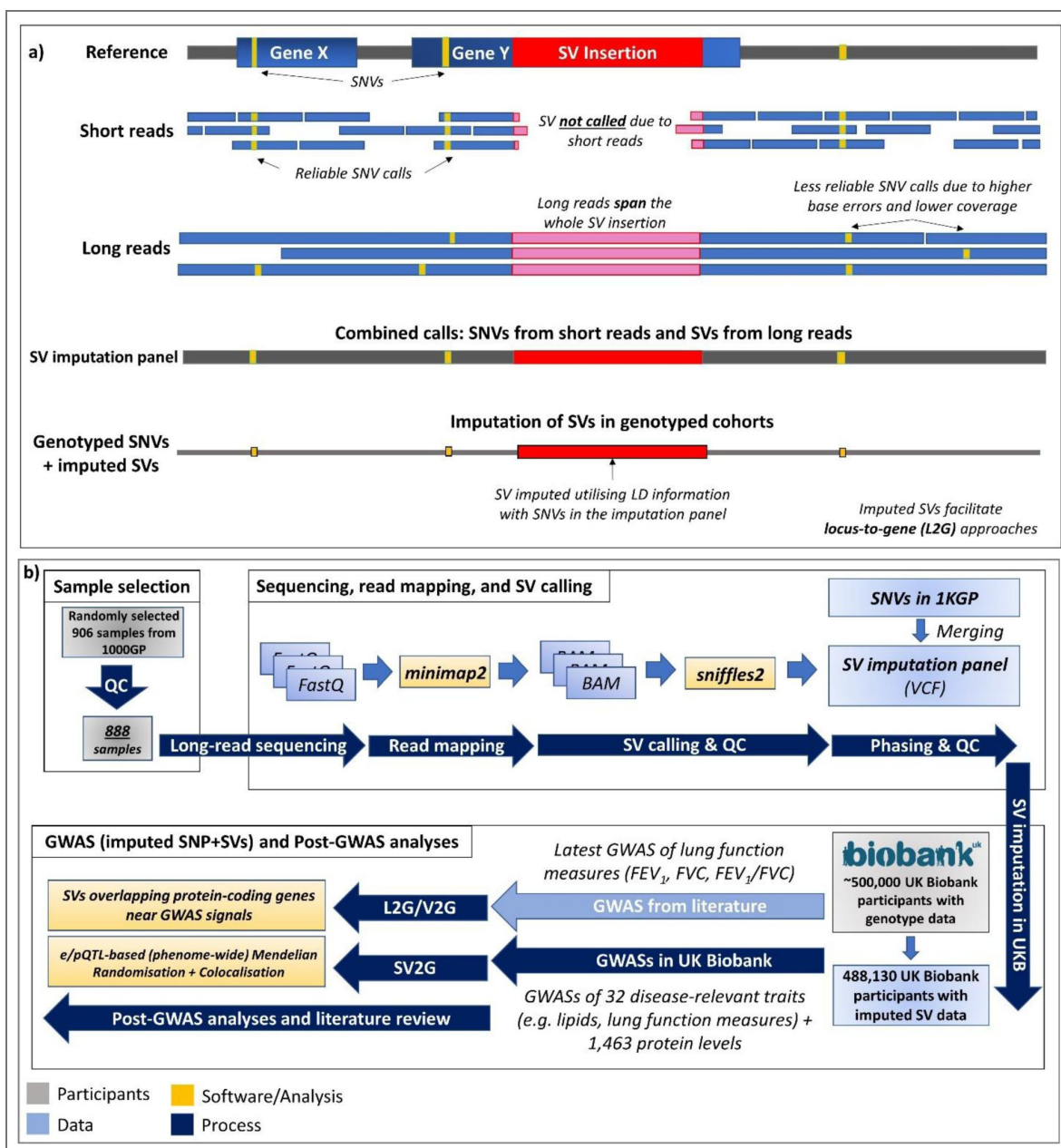


Figure 1. Study design and visual abstract of purpose and benefits of long-read sequencing and a multi-ancestry imputation panel.

a) Long-reads that span large SVs enable identification of novel SVs and improve calling of known SVs, which better inform locus-to-gene (L2G) approaches. b) Study design including the sample selection, SV calling, SV imputation in UK Biobank, GWAS, and post-GWAS stages (not exhaustive). For the GWASs in UK Biobank, separate analyses using the same phenotype definitions were carried out using (i) the imputed SNVs, and (ii) imputed SVs (i.e., SV-WAS). The two GWAS summary statistics were then combined for the relevant post-GWAS analyses. L2G/V2G: locus-to-gene/variant-to-gene (carried out on top SNVs identified by Shrine et al., which have an SV deletion within 500 kb); SV2G: SV-to-gene (carried out **only** on SV deletions that were the most significant variants in the associated loci in the GWASs we carried out in UK Biobank, which also overlap with protein-coding genes); 1KGP: 1000 Genomes Project (Phase 1); UKB: UK Biobank.

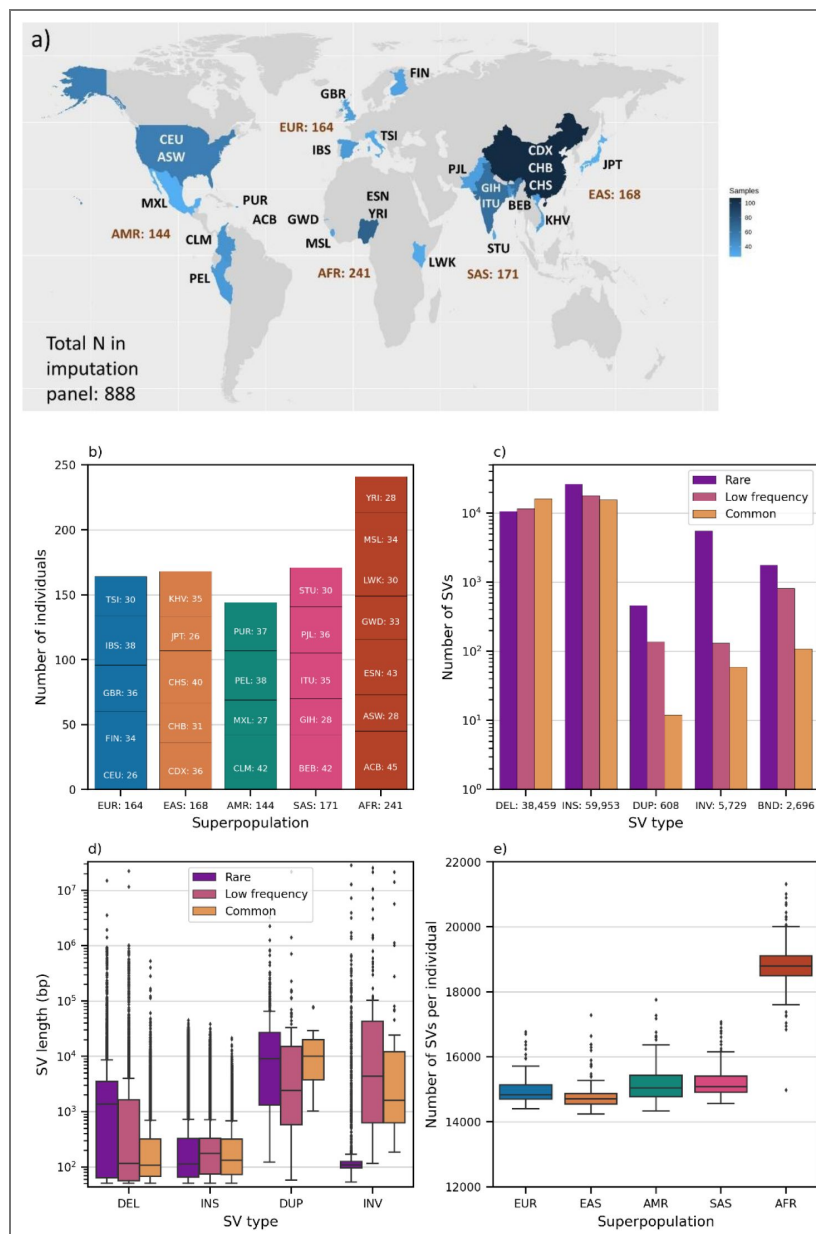


Figure 2. Characterisation of the quality-controlled imputation panel

(N=888 samples and n=107,445 SVs). **a)** Map of the 888 samples from the 1000 Genomes project, mapping the samples to countries and not indicating detailed geographical origins of populations. AMR: Admixed American, AFR: African, EUR: White European, EAS: East Asian, SAS: South Asian ancestries. **b)** Sample counts by superpopulation and population codes (population code description in Supplementary Table 3), using abbreviations from 1000 Genomes Project. **c)** Number of SVs by variant type (DEL: deletion, INS: insertion, DUP: duplication, INV: inversion, BND: break-end) and frequency class (common: $0.05 < \text{MAF} \leq 0.5$, low frequency: $0.005 < \text{MAF} \leq 0.05$, rare: $0 < \text{MAF} \leq 0.005$), with total SV counts by SV type. **d)** SV size distributions by SV type and frequency class, excluding break-ends, which do not have a size. **e)** Number of minor SV alleles per individual by superpopulation in the imputation panel.

three categories: “novel”, SVs not detected in the NYGC dataset (75,192 or 70%); “known”, SVs previously identified in at least one individual (29,119 or 27.1%); and “not compared”, SVs ignored by the Truvari bench tool (3,134 or 2.9%) (see Methods).

The GIAB consortium described which genomic regions confidently produce reliable genotypes¹³, and 41.9% of our SVs mapped to high-confidence (‘confident’ henceforth) regions. The most common SV types were insertions (55.8%), followed by deletions (35.8%), inversions (5.3%), break-ends (typically unresolved SVs; 2.5%), and duplications (<0.06%) (Figure 2c [↗](#)). Most SVs were rare (minor allele frequency (MAF) <0.005: 41.8%), the remainder being low-frequency (MAF 0.005-0.05: 28.7%) and common (MAF >0.05: 29.6%) variants, with the latter being enriched for insertions and deletions. Note that the panel contains a larger proportion of rare SNVs than rare SVs because a) rare SVs are more difficult to detect reliably compared to rare SNVs and b) the SNVs were originally called in the larger, full 1000 Genomes cohort.

Sizes of most SV deletions and insertions were between 50 and 1000 bp (Figure 2d [↗](#)). Duplications and inversions predominantly ranged from 1 to 30 kbp. A considerable number of SVs extended up to 1 Mbp and beyond. However, SVs that are longer than 1 Mbp are likely to be artefacts of the SV calling process (see Methods).

Of the 107,445 SVs, 4,406 SVs were predicted using SnpEff¹⁴ to have high functional impact. Of these, 1,198 were predicted to cause frameshifts, and 581 the complete loss of exons. Based on GWAS Catalog¹⁵, 2,465 SVs overlapped positions of variants with known GWAS associations (Supplementary Table 5 [↗](#)).

On average, individuals carried 16,065 SVs, with individuals of African ancestry exhibiting the highest (mean: 18,822 SVs), and East Asian-ancestry individuals exhibiting the lowest (14,729) SV diversity (Figure 2e [↗](#), Supplementary Table 6 [↗](#))^{3,16}. This observation was most pronounced for insertions and deletions, which are more frequent and can be called with high confidence. Only 36.6% (39,273) of SVs were shared across all superpopulations with individuals of African ancestry exhibiting the most SVs not shared with other populations (13,153) and (Admixed) Americans the least (841) (Supplementary Table 7 [↗](#)).

Generation of SV imputation panel and application to UK Biobank

To enable the imputation of SVs in genotyped studies, we constructed a haplotype reference panel consisting of 888 samples. To this end, we merged the 107,445 SVs generated in the present study with the ~45M short variants identified from short-read sequencing data of the 1000 Genomes Project Phase 3 release¹¹ (see Methods).

We assessed the suitability of the panel for SV imputation into UKB by performing leave-one-out imputation for all individuals in the panel. Here, we specifically emulated the process of SV imputation in UKB by using all genotyped UKB SNVs present in the reference panel (~702K SNVs) as the basis for imputation. To facilitate benchmarking, we not only imputed SVs but also ~58K randomly selected SNVs from the panel. We calculated the per-variant non-reference genotype concordance¹⁷ and the imputation quality metric r^2_{imp} (see Figure 3b [↗](#); Methods; Supplementary Table 9 [↗](#), Supplementary Figure 6 [↗](#)). Both metrics varied based on MAF, GIAB region type, and variant type. Generally, rare SVs showed poorer imputation quality than common insertions and deletions, which produced high-quality imputation metrics. The mean non-reference concordance for common insertions and deletions was 0.696 and 0.693, respectively (both concordances were 0.846 in confident regions), with mean r^2_{imp} = 0.718 and 0.721 (0.858 and 0.852 in confident regions, see Supplementary Table 9 [↗](#) and Figure 3b [↗](#)). These metrics demonstrate sufficient imputation quality for conducting GWAS of common variants. As anticipated, the imputation quality of SVs was lower than that of imputed SNVs across GIAB region types and MAF classes, due to the greater difficulty in reliably calling SVs; however, this difference was not substantial in confident regions (Figure 3b [↗](#)). Note that, for our SV association analyses (see below), we only evaluated results for high-quality SVs.

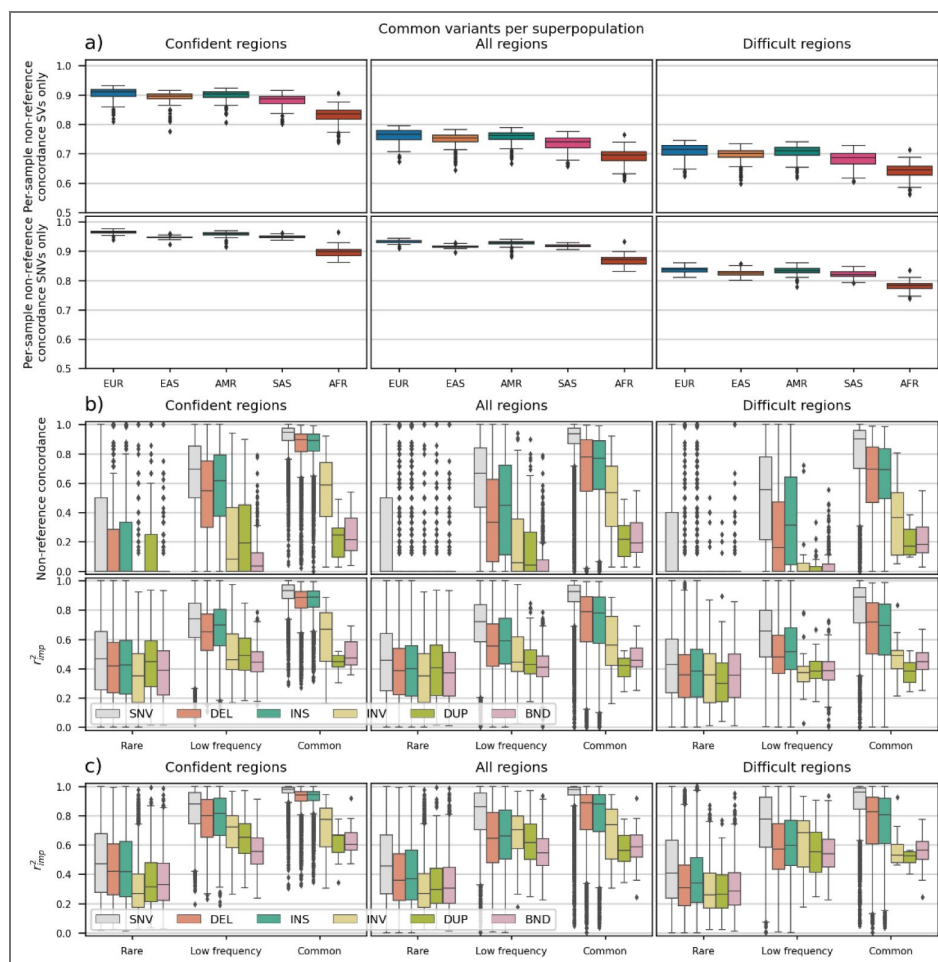


Figure 3. Evaluation of imputation quality:

a) Leave-one-out imputation: sample-wise non-reference concordance stratified by superpopulation and GIAB region type, based on imputed common SVs (top) and imputed common SNVs (bottom). b) Leave-one-out imputation: variant-wise non-reference concordance (top) and r^2_{imp} score (bottom) stratified by variant type, MAF class, and GIAB region type. c) Imputation into UK Biobank: r^2_{imp} score by variant type and MAF class, for confident regions, all regions, and difficult regions.

We also computed the per-individual non-reference genotype concordance for each superpopulation, based on common SVs and SNVs stratified by GIAB region type (Figure 3a; Supplementary Table 8). African-ancestry individuals, being the most diverse, had a mean non-reference concordance of 0.690 based on all common SVs, which was the lowest amongst all superpopulations. By comparison, individuals of European ancestry showed a mean non-reference concordance of 0.760 for common SVs. Overall, we did not observe significant outliers in the per-individual concordance values, indicating that there were no issues with sequencing and data processing of the samples. In conclusion, the SV reference panel offers a robust foundation for imputing SVs, particularly common insertions and deletions, into UKB and for performing subsequent GWAS.

We used our imputation panel to impute SVs into the genomes of 488,130 UKB participants (see Methods). Imputation quality in UKB mirrored our observations from the leave-one-out validation in the 1000 Genomes study: r^2_{imp} was higher for common variants and in high-confidence regions. For example, common deletions and insertions in high-confidence regions showed mean r^2_{imp} = 0.915, compared to 0.961 for common SNVs in the same regions (Supplementary Tables 10 and 11; Figure 3c). Notably, UKB predominantly consists of European-ancestry individuals, for which we observed higher average imputation qualities in the 1000 Genomes leave-one-out analyses.

Finally, as a proof of principle, we manually compared one deletion (Sniffles2.DEL.3639MF) imputed from the genotype data to the recent UKB short-read based whole-genome sequencing (WGS) data¹⁸. To this end, we analysed the read coverage from *vcf* files of the respective genomic region to determine chromosomes carrying the deletion allele (Methods and Supplementary Figure 7). Thus, we could group individuals based on their ‘Sniffles2.DEL.3639MF’ deletion genotype and compare them to the imputed genotypes (Supplementary Figure 8). We found a high concordance rate of 98.7% between the WGS and imputed genotypes, exceeding the accuracy estimations from the leave-one-out analyses. We chose this specific deletion for the proof-of-principle analysis because it a) has an interesting disease association and is thus relevant (see Figure 4) and b) our calling approach for SVs from WGS data only works for deletions in non-repetitive regions (see Methods).

Proof-of-principle SV imputation and genome-wide association studies in UK Biobank

To demonstrate the added value of the imputation panel, we selected 19 exemplary continuous traits and 13 binary traits available in UKB relevant to respiratory (n=6 traits), metabolic (n=16), and liver (n=10) diseases (Supplementary Tables 12 and 13). On these phenotypes, we performed GWASs of the imputed SVs (SV-WAS) in up to 453,754 European UKB participants¹⁹ (Supplementary Figure 9).

Filtering variants by imputation metric INFO>0.7 and MAF>0.01, 3,858 SV associations (in 1,898 unique SVs, 1,258 of which were in the *novel* category) passed the established genome-wide significance threshold of $p < 5 \times 10^{-8}$ in any phenotype (Supplementary Table 14). For an in-depth assessment of potential causal genes, we selected all significantly associated SVs from these SV-WASs that overlapped functional protein-coding genes (some SVs spanned several genes), thereby prioritising 689 unique genes (Figure 4a; Supplementary Table 14; Methods). We also analysed how SVs influence the levels of 1,463 proteins in blood plasma. Here, we identified 10,518 significant ($p < 5 \times 10^{-8}/1463 = 3.4 \times 10^{-11}$, INFO>0.7, MAF>0.01) SV-based pQTLs (3,723 unique SVs, 2,495 of which were novel) for 1,101 proteins (Supplementary Table 15), including 84 (excluding the major histocompatibility complex region, *i.e.*, *6p21*) that were conditionally independent of nearby SNVs (Methods; Supplementary Table 16).

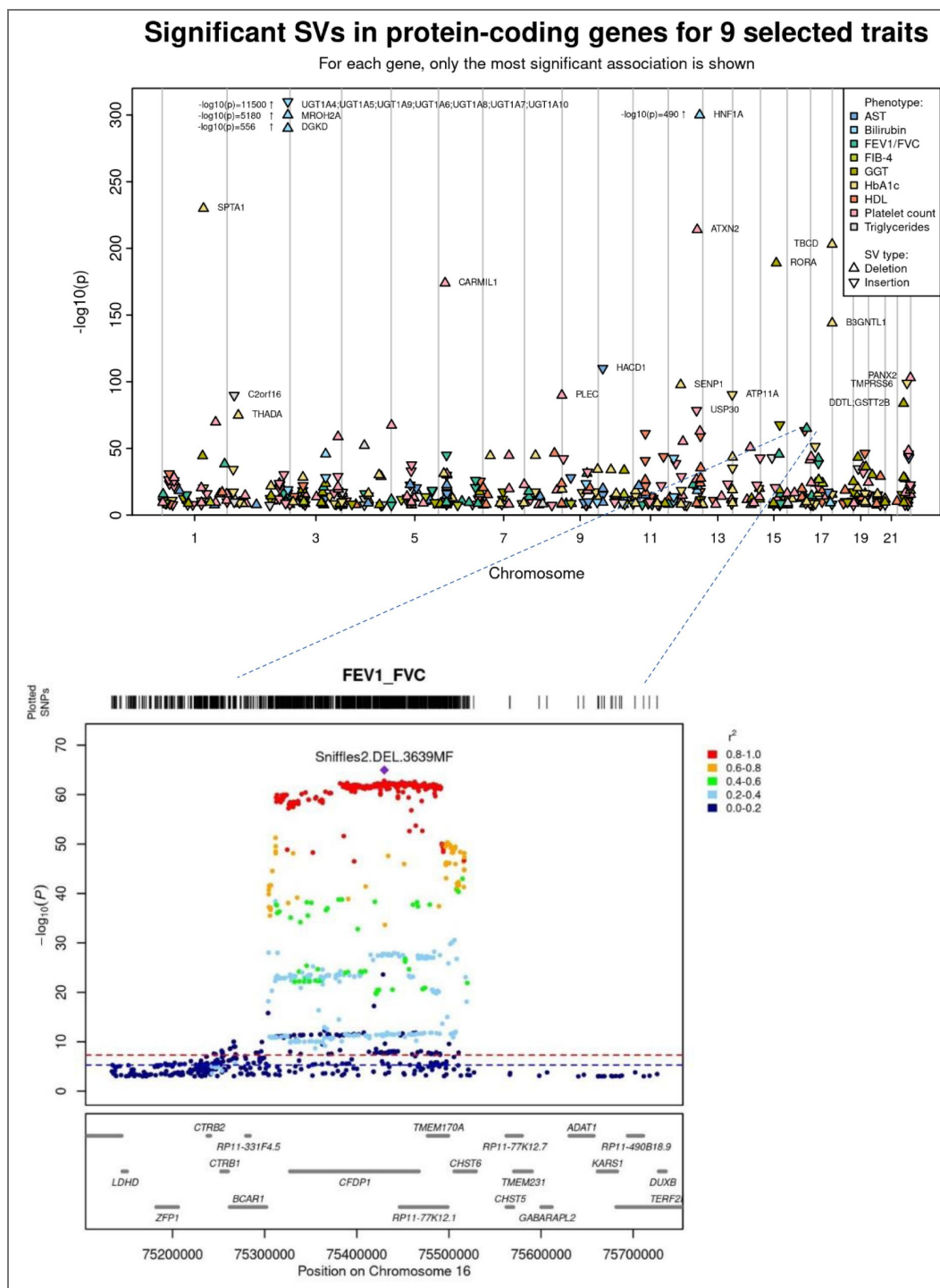


Figure 4.

a) Manhattan plot of 9 selected traits, with only GW-significant SVs overlapping with protein-coding genes shown. Phenotype definitions and full list of GW-significant SVs are available in Supplementary Tables 12-13 and 14 [\[13\]](#), respectively. **b)** Region plot showing the FEV₁/FVC association of Sniffles2.DEL.3639MF (841-base SV deletion in an intron of *CFDP1*) and nearby SNVs in a GWAS of UK Biobank participants. LD information between the variants was calculated using the same population utilised in the UKB GWASs. Additional details on the SV can be found in Supplementary Table 5 [\[13\]](#).

Systematic analysis of SV information in identifying novel associations and causal genes

To follow up on the significant SV associations identified in the UKB SV-WAS and pQTL analyses, we examined whether these SV associations provide added value compared to analysing SNVs alone. To this end, we carried out association analyses on imputed UKB SNVs mapping to each locus containing significant SVs, using the same sample inclusion criteria, phenotype transformations, and covariates as in the SV analyses. To identify whether SVs showed more robust evidence for an association, *i.e.*, lower *p*-values, at the SNV-based GWAS- and pQTL-associated loci, we combined the SNV and SV association results. At 55 genome-wide significant GWAS loci, an SV showed the lowest *p*-value (Supplementary Table 17 [↗](#)). Conditional analyses revealed that SVs constituted a secondary, independent signal at 23 further loci; 38 of these 78 SVs overlapped with protein-coding genes.

We then systematically assessed the contribution of SV associations to the prioritisation of causal genes by comparing genes implicated by SVs to the putatively causal genes reported by the latest GWAS of lung function measures published by Shrine *et al.*²⁰. The authors used several post-GWAS approaches (e.g., nearest gene, e/pQTLs, polygenic priority score (PoPS, a tool for gene prioritisation²¹), and rare respiratory disease causal genes) to map associated loci to genes, thereby prioritising on average ~3 genes per locus. We examined the reported autosomal loci associated with ≥1 of the three lung function phenotypes (*i.e.*, FEV₁, FVC, or FEV₁/FVC – the latter a clinically important lung function measure utilised in diagnosing chronic obstructive pulmonary disease). Seventy of these reported loci harboured SVs (*i.e.*, an SV mapped within ±500kb of the top SNV) that were significantly associated with the same primary lung function trait in our UKB-based SV-WASs (Supplementary Table 18 [↗](#)). At 55 of these 70 loci, the gene implicated by our SV analyses was also implicated by ≥1 of the post-GWAS methods utilised by Shrine *et al.* – strengthening the evidence for causality for the respective genes. In the remaining 15 loci where SV-implicated genes had not been prioritised yet, SVs pointed to genes very likely involved in lung health via influencing smoking behaviour (e.g., *SLC1A2* - highly expressed in the brain; Supplementary Figure 10 [↗](#)) or other lung-disease causal mechanisms such as abnormal respiratory ciliary function²² (e.g., *DNAH12* and *DYNLRB1*) and promotion of inflammation²³ (e.g., *PRDX1*).

Added value of SV information to gene prioritisation: examples from GWASs of lung function measures

In addition to the identification of the above-mentioned genes, the added value of SV information for identifying the causal gene, *i.e.*, improving locus-to-gene (L2G) prioritisation pipelines at the loci identified by Shrine *et al.*²⁰, can be demonstrated using four representative examples:

In our association analyses of quantitative lung function measures in UKB, SVs constituted the conditionally independent variant in primary or secondary signals of 14 loci, thereby facilitating identification of the respective causal genes (Supplementary Table 17 [↗](#)). One of these FEV₁/FVC-associated loci contained an 841-base SV deletion in an intron of *CFDP1* (Sniffles2.DEL.3639MF; $p=1.1\times 10^{-65}$; MAF=0.413; Figure 4b [↗](#)). We confirmed this deletion using WGS data. Shrine *et al.* prioritised three potentially causal genes at this locus (including *CTRB1* and *BCAR1*), with *CFDP1* only being implicated by its proximity to the top SNV (rs11864587) without any functional or regulatory support. Follow-up analyses on this gene, including phenome-wide cis-eQTL-based Mendelian Randomisation (MR) and colocalisation analyses (not carried out by Shrine *et al.*) provided strong evidence (min MR $p=1.1\times 10^{-22}$; colocalisation posterior probability (PP) >90%) for putatively causally linking *CFDP1* expression in various (bulk GTEx) tissues including fibroblasts – a cell type relevant for respiratory disease – to a lower FEV₁/FVC measure (Supplementary Figure 11 [↗](#); Supplementary Table 19 [↗](#)). *CFDP1* and *BCAR1* protein expression levels were not assayed by the UKB-PPP²⁴, thus information from SV-based pQTLs could not be utilised to distinguish between the genes.

At a second FEV₁/FVC-associated locus (top-associated SNV: rs947350), Shrine *et al.* prioritised ten genes, including *MEGF6* with suggestive evidence from rare coding variants. In our study, an FEV₁/FVC-associated SV deletion (Sniffles2.DEL.6E8M0; $p=2.1\times 10^{-16}$; MAF=0.069) mapped only to *MEGF6*. Similar to the *CFDP1* example, phenome-wide MR and colocalisation analyses provided strong evidence (min MR $p=6.3\times 10^{-7}$; colocalisation PP>97%) putatively causally linking *MEGF6* expression in various (bulk GTEx) tissues including skeletal muscle and heart tissue (a likely respiratory disease-relevant cell type as a smooth muscle-containing tissue) to a lower FEV₁/FVC measure (Supplementary Figure 12 [↗](#); Supplementary Table 20 [↗](#)).

Third, *AAGAB* was solely implicated by two different FVC-associated SV deletions (Sniffles2.DEL.268AME; $p=1.6\times 10^{-12}$; MAF=0.231; and Sniffles2.DEL.2689ME; $p=2.6\times 10^{-12}$; MAF=0.233) but not by any of the post-GWAS methods utilised by Shrine *et al.* Similar to the above examples, phenome-wide MR and colocalisation analyses provided strong evidence (min MR $p=1.7\times 10^{-15}$; colocalisation PP>90%) putatively causally linking *AAGAB* expression in various (bulk GTEx) tissues, including lung tissue, directly to lung function measures (Supplementary Figure 13 [↗](#); Supplementary Table 21 [↗](#)).

As a final example, Shrine *et al.* identified rs7108992 to be associated with FVC and FEV₁. They prioritised two genes, *ETS1* (by proximity) and *FLI1* (via PoPS). In our analysis, we identified an FVC-associated SV deletion only mapping to *FLI1* (Sniffles2.DEL.5C9EMA; $p=2.0\times 10^{-30}$; MAF=0.322). Phenome-wide MR and colocalisation analyses strongly linked *FLI1* expression in lung and smooth muscle-containing tissues to lung function measures (MR $p=3.9\times 10^{-4}$; colocalisation PP=96%), pulmonary heart disease risk (MR $p=8.8\times 10^{-7}$; colocalisation PP=80%), *FGF10*²⁵ (MR $p=7.5\times 10^{-5}$; colocalisation PP=89%), and *LRP1* protein levels²⁶ (MR $p=2.8\times 10^{-5}$; colocalisation PP=87%) – known factors contributing to respiratory diseases²⁵ (Supplementary Figure 14 [↗](#)).

Taken together, these examples illustrate the potential added value of SV information for the identification of novel gene-disease associations, and for improving gene prioritisation pipelines applied to GWAS summary statistics (e.g., the composite locus-to-gene (L2G) score calculated in Open Targets Genetics²⁷).

Discussion

Long-read sequencing holds the promise of conducting reliable association studies of SVs in large cohorts, but its widespread adoption is impeded by its significant cost. For example, comprehensive long-read sequencing of all UKB participants would cost approximately 0.5 billion USD, based on an estimate of 1,000 USD per whole genome. It was recently suggested to reduce costs by sequencing only a limited number of SVs²⁸. By contrast, our SV imputation approach allows for analyses of a comprehensive, genome-wide SV panel without additional sequencing costs. Therefore, use of an SV imputation panel constitutes a practical and cost-effective solution for the robust analysis of common SVs. The multi-ancestry imputation panel applied in the present study to Europeans from UKB can also be used to impute SVs in diverse ancestries, e.g., from BioBank Japan²⁹, China Kadoorie Biobank³⁰, or the Singapore Precision Medicine Programme³¹.

In this context, we observed that the number of SVs detected per individual differed between superpopulations. We identified the highest average number of SVs in individuals of African descent and a slightly lower average in East Asians, compared to the other superpopulations. While we included a higher number of African ancestry individuals, the number of East Asian individuals included in our reference panel was comparable to the number of individuals from other, non-African ancestries. In fact, it was even larger than the number of European ancestry individuals (AFR $n=241$, SAS $n=171$; EAS $n=168$; EUR $n=164$; AMR $n=144$). Therefore, we do not expect a major bias from underrepresentation of any superpopulation in the training dataset. It is well established that African ancestry is more diverse than is the case for other superpopulations^{32,33} and previous studies indicate that East Asian populations tend to exhibit slightly lower genetic diversity compared to European populations³⁴, which is consistent with the lower observed SV counts per genome.

We benchmarked SV calls of the individual NA12878 against GIAB's PacBio-based calls, achieving both high precision and recall, and found that our SV calls, compared to Illumina's short-read data, detected most known but also many additional ('novel') SVs, with validation from GIAB's PacBio dataset. Seventy percent of the SVs contained in our imputation panel can be considered as novel, demonstrating the value of long-read sequencing in SV identification efforts. The size of the sample used for our long-read sequencing was smaller than some of the short-read sequencing samples available. However, our analyses have demonstrated that long-read sequencing is required to detect most of the novel SVs identified in the present study.

We observed that the quality of SV calling and imputation is strongly stratified by variant type, frequency, and the complexity of genomic regions. Using genome stratification files provided by the GIAB consortium, we identified robust SV calls and showed that deletions and insertions with a $MAF > 0.05$ were most reliable. These insights motivate future refinement of variant stratification methods.

We demonstrated the utility of our SV imputation panel by imputing SVs in UKB and conducting SV-WASs on 32 disease-relevant traits. Here, 907 SVs were significantly associated, mapping to 689 functional protein-coding genes. We also ran association analyses of $\sim 1.5k$ plasma protein levels and identified 1720 significant SVs overlapping with 1197 genes. The majority of SVs carrying a significant association with the traits or protein levels were novel, *i.e.*, not detected in 1000 Genomes short-read sequencing. Using selected examples, we highlighted the relevance of the identified potentially causal genes to respiratory diseases.

In the present proof-of-principle study, we focused on SVs overlapping with the coding sequence of genes. In future applications of our SV imputation panel, more refined gene mapping approaches could be employed, *e.g.*, including SVs overlapping enhancer regions or epigenetic marks. Such an enhanced mapping would increase the number of identified associated genes and thus provide additional insights into regulatory mechanisms and disease biology.

Given the sample size of our imputation panel, the resource might be more useful in fine-mapping signals at known GWAS-associated loci than to identify novel disease-associated loci. Our imputation panel facilitates analyses of SVs overlapping exons, splice-sites, promoters, or enhancers of protein-coding genes at GWAS-associated loci. Therefore, it has the potential to become a routine component of post-GWAS gene prioritisation workflows. Structural variants can influence complex disease biology through either the disruption of coding sequence or an altered regulation of gene expression. Such effects may not be well captured by short variants alone. Incorporating SVs into GWAS and follow-up analyses would thus provide more accurate disease risk prediction, uncover underlying pathomechanisms by highlighting actionable pathways and targets, and support precision medicine by providing biomarkers for patient stratification. We also envisage that our SV imputation panel will enable diverse applications of integrating SVs with other *omics* data, *e.g.*, in machine learning-based frameworks. For example, genome-wide SV-based features could be used in models for high-throughput post-GWAS gene prioritisation²⁷, disease subtyping/precision medicine^{35,36}, and drug response/adverse event prediction³⁰.

As next steps in applying our SV panel resource, we suggest, first, to impute SVs in other international biobanks, making more SV data available to the research community. Second, to carry out GWASs utilising SVs to identify novel associations and disease-causal genes. Finally, to carry out comprehensive analyses of the effects of SV inversions (*e.g.*, spanning transcription factor binding sites), translocations, and duplications (*e.g.*, to analyse dosage-dependent effects of SNVs affected by the duplications) for which more research is needed. Thereby, this resource will strongly advance our knowledge of the genetic underpinnings of disease.

Methods

Oxford nanopore long-read sequencing

Sample sequencing was executed using Oxford Nanopore Technologies' (ONT) PromethION P48. After DNA extraction, libraries were prepared according to ONT ligation sequencing kit SQK-LSK110 protocols with r9.4.1 flowcells (further details can be found in the Qiagen Gentra Puregene Handbook).

For base calling, we used guppy v6.0.1³⁷ with the super high accuracy model (SUP). Sequencing statistics for the individual *fastq* files were generated with the NanoStat³⁸ software (Supplementary Table 1 [↗](#); Supplementary Figure 1 [↗](#)).

The data generation process for the complete long-read dataset is detailed in Schloissnig et al.⁷

Alignment and calling

The alignment was performed against the GRCh38 assembly using minimap2 v2.24³⁹ and recommended default parameters for Oxford Nanopore long reads (“-ax map-ont”).

Alignment metrics were computed with Picard's⁴⁰ CollectWgsMetrics tool, with modified parameters to adapt for Oxford Nanopore reads and our analysis setup. Specifically, Picard was instructed to count unpaired reads (“--COUNT_UNPAIRED true”), all base qualities (“--MINIMUM_BASE_QUALITY 0”), and only coverages at sites in the reference on autosomal chromosomes (interval file via “--INTERVALS”). For descriptive statistics see [Supplementary Table 2 \[↗\]\(#\)](#).

Joint variant calling was conducted with Sniffles2⁹ v2.0.7, supplemented by tandem repeat annotations to improve variant calls in these regions, using the default annotation file provided by the tool's authors (“--tandem-repeats human_GRCh38_no_alt_analysis_set.trf.bed”).

Sample identity and quality verification

We performed extensive sample identity verification. We generated long-read ONT sequencing data for 906 DNA samples from the 1000 Genomes project. To confirm the identity of 906 ONT-sequenced individuals, we used the aligned reads to call genotypes for each DNA sample at approximately 12,000 common SNVs (with AF between 0.45 and 0.55) across the genome. We ran ‘bcftools mpileup’ on each sample bam file with the following settings:

```
bcftools mpileup $bam -f $hg38_fasta --skip-indels -a FORMAT/AD -B -R SNV_coordinates.txt -X ont
```

The file ‘SNV_coordinates.txt’ contained the genome coordinates of the SNVs.

Next, we extracted the genotypes of the same SNVs from 3,202 individuals of the 1000 Genomes Project Phase 3 release *vcf* files using ‘bcftools view’. For each ONT-sequenced DNA sample, we then generated a vector of read ratios, where each vector element was a number between 0 and 1 – the number of reads carrying the non-reference base at the SNV site divided by the total number of reads. We converted the 1000 Genomes Project Phase 3 genotypes to comparable numerical vectors by substituting homozygous reference, heterozygous, and homozygous non-reference genotypes with 0, 0.5, and 1, respectively.

To compare the SNV genotypes of ONT-sequenced and 1000 Genomes phase 3 individuals, we calculated the pair-wise distance between ONT and 1000 Genomes phase 3 genotypes. As a distance measure, we used the square of the Euclidean distance divided by the number of SNPs. Since the AFs of the selected SNVs were close to 0.5, the distance measure between unrelated samples is close to 1, around 0 for identical samples, and around 0.5 between first-degree relatives. Moreover, these distances also allowed us to identify contaminated sequencing runs within the ONT data. In case of contamination, the observed distance between DNA samples from supposedly unrelated individuals ranges between 0 and 1, depending on the extent of the contamination.

Sample exclusions and corrections

In total, 18 sequenced DNA samples were removed and the labels of three were corrected: After careful analysis, seven samples were found to be contaminated, and therefore removed from the panel: HG02813 (contaminated with HG02807), HG02870 (contaminated with HG02813), HG02888 (contaminated with HG02870), HG02890 (contaminated with HG02888), HG03778 (contaminated with HG03793), and HG00138 (contaminated with HG0017), and HG02895 (contaminated with HG02890). Furthermore, we found that the DNA sample labelled as HG02807 belonged to individual HG02895 and corrected the label accordingly. Individuals HG01951 and HG01983 were swapped during sequencing and we swapped the labels back.

Seven samples were removed because they were descendants of other individuals in the panel, namely: NA19828, HG03804, HG00420, HG01258, NA19129, NA12877, and NA12878. Note that the SV calls of NA12878 were still used for SV calling benchmarking, see below. Two additional individuals were removed due to sequencing problems, including HG00128, which had 50.1% duplications, and NA18966, which had 34.5% duplications in one of the two *fastq* files available for the sample.

In addition, individual NA15728 was found to be not from the 1000 Genomes dataset and was therefore removed from the panel. Finally, individual HG03127 was removed from the panel due to low quality and low performance metrics in the preliminary leave-one-out analysis. Thus, 18 DNA samples were removed, resulting in 888 individuals included in the imputation reference panel.

Benchmarking SV calls

We generated a benchmark of our SV calls for the NA12878 individual (also known as HG001) against the Genome in a Bottle (GIAB) set of SV calls for the same individual based on high coverage 15kb and 20kb PacBio HIFI sequencing (see https://ftp-trace.ncbi.nlm.nih.gov/giab/ftp/data/NA12878/analysis/PacBio_CCS_15kb_20kb_chemistry2_042021/).

We ran the Truvari bench v4.2.2 tool, with settings ‘--passonly --pctseq 0 --pctsize 0.5 --refdist 500 --dup-to-ins’ to compare our SV calls for NA12878 (the comparison set) to the SV calls included in a file HG001.GRCh38.pbsv.vcf.gz from the link above (the base set). The true positive set accounted for 17,571 variants, while the false positive and false negative sets included 6,899 and 5,464 SVs, respectively. This resulted in a precision of 71.8%, a recall of 76.3%, and an F1 score of 74.0%.

We also generated the benchmark excluding tandem repeat regions, which are known to be challenging for accurate SV calling, because the same SV can be placed in different places in the tandem repeat. The benchmark achieved a very high precision (90.4%), recall (91.5%), and F1 score (91.0%), with true positive, false positive, false negative sets comprising 8486, 897, 789 SVs, respectively. The definition of tandem repeat regions that was used can be found in https://raw.githubusercontent.com/PacificBiosciences/pbsv/master/annotations/human_GRCh38_no_alt_analysis_set.trf.bed. We only excluded tandem repeats longer than 200 bases, there were 172,364 such repeats. The non-tandem repeat regions were included in the Truvari bench run by adding the ‘--includebed human_GRCh38_no_alt_no_trf200.bed.gz --extend 500’ parameters. The *bed* file used here was generated with the command ‘bedtools complement’ for tandem repeats longer than 200 bases.

We also compared our SV calls of NA12878 against the NYGC SV call set based on short-read high-coverage Illumina sequencing from Byrska-Bishop *et al.*¹¹ To obtain the published call set for comparison, we downloaded the VCF files containing the phased SNVs, InDels, and SVs (https://ftp.1000genomes.ebi.ac.uk/vol1/ftp/data_collections/1000G_2504_high_coverage/working/2020422_3202_phased_SNV_INDEL_SV/), extracted SVs (variant IDs starting with ‘HGSV’) present in the NA12878 sample, and ran ‘Truvari bench’ with the same settings as above, now using the short-read-based SV callset as the base set.

The benchmark identified 3,675 true positive SVs, 19,446 false positives, and 630 false negatives, resulting in a precision of 15.9%, a very high recall of 85.4% and an F1 score of 26.8%. This indicates that we successfully detected most SVs identified by short-read sequencing and also identified many additional SVs. Notably, a large proportion (13,127 out of 19,446, or 67.5%) of these additional SVs were also present in the PacBio-based GIAB call set for NA12878 described above, strongly supporting the validity of our calls.

We used the same procedure to benchmark our SV calls for all the 888 samples in the panel, comparing each sample individually with the NYGC SV calls in the same individual. The recall and precision numbers were similar to that of NA12878 (Supplementary Figure 3 [A](#)), with recall rates around 83-87% and precision around 16-17% for the majority of non-African ancestry samples. For African ancestry individuals, recall was around 80% and precision around 18-19%. Excluding SVs in tandem repeats longer than 200 bases from the comparison increased the recall to around 90% (Supplementary Figure 3 [B](#)) and doubled the precision to around 34% for non-African ancestry samples. The difference in the recall and precision rates between African ancestry and other individuals (the two clusters on Supplementary Figures 3A and B [C](#)) can be attributed to the higher number of SVs per individual in case of African ancestry, as shown in Figure 2e [C](#).

Following the SV comparison analysis described above, we annotated each SV in the panel as one of three categories: “novel” – not present in the true positive set for any of the 888 individual-level comparisons with the NYGC-generated SVs; “not compared” – SVs that are ignored by the Truvari bench tool, *i.e.*, break-end (BND)-type SVs or SVs longer than the default size cut-off of 50 kbp; or “known” – SVs classified as true positives in the benchmark for at least one individual.

Individual-level genotype missingness rates

The individual missingness rates exhibited substantial differences mainly due to variation in coverage, with over 10% missing data in 65 samples (Supplementary Figure 4 [C](#)). In this context, we adopted a lenient QC approach for several reasons: Primarily, our aim was to enhance sensitivity and we observed a significant stratification of missingness rates by variant type and region *difficulty* (Supplementary Figure 5 [C](#)). For instance, the sample-wise median missingness rates for all, confident, and difficult regions were 0.42%, 0.11%, and 0.63%, respectively. In contrast, the average missingness rates for the same regions were 2.60%, 2.42%, and 2.74%, respectively.

By adopting this approach, we were able to leverage the data from samples with higher missingness rates in regions where successful genotyping was achieved. Furthermore, we could filter for SVs with elevated missingness rates in our final association results, thereby optimising the use of available data.

Variant-level missingness rates

Overall, raw genotype calls exhibited elevated missingness rates. However, these rates were noticeably stratified by SV type and region type (*i.e.*, difficult and confident regions, see Supplementary Figure 5 [C](#), Supplementary Table 22 [C](#)). Deletions exhibited the highest mean missingness rate at 0.166, closely followed by insertions with a mean missingness rate of 0.086. In contrast, the missingness rates within confident regions were considerably lower, registering at only 0.023 and 0.025 for deletions and insertions, respectively. Other SV types within confident regions also maintained low missingness rates, ranging from 0.027 to 0.029.

We removed SVs with a genotype missingness rate >0.2, thereby excluding 15,830 SVs. Thereafter, the mean missingness rates were uniform across SV types: they ranged from 0.023 to 0.028 in confident regions and had slightly higher values of 0.026 to 0.033 for difficult regions. In both region types, deletions showed the highest and break-ends the lowest call rates.

These results demonstrate that although genotype missingness rates can be elevated in SV calls, problematic calls are predominantly localised within genome sections that are challenging to sequence. Therefore, high missingness rates can be effectively managed by stratifying SVs by region type and excluding genotypes with very low call rates.

False-positive calls of very large SVs generated by Sniffles2

It is known that Sniffles2 (as of version 2.0.7) generates false-positive calls for inversions (INV), duplications (DUP), and deletions (DEL), with apparent lengths up to the size of chromosomes. This issue has been previously reported on the Sniffles2 GitHub issue tracker^{41,42}. Examples in our dataset include a 149 Mbp duplication on chromosome 4, a 131 Mbp inversion on chromosome 1, and a 64.5 Mbp deletion on chromosome 3. These unreasonably large (and hence false-positive) SV calls are a consequence of the internal logic of Sniffles2: Sniffles2 cannot handle edge cases with complex rearrangements, e.g. cut-and-paste insertions with simultaneous inversion of the insertion sequence. It is currently unknown whether insertions (INS) or break-ends (BND) can also be affected by similar issues. Although this problem was discovered because of implausible sizes of some erroneous calls, there is no reason to assume that calls with smaller SV sizes could not be affected, too. We are not aware of any method to verify whether SV calls are affected by this issue, except for direct visual inspection of sequence alignments. We therefore assume that any large SV call-set generated by Sniffles2 is likely to contain a certain proportion of mischaracterised SV calls. Thus, SV calls may require visual or experimental validation.

In our systematic quality control, we removed only Sniffles2 calls considered as highly likely problematic because of their large sizes. Although the maximum expected size of SVs in healthy humans cannot be estimated with certainty, we decided on a lenient threshold of 30 Mbp, guided by a 21.3 Mbp duplication in a healthy human adult, previously reported in the literature⁴³. Only 23 variants were longer than 30 Mbp and subsequently removed (see below).

Reference panel generation

The imputation reference panel was constructed from a combination of the long-read derived SV calls as well as short-read derived SNV and short InDel calls generated from the NYGC's resequencing efforts of the 1000 Genomes project¹¹.

The SV calls were generated as described above, and were filtered to remove variants that did not have the term PASS in the FILTER column of the VCF file, were present in just 1 individual (private singletons or doubletons), were annotated as IMPRECISE, were <50 bp or >30 Mbp, or had genotypes missing in >20% of samples. See Supplementary Table 22 [↗](#) and Supplementary Figures 4 and 5 [↗](#) for statistics on raw and post-QC SV calls. In total, 107,445 SVs remained in the imputation panel after QC.

To generate the short variant calls, we downloaded the sequences for all 3202 samples from the International Genome Sample Resource (IGSR; <https://www.internationalgenome.org/data-portal/data-collection/30x-grch38> [↗](#)), aligned them to GRCh38, and jointly called them using the Sentieon DNaseq pipeline⁴⁴ according to GATK best practices⁴⁵.

We then filtered out InDels that satisfied “INFO/QD < 2.0 || INFO/FS > 200.0 || INFO/ReadPosRankSum < -20.0 || INFO/SOR > 10.0” and SNVs with “INFO/QD < 2.0 || INFO/FS > 60.0 || INFO/MQ < 40.0 || INFO/MQRankSum < -12.5 || INFO/ReadPosRankSum < -8.0”. Next, we subset the filtered, jointly called dataset to the 888 samples with long-read data and filtered out any variants present in <2 individuals.

The filtered long-read-derived SV callset was then merged using bcftools with the filtered short-read-derived short-variant callset, resulting in an unphased VCF file with sporadic missingness. Next, these sporadically missing genotypes were imputed using Beagle4⁴⁶ and phased using Beagle5⁴⁷, resulting in an imputation-ready reference panel VCF file. Default Beagle parameters were used except for the “-gtgl” flag instead of “-gl” during imputation because genotype likelihoods were absent for SVs. Beagle was run on chunks of 55,000 variants each, flanked by 3,000 variants up- and downstream.

Reference panel characterisation

The generated reference panel of 888 unrelated 1000 Genomes Project individuals consisted of 45 million short variants (SNVs and InDels) and 107,445 SVs. In addition to stratification of the SVs by genomic region and variant types, we used frequency classes (MAF bins) for a more nuanced characterisation of SVs in the imputation panel (Supplementary Table 4 [4](#)).

We calculated the average counts of SV minor alleles per individual, stratified by superpopulation, SV type, and frequency class (Supplementary Table 6 [6](#)). Overall, individuals from the African superpopulation had the most diverse SVs, carrying on average 18,822 minor SV alleles, while individuals with East Asian ethnicities were least diverse, showing on average 14,792 SV minor alleles.

We analysed how SVs were shared across populations, stratified by SV type, frequency class, and region type. We counted SVs exclusive to specific populations and SVs shared between either two, three, four, or all five ancestry groups (Supplementary Table 7 [7](#)).

Functional annotation of SVs

This section describes the annotation of SVs as shown in Supplementary Table 5 [5](#).

For the purposes of filtering and functional interpretation of association results, SVs mapping to significant loci were annotated with the GIAB *genome stratification files*¹³ by intersecting bed files containing SV coordinates with genome stratification bed files (intersect -a $\{\text{SVbed}\}$ -b $\{\text{genome_stratification.bed}\}$ -c) and merging the genome stratification labels overlapping with SVs into a single column (column *GIAB_annotations*).

SnpEff¹⁴ v5.1d was used to annotate genes and functional consequences based on canonical transcripts (*command*: snpEff/exec/snpEff GRCh38.105 $\{\text{vcf}\}$ -canon) and SnpSift⁴⁸ v5.1d for extracting fields of interest from the annotated VCF file:

```
gunzip -c  $\{\text{vcf\_annotated}\}$  | snpEff/scripts/vcfEffOnePerLine.pl | snpEff/exec/snpSift extractFields - CHROM POS ID "ANN[*].GENE" "ANN[*].GENEID" "ANN[*].FEATURE" "ANN[*].FEATUREID" "ANN[*].BIOTYPE" "ANN[*].IMPACT" "ANN[*].EFFECT" "ANN[*].ERRORS"
```

For each SV, all fields in the output were merged into one row and added to the table (columns starting with *SNPeff*).

For annotating GWAS catalogue variants, we used the *All associations v1.0.2* file⁴⁹ called *gwas_catalog_v1.0.2-associations_e109_r2023-06-03.tsv* and SnpSift⁵⁰ (*command*: snpEff/exec/snpSift gwasCat -db $\{\text{GWAScatalog}\}$ $\{\text{vcf}\}$). We extracted the fields of interest with bcftools using the query:

```
-f "%ID %GWASCAT_TRAIT %GWASCAT_P_VALUE %GWASCAT_OR_BETA %GWASCAT_REPORTED_GENE %GWASCAT_PUBMED_ID\n"
```

These columns were added to the table with the prefix *GWASCAT*.

In addition, we used the Ensembl BioMart functionality to export regions from the *Ensembl Regulation 109* database, including the datasets "Human Regulatory Features", "Human miRNA Target Regions", and "Human Other Regulatory Regions" (corresponding to PHANTOM regulatory region predictions). The files exported from BioMart were converted to the *bed* format and intersected with SVs (intersect -a $\{\text{SVbed}\}$ -b $\{\text{genome_stratification.bed}\}$ -loj). For each SV, the annotations were added to the table as "regulatory_feature", "predicted_regulatory_feature", and "targeted_by_miRNA", each with the prefix *Ensembl*.

Generating a reduced panel for SV imputation

We created a backbone for imputation by using short variants available in the UK Biobank genotype data. To this end, the genomic positions of 805,426 directly genotyped UK Biobank variants were lifted over from the version GRCh37 to the GRCh38 human genome assembly using the DNAnexus pipeline⁵¹ based on Picard's⁴⁰ LiftOverVcf tool. 803,700 (99.8%) variants were

successfully lifted over, 764,685 of which were SNVs. 702,464 of 764,685 (91.9%) of the genotyped SNVs could be matched to SNVs in our reference panel, serving as the basis for imputation into UKB (see below).

To facilitate imputation of SVs, we generated a reduced panel by removing millions of short variants not directly required for this purpose. The full reference panel containing ~45M short variants and 107,445 SVs was reduced to include only the 702,464 SNVs directly genotyped in UKB, the 107,445 SVs, and 77,350 randomly selected short variants. The 77,350 random short variants (of which 57,733 were SNVs) were imputed along the SVs for benchmarking purposes. The filtering of the panel VCF file was performed using ‘awk’, retaining preannotated genotyped SNVs and SVs and a fraction of additional variants selected using awk’s `rand()` function. All the imputation steps described in the present study were performed using the reduced panel.

Leave-one-out imputation performance

To assess whether the reduced reference panel is suitable for imputation, we performed leave-one-out cross-validation analyses: We excluded one individual from the panel and imputed SVs for this individual using the panel of the remaining 887 samples, applying exactly the same pipeline and settings as those later used for SV imputation into UK Biobank (see below). Then we compared the imputed SV genotypes to the sequenced genotypes of the excluded individual. To this end, we computed genotype concordance, non-reference concordance⁵², minor allele genotype concordance, and r^2_{imp} ⁵³ from allele probabilities produced by Beagle v5.4. We repeated this leave-one-out analysis for each of the 888 individuals. Next, we computed individual-level mean concordances from SVs stratified by MAF, variant, and region class, and aggregated the metrics across superpopulations (Supplementary Table 8 [↗](#)).

We also calculated the variant-level concordance and r^2_{imp} as aggregated imputation performance metrics for SVs (Supplementary Table 9 [↗](#)). Across all 107,445 analysed SVs, the mean genotype concordance was 0.940 and r^2_{imp} = 0.538. However, aggregate metrics are not meaningful for the interpretation of the overall imputation quality. Instead, the imputation quality needs to be stratified by SV type, MAF, and genomic region.

The interpretability of the concordance metric is limited in the presence of rare variants. For a rare allele m , most imputed genotypes have the genotype MM (i.e., the homozygous genotype MM is most common). Consequently, a significant portion of genotypes may match between the ground truth and the imputed data purely by random chance. To better capture the imputation quality specifically of the minor alleles, the non-reference concordance metric was introduced in the literature⁵². Here, only the proportion of correctly imputed genotypes that contain non-reference alleles is assessed, ignoring homozygous reference alleles. However, a flaw in this metric becomes apparent when the reference allele corresponds to the minor allele. This issue can occur in genotype data, depending on the reference sequence and the composition of the cohort. In such cases, the intention of the non-reference concordance metric does not align with the actual computation since genotypes containing major alleles would be considered instead of genotypes containing minor alleles.

Consequently, this conceptual flaw skews the results of the non-reference concordance metric. For example, when calculating the metric on a sample-wise basis, the use of rare variants may lead to seemingly better performance compared to low-frequency variants (Supplementary Table 8 [↗](#)). This misrepresents the fact that it is inherently more challenging to impute rare alleles accurately.

To mitigate this issue, we introduced the minor allele genotype concordance metric. This metric assesses the proportion of correctly imputed genotypes containing minor alleles, providing a more interpretable and intuitive measure compared to the non-reference concordance metric. By explicitly accounting for genotypes with minor alleles, this metric ensures a more accurate and interpretable evaluation of imputation performance, particularly for rare variants (Supplementary Tables 8-9 [↗](#)).

Preprocessing and imputation of SVs into UK Biobank

Imputation of SVs was performed in UK Biobank¹⁹ under the approved research proposal 57952.

Phasing of genotyped SNVs

After running liftover, we converted the output to the *VCF* format. Using PLINK2⁵⁴ v2.00a3.8, we filtered the file to keep only SNVs (`--snps-only`) and aligned the files to a reference file (*Homo_sapiens.GRCh38.dna.primary_assembly.fa.gz* downloaded from Ensembl) using the “`--ref-from-fa`” parameter (to ensure the order of REF and ALT alleles remained compatible with the GRCh38 genome assembly). The resulting dataset of UKB-genotyped SNVs for 488,130 UKB individuals was then phased using SHAPEIT4 v4.2.2⁵⁵. Each chromosome was phased separately using SHAPEIT4⁵⁶ after splitting the chromosomes at the centromere region. We increased the number of Markov Chain Monte Carlo (MCMC) iterations for better phasing performance at cost of computational time, as recommended by the SHAPEIT4 authors (`--mcmc-iterations 10b,1p,1b,1p,1b,1p,1b,1p,10m`) and increased the number of conditioning neighbours in the Positional Burrows-Wheeler Transform (PBWT) to increase accuracy (`--pbwt-depth 8`). The resulting phased *VCF* files were concatenated per chromosome using the *concat* command of bcftools v1.16⁵⁷.

Imputation of UK Biobank data

The imputation of SVs into UKB genotype data was performed with Beagle v5.4⁴⁷. To this end, we converted our reduced panel from *VCF* format to Beagle’s *bref3* format. We could not impute the full set of UKB variants for each chromosome simultaneously due to limited computational resources (memory). Therefore, we divided the data for each chromosome in chunks of 10,000 individuals and then imputed each of these chunks independently. The imputation was conducted with Beagle v5.4 (file: *beagle.22Jul22.46e.jar*; downloaded from <http://faculty.washington.edu/browning/beagle/beagle.html>⁴⁸) using 32 cores (`nthreads=32`) and the genetic maps provided by the Beagle authors (`map="plink.chr${chr}.GRCh38.map"`)⁵⁸. We applied parameters to produce genotype and allele probabilities (`gp=true, ap=true`). Subsequently, the files per chromosome containing the imputed data for 10,000 individuals each were merged using the merge command of bcftools with the option `--merge none`. Finally, r^2_{imp} scores were re-computed based on all individuals, as outlined in the literature⁵³ (Supplementary Tables 10-11⁴⁹).

SV imputation verification using UK Biobank whole-genome sequencing

As a proof-of-principle verification, we used the recently published short-read whole-genome sequencing (WGS) data from UKB¹⁸ to detect the Sniffles2.DEL.3639MF deletion in UKB individuals. The read coverage was extracted from the population-level *VCF* files (UKB field 24310) by adding up the local allelic depth numbers (‘LAD’ field in the *VCF* files). We saw a clear drop in read coverage in the SV region chr16:75395953-75396795 in individuals that had one or both alleles deleted (Supplementary Figure 7⁵⁰). We compared the mean coverage at the SV site to the mean coverage on the flanking regions to detect the deletion genotype in each UKB individual. The individuals clustered very well into three groups, corresponding to 0, 1 or 2 deleted alleles (see Supplementary Figure 8⁵¹). We compared this genotype to the imputed genotype for 485,129 samples present in both the WGS and imputed datasets. Excluding 1553 individuals for which the WGS calls were not confident (the ratio of coverage at and around Sniffles2.DEL.3639MF was in between the “no deletion” and “one allele deleted” clusters), the genotypes coincided for 98.7% of individuals (477086 out of 483576), exceeding our accuracy estimations based on the leave-one-out analysis (concordance of 92.3%). We only used this simple read coverage method to verify one imputed SV since it requires a significant manual effort and can only be used to detect deletions in non-repetitive areas of the genome. A full-scale SV calling from the UKB short-read WGS data is beyond the scope of this study.

Genome-wide association study utilising imputed structural variants in UK Biobank

UK Biobank sample selection for association analyses

We performed SV genome-wide association analyses (SV-WAS) in UKB participants of European ancestry. For the selection of samples with “White European” ancestry, we adhered to a strategy previously established in the literature^{59,60}. In summary, we applied k-means clustering (with $k=6$) on the first two genetic principal components from the principal component analysis (PCA) results provided by UKB (“ukb_sqc_v2_pca.txt”). We excluded individuals who had withdrawn consent or did not map to the European cluster (Supplementary Figure 9⁶¹). We identified 171 duplicate pairs among the selected samples using KING v2.3.0 (--related). For each pair, we excluded the sample with the most missing phenotypes of interest. In cases where no significant differences were found, we randomly selected a sample from the pair to exclude. This resulted in a dataset of 455,589 samples for association analyses.

SV-wide association studies in UK Biobank

SV-WASs and SV-based pQTLs were calculated in UKB using imputed SV dosages and a linear mixed model, incorporating a leave-one-chromosome-out whole-genome regression model to account for population stratification as implemented in Regenie v3⁶¹. The effective number of individuals included in each analysis depended on the availability of phenotype and protein-level measurements. The number varied from 363K to 454K for quantitative and binary phenotypes (see Supplementary Tables 12 and 13⁶²) and was 48,205 for pQTL analyses. We computed the ancestry principal components separately for each phenotype set of individuals and for the pQTL set of individuals using PLINK2 with parameters ‘--pca approx 20 --chr 1-22 --snps-only’ based on a pruned set of 255,369 directly genotyped UKB variants. The pruning was also done with PLINK2 using ‘--maf 0.005 --indep-pairwise 200kb 0.5’ on the set of all directly genotyped UKB variants.

This set of 255,369 pruned variants was used in the first step of the Regenie run. The VCF files of imputed variants generated by Beagle were first filtered for SVs and then transformed to the PLINK2 format (pgen) using the PLINK2 command ‘--make-pgen --vcf \$filtered_vcf dosage=DS’, with the dosage information kept. Regenie uses dosages for association analysis whenever they are available. The first and the second step of Regenie were run on each phenotype and each of the 1463 protein levels with the following settings: ‘--bsize 1000 --loocv --gz --lowmem --lowmem-prefix’ for the first step, and ‘--bsize 200 --firth --firth-se --approx --pThresh 0.05 --minMAC 3’ at the second step. In addition, the ‘--bt’ flag was used for binary phenotypes and ‘--qt’ otherwise. For all 19 quantitative phenotypes, the measurement at the first examination was used to maximise sample size. Quantitative phenotype values and protein levels were transformed with rank-based inverse normal transformation in R (‘function(x) qnorm((rank(x, na.last = “keep”) - 0.5) / sum(!is.na(x)))’) before association analyses. We used the following covariates: sex and genotype array (UKB data field 22000) as binary covariates, age, age², and the first 20 genetic principal components as quantitative covariates. For lung function-related traits (FEV₁, FVC, FEV₁/FVC), ‘ever smoked’ (data field 20160) was added as a binary covariate. Note that the Regenie output includes the statistical attributes of the association test for each variant, together with the allele frequency and INFO score for the set of participants in the given association analysis.

In the exploratory SV-WAS, we used the standard threshold for genome-wide significance of $p < 5 \times 10^{-8}$. For the pQTL analyses, we applied Bonferroni correction for multiple testing on top of that genome-wide threshold, correcting for the number of tested protein levels ($n=1463$): $p < 5 \times 10^{-8} / 1463 = 3.4 \times 10^{-11}$. We extracted 4,528 significant ($p < 5 \times 10^{-8}$) SV associations in all the SV-WASs performed, which were all included in Supplementary Table 14⁶³. Filtering down by Regenie metrics INFO>0.7 and MAF>0.01, 3,858 SV associations remained, corresponding to 1,898 unique SVs.

A similar table for pQTL SV associations (Supplementary Table 15⁶⁴) includes 23,810 associations at the $p < 5 \times 10^{-8}$ significance level. Using a Bonferroni corrected p -value threshold $p < 3.4 \times 10^{-11}$ and the INFO>0.7, MAF>0.01 criteria, 10,518 significant SV associations remained (for 3,723 unique SVs).

for 1,101 proteins.

Signal selection and conditional analyses

For each phenotype and protein level, the significant SVs (GWAS p -value $< 5 \times 10^{-8}$; pQTL p -value $< 5 \times 10^{-8}/1463$) were split into sets separated by a genomic distance ≥ 100 kb. For each set, we performed conditional analyses: First, SV coordinates were lifted from GRCh38 to GRCh37, then the SV imputed data was merged with the standard UKB short-variant imputed genotype data (data field 22828) at the locus with ± 500 kb flanking regions on each side. Next, we ran an iterative conditional analysis on the merged dataset to select independent association signals, conditioning on the set of top variants and adding the identified independent variants in a stepwise manner, stopping when the p -value of the top variant reached $> 5 \times 10^{-6}$ ($> 5 \times 10^{-6}/1463$ for pQTLs) or after five iterations. We considered the SV as an independent and novel signal if it appeared as a top variant in any iteration of the conditional analysis.

Utilising SV information to identify novel disease-relevant genes (SV2G) and improve post-GWAS locus-to-gene (L2G) prioritisation

From identified SV associations to causal genes (SV2G)

We conducted the following steps to identify the putatively causal gene(s) in the associated regions where an SV overlapping a protein-coding gene was the top variant (see ‘SV2G’ in Figure 1b): a) Examination of the literature, b) cis-e/pQTL-based phenome-wide Mendelian randomisation (MR) analyses (using the TwoSampleMR package⁶²), c) finemapping and colocalisation (coloc v5)⁶³, for all genes implicated by the SVs at each locus. GCTA-COJO⁶⁴ and LD pruning ($r^2 < 0.01$) were used to select cis-e/pQTLs (within ± 1 Mbp of the gene’s transcription start site) as instrumental variables (IVs) for each exposure (list of exposures in Supplementary Table 23a). Where we identified a strong MR association between gene expression and a trait (e.g., between *MEGF6* expression and a lung-disease-relevant trait), we manually checked the regional Manhattan plots to ensure that there was strong evidence for local colocalisation between the cis-e/pQTL signals used as the MR instruments and the outcome/trait GWAS (examples in Supplementary Figures 10–14). The outcome GWASs ($n = 7,429$; Supplementary Table 23b) consisted mostly of a manually curated list from IEU OpenGWAS⁶⁵ but also from internally conducted GWASs on 44 endophenotypes and outcomes related to respiratory and/or fibrotic diseases^{66,67}. A Bayesian method was used to fine-map each associated locus to a set of variants that – assuming the causal variant was also included in the analysis – contains the underlying causal variant with 95% probability⁶⁸. We set the parameter W (i.e., the variance of the prior distribution of effect sizes) to 0.04 ($\approx 0.21^2$) in the approximate Bayes factor formula – which equates to a 95% belief that the absolute relative risk is $< \frac{3}{2}$. To estimate the probability that a single variant explains both the cis-e/pQTL signal and the signal in the trait/outcome GWAS, we manually inspected the regional Manhattan plots in addition to using the “coloc.abf” function⁶⁹.

From SNV-based GWAS signals to causal genes utilising SV information (L2G)

To systematically analyse the contribution of SV information to post-GWAS locus-to-gene (see ‘L2G’ in Figure 1b) prioritisation approaches, we utilised the list of independent associations identified by the most recent GWAS of lung function carried out by Shrine *et al.*²⁰ and the table of implicated genes (consisting of genes implicated by various post-GWAS methods such as e/pQTL association, PoPS²¹, and rare respiratory disease causal genes) for each associated locus (Column AB in Supplementary Table 18). First, we lifted over the SNV coordinates reported by Shrine *et al.* (in GRCh37) to GRCh38 positions. We then looked for SVs that were (i) within 500 kb of the independent SNVs reported by Shrine *et al.*, (ii) significantly associated ($p < 5 \times 10^{-8}$ in our UKB-based SV-WAS) with the primary lung function measure reported by Shrine *et al.*²⁰ (e.g. FEV₁/FVC for rs2355210), and (iii) which spanned protein-coding genes (as defined by GENCODE v43). As a pragmatic approach to focus on the SVs likely to be real and functional, we restricted the analyses to a set of well-imputed SVs (INFO metric > 0.7) not mapping to ‘difficult’ regions.

Data availability

Raw SV calls, the long-read sequencing-based SV imputation panel, and the SV summary statistics from 32 SV-wide association studies are available through the OpnMe initiative of Boehringer Ingelheim GmbH (<https://opnme.com/genomiclens>). The raw long-read sequencing data (FASTQ files) for the 1000 Genomes Project samples included in this study are accessible via the European Nucleotide Archive under accession number PRJEB89727 (<https://www.ebi.ac.uk/ena/browser/view/PRJEB89727>). The dataset analysed here constitutes a subset of this broader collection. Imputed SVs for UK Biobank participants will be made available through the UKB Research Analysis Platform (RAP).

Acknowledgements

This research has been conducted using the UK Biobank, a major biomedical database (www.ukbiobank.ac.uk). We thank all UK Biobank and 1000 Genomes Project participants, without whom this project would not have been possible.

We express our gratitude to the MARVL Initiative – a collaboration between the Research Institute of Molecular Pathology (IMP), BI X and gCBDS (Boehringer Ingelheim). In particular: Siegfried Schloissnig, Klaus Ehrlinger, Julien Charest, Mila Asparuhova and Patrick Hüther for sequencing the samples.

We thank Gilean McVean and Jan Korbel for guidance during the project and providing critical feedback on the manuscript. We would like to thank the anonymous reviewers of the first version of this manuscript.

Additional information

Author contributions

Study conception and design: ZD, JS, SO, AME, BN, JNJ, NP, JK. Data governance/infrastructure, statistical analysis, and/or interpretation: BN, JS, SO, DD, AME, TFMA, CM, JCBL, BAB, CB, LS, SM, AK, JHL, GMB, IB. Preparation of the manuscript: ZD, AME, BN, TFMA, DD, SO, JS, JKP, JHL, GMB, JA, IB. All authors (incl. all under the ‘gCBDS’ banner) have critically reviewed and approved the final version of this paper, including the authorship statement.

‘Boehringer Ingelheim – Global Computational Biology and Digital Sciences’ authors

Johann de Jong, Fidel Ramirez, Frank Li, Thomas Kandl, James Cai, George Okafo.

Address: Global Computational Biology and Digital Sciences (gCBDS), Boehringer Ingelheim Pharma GmbH & Co. KG, Biberach an der Riss, Germany.

Ethics declarations

This research has been conducted using the UK Biobank Resource under Application Number 57952.

Author ORCID iDs

Zhihao Ding:  <https://orcid.org/0000-0003-0961-2284>

Author notes

Boris Noyvert, A Mesut Erzurumluoglu, Dmitriy Drichel, Steffen Omland, Till FM Andlauer:

These authors contributed equally to this work

Jan Stutzki, Zhihao Ding: Joint senior authors

Competing interests: Boehringer Ingelheim, a privately-owned pharmaceutical company, funded this initiative. DD and LS are independent contractors and declared no conflicts of interest. GMB,

JHL, and JKP are employees of Gencove and declared no conflicts of interest.
 Funding Statement: This study was funded by Boehringer Ingelheim, a privately-owned pharmaceutical company.

Additional files

[Supplementary Figures and Tables](#)

[Supplementary Tables 1-2, 5, 12-21, 23-25](#)

References

1. Sedlazeck FJ, et al. (2018) Accurate detection of complex structural variations using single-molecule sequencing. *Nat Methods* **15**:461-468 <https://doi.org/10.1038/s41592-018-0001-7> | PubMed
2. Jakubosky D, et al. (2020) Properties of structural variants and short tandem repeats associated with gene expression and complex traits. *Nat Commun* **11**:2927 <https://doi.org/10.1038/s41467-020-16482-4> | PubMed
3. Sudmant PH, et al. (2015) An integrated map of structural variation in 2,504 human genomes. *Nature* **526**:75-81 <https://doi.org/10.1038/nature15394> | PubMed
4. Beyter D, et al. (2021) Long-read sequencing of 3,622 Icelanders provides insight into the role of structural variants in human diseases and other traits. *Nat Genet* **53**:779-786 <https://doi.org/10.1038/s41588-021-00865-4> | PubMed
5. Wu Z, et al. (2021) Structural variants in the Chinese population and their impact on phenotypes, diseases and population adaptation. *Nat Commun* **12**:6501 <https://doi.org/10.1038/s41467-021-26856-x> | PubMed
6. Abel HJ, et al. (2020) Mapping and characterization of structural variation in 17,795 human genomes. *Nature* **583**:83-89 <https://doi.org/10.1038/s41586-020-2371-0> | PubMed
7. Schloissnig S, et al. (2024) Long-read sequencing and structural variant characterization in 1,019 samples from the 1000 Genomes Project. *bioRxiv* 2024.04.18.590093 <https://doi.org/10.1101/2024.04.18.590093> | PubMed
8. Gustafson JA, et al. (2024) High-coverage nanopore sequencing of samples from the 1000 Genomes Project to build a comprehensive catalog of human genetic variation. *Genome Res* **34**:gr.279273.124 <https://doi.org/10.1101/gr.279273.124> | PubMed
9. Smolka M, et al. (2024) Detection of mosaic and population-level structural variants with Sniffles2. *Nat. Biotechnol* **42**:1571-1580 <https://doi.org/10.1038/s41587-023-02024-y> | PubMed
10. PacBio HIFI long-read sequencing (no date) https://ftp-trace.ncbi.nlm.nih.gov/giab/ftp/data/NA12878/analysis/PacBio_CCS_15kb_20kb_chemistry2_042021/
11. Byrska-Bishop M, et al. (2022) High-coverage whole-genome sequencing of the expanded 1000 Genomes Project cohort including 602 trios. *Cell* **185**:3426-3440.e19 <https://doi.org/10.1016/j.cell.2022.08.004> | PubMed
12. English AC, Menon VK, Gibbs RA, Metcalf GA, Sedlazeck FJ (2022) Truvari: refined structural variant comparison preserves allelic diversity. *Genome Biol* **23**:271 <https://doi.org/10.1186/s13059-022-02840-6> | PubMed
13. GIAB consortium (2022) genome-stratifications. GitHub. <https://github.com/genome-in-a-bottle/genome-stratifications>
14. Cingolani P, et al. (2012) A program for annotating and predicting the effects of single nucleotide polymorphisms, SnpEff. *Fly* **6**:80-92 <https://doi.org/10.4161/fly.19695> | PubMed
15. Welter D, et al. (2014) The NHGRI GWAS Catalog, a curated resource of SNP-trait associations. *Nucleic Acids Res* **42**:D1001-D1006 <https://doi.org/10.1093/nar/gkt1229> | PubMed
16. Abel HJ, et al. (2020) Mapping and characterization of structural variation in 17,795 human genomes. *Nature* **583**:83-89 <https://doi.org/10.1038/s41586-020-2371-0> | PubMed

17. Linderman MD, et al. (2014) Analytical validation of whole exome and whole genome sequencing for clinical applications. *Bmc Med Genomics* **7** <https://doi.org/10.1186/1755-8794-7-20> | PubMed
18. Li S, Carss KJ, Halldorsson BV, Cortes A, Consortium UBW.-GS (2023) Whole-genome sequencing of half-a-million UK Biobank participants. *medRxiv* 2023.12.06.23299426 <https://doi.org/10.1101/2023.12.06.23299426>
19. Bycroft C, et al. (2018) The UK Biobank resource with deep phenotyping and genomic data. *Nature* **562**:203-209 <https://doi.org/10.1038/s41586-018-0579-z> | PubMed
20. Shrine N, et al. (2023) Multi-ancestry genome-wide association analyses improve resolution of genes and pathways influencing lung function and chronic obstructive pulmonary disease risk. *Nat Genet* **55**:410-422 <https://doi.org/10.1038/s41588-023-01314-0> | PubMed
21. Weeks EM, et al. (2023) Leveraging polygenic enrichments of gene features to predict genes underlying complex traits and diseases. *Nat. Genet* **55**:1267-1276 <https://doi.org/10.1038/s41588-023-01443-6> | PubMed
22. Leigh MW, et al. (2009) Clinical and genetic aspects of primary ciliary dyskinesia/Kartagener syndrome. *Genet Med* **11**:473-487 <https://doi.org/10.1097/gim.0b013e3181a53562> | PubMed
23. Liu D, et al. (2014) Proteomic Analysis of Lung Tissue in a Rat Acute Lung Injury Model: Identification of PRDX1 as a Promoter of Inflammation. *Mediat Inflamm* **2014**:469358 <https://doi.org/10.1155/2014/469358> | PubMed
24. Sun BB, et al. (2023) Plasma proteomic associations with genetics and health in the UK Biobank. *Nature* **622**:329-338 <https://doi.org/10.1038/s41586-023-06592-6> | PubMed
25. Jiang T, et al. (2022) Fibroblast growth factor 10 attenuates chronic obstructive pulmonary disease by protecting against glycocalyx impairment and endothelial apoptosis. *Respir Res* **23**:269 <https://doi.org/10.1186/s12931-022-02193-5> | PubMed
26. Nichols CE, et al. (2020) Lrp1 Regulation of Pulmonary Function. Follow-Up of Human GWAS in Mice. *Am J Resp Cell Mol* **64**:368-378 <https://doi.org/10.1165/rcmb.2019-0444oc> | PubMed
27. Mountjoy E, et al. (2021) An open approach to systematically prioritize causal variants and genes at all published human GWAS trait-associated loci. *Nat Genet* **53**:1527-1533 <https://doi.org/10.1038/s41588-021-00945-5> | PubMed
28. Auwerx C, et al. (2022) The individual and global impact of copy-number variants on complex human traits. *The Am J Hum Genet* **109**:647-668 <https://doi.org/10.1016/j.ajhg.2022.02.010> | PubMed
29. Nagai A, et al. (2017) Overview of the BioBank Japan Project: Study design and profile. *J Epidemiol* **27**:S2-S8 <https://doi.org/10.1016/j.je.2016.12.005> | PubMed
30. Walters RG, et al. (2022) Genotyping and population structure of the China Kadoorie Biobank. *medRxiv* 2022.05.02.22274487 <https://doi.org/10.1101/2022.05.02.22274487>
31. Wong E, et al. (2023) The Singapore National Precision Medicine Strategy. *Nat Genet* **1**-9 <https://doi.org/10.1038/s41588-022-01274-x> | PubMed
32. Auton A, et al. (2015) A global reference for human genetic variation. *Nature* **526**:68-74 <https://doi.org/10.1038/nature15393> | PubMed
33. Pereira L, Mutesa L, Tindana P, Ramsay M (2021) African genetic diversity and adaptation inform a precision medicine agenda. *Nat. Rev. Genet* **22**:284-306 <https://doi.org/10.1038/s41576-020-00306-8> | PubMed
34. Keinan A, Mullikin JC, Patterson N, Reich D (2007) Measurement of the human allele frequency spectrum demonstrates greater genetic drift in East Asians than in Europeans. *Nat. Genet* **39**:1251-1255 <https://doi.org/10.1038/ng2116> | PubMed
35. Chandak P, Huang K, Zitnik M (2023) Building a knowledge graph to enable precision medicine. *Sci Data* **10**:67 <https://doi.org/10.1038/s41597-023-01960-3> | PubMed
36. Jong J. de, et al. (2021) Towards realizing the vision of precision medicine: AI based prediction of clinical drug response. *Brain* **144**:1738-1750 <https://doi.org/10.1093/brain/awab108> | PubMed

37. Oxford Nanopore Technologies plc (no date) Guppy. <https://nanoporetech.com/document/Guppy-protocol>
38. Coster WD, D'Hert S, Schultz DT, Cruts M, Broeckhoven CV (2018) NanoPack: visualizing and processing long-read sequencing data. *Bioinformatics* **34**:2666-2669 <https://doi.org/10.1093/bioinformatics/bty149> | PubMed
39. Li H (2018) Minimap2: pairwise alignment for nucleotide sequences. *Bioinformatics* **34**:3094-3100 <https://doi.org/10.1093/bioinformatics/bty191> | PubMed
40. Picard Toolkit (2019) <https://broadinstitute.github.io/picard/>
41. ddrichel (2020) Sniffles issue #235. GitHub. <https://github.com/fritzsedlazeck/Sniffles/issues/235>
42. ddrichel (2023) Sniffles issue #387. GitHub. <https://github.com/fritzsedlazeck/Sniffles/issues/387>
43. Consortium, Gte, et al. (2017) The impact of structural variation on human gene expression. *Nat Genet* **49**:692-699 <https://doi.org/10.1038/ng.3834> | PubMed
44. Freed D, Aldana R, Weber JA, Edwards JS (2017) The Sentieon Genomics Tools - A fast and accurate solution to variant calling from next-generation sequence data. *bioRxiv* 115717 <https://doi.org/10.1101/115717>
45. Borad Institute (no date) GATK: Best Practices Workflows. <https://gatk.broadinstitute.org/hc/en-us/sections/360007226651-Best-Practices-Workflows>
46. Browning SR, Browning BL (2007) Rapid and Accurate Haplotype Phasing and Missing-Data Inference for Whole-Genome Association Studies By Use of Localized Haplotype Clustering. *Am. J. Hum. Genet* **81**:1084-1097 <https://doi.org/10.1086/521987> | PubMed
47. Browning BL, Zhou Y, Browning SR (2018) A One-Penny Imputed Genome from Next-Generation Reference Panels. *Am J Hum Genetics* **103**:338-348 <https://doi.org/10.1016/j.ajhg.2018.07.015> | PubMed
48. Cingolani P, et al. (2012) Using Drosophila melanogaster as a Model for Genotoxic Chemical Mutational Studies with a New Program, SnpSift. *Front. Genet* **3**:35 <https://doi.org/10.3389/fgene.2012.00035> | PubMed
49. EMBL EBI (no date) GWAS Catalog. <https://www.ebi.ac.uk/gwas/docs/file-downloads>
50. Cingolani P, et al. (2012) Using Drosophila melanogaster as a Model for Genotoxic Chemical Mutational Studies with a New Program, SnpSift. *Front. Genet* **3**:35 <https://doi.org/10.3389/fgene.2012.00035> | PubMed
51. DNAnexus (2022) liftover_plink_beds. GitHub. https://github.com/dnanexus-rnd/liftover_plink_beds
52. Linderman MD, et al. (2014) Analytical validation of whole exome and whole genome sequencing for clinical applications. *Bmc Med Genomics* **7**:20 <https://doi.org/10.1186/1755-8794-7-20> | PubMed
53. Das S, Abecasis GR, Browning BL (2018) Genotype Imputation from Large Reference Panels. *Annu Rev Genom Hum G* **19**:1-24 <https://doi.org/10.1146/annurev-genom-083117-021602> | PubMed
54. Chang CC, et al. (2015) Second-generation PLINK: rising to the challenge of larger and richer datasets. *Gigascience* **4**:1-16 <https://doi.org/10.1186/s13742-015-0047-8> | PubMed
55. Delaneau O, Zagury J.-F, Robinson MR, Marchini JL, Dermitzakis ET (2019) Accurate, scalable and integrative haplotype estimation. *Nat Commun* **10**:5436 <https://doi.org/10.1038/s41467-019-13225-y> | PubMed
56. Delaneau O (2026) SHAPEIT 4 genetic maps. GitHub. <https://github.com/odelaneau/shapeit4/tree/master/maps>
57. Danecek P, et al. (2021) Twelve years of SAMtools and BCFtools. *Gigascience* **10**:giab008 <https://doi.org/10.1093/gigascience/giab008> | PubMed
58. Browning B (no date) Beagle genetic maps. https://bochet.gcc.biostat.washington.edu/beagle/genetic_maps/

59. Wain LV, et al. (2017) Genome-wide association analyses for lung function and chronic obstructive pulmonary disease identify new loci and potential druggable targets. *Nat Genet* **49**:416-425 <https://doi.org/10.1038/ng.3787> | PubMed
 60. Shrine N, et al. (2019) New genetic signals for lung function highlight pathways and chronic obstructive pulmonary disease associations across multiple ancestries. *Nat Genet* **51**:481-493 <https://doi.org/10.1038/s41588-018-0321-7> | PubMed
 61. Mbatchou J, et al. (2021) Computationally efficient whole-genome regression for quantitative and binary traits. *Nat Genet* **53**:1097-1103 <https://doi.org/10.1038/s41588-021-00870-7> | PubMed
 62. Hemani G, et al. (2018) The MR-Base platform supports systematic causal inference across the human phenome. *eLife* **7**:e34408 <https://doi.org/10.7554/eLife.34408> | PubMed
 63. Wallace C (2021) A more accurate method for colocalisation analysis allowing for multiple causal variants. *PLoS Genet* **17**:e1009440 <https://doi.org/10.1371/journal.pgen.1009440> | PubMed
 64. Yang J, et al. (2012) Conditional and joint multiple-SNP analysis of GWAS summary statistics identifies additional variants influencing complex traits. *Nat. Genet* **44**:369-375 <https://doi.org/10.1038/ng.2213> | PubMed
 65. Elsworth B, et al. (2020) The MRC IEU OpenGWAS data infrastructure. *Biorxiv* <https://doi.org/10.1101/2020.08.10.244293>
 66. Hemani G, Tilling K, Smith GD (2017) Orienting the causal relationship between imprecisely measured traits using GWAS summary data. *PLoS Genet* **13**:e1007081 <https://doi.org/10.1371/journal.pgen.1007081> | PubMed
 67. Bowden J, Smith GD, Burgess S (2015) Mendelian randomization with invalid instruments: effect estimation and bias detection through Egger regression. *Int. J. Epidemiology* **44**:512-525 <https://doi.org/10.1093/ije/dyv080> | PubMed
 68. Wakefield J (2007) A Bayesian Measure of the Probability of False Discovery in Genetic Epidemiology Studies. *Am J Hum Genetics* **81**:208-227 <https://doi.org/10.1086/519024> | PubMed
 69. Giambartolomei C, et al. (2014) Bayesian Test for Colocalisation between Pairs of Genetic Association Studies Using Summary Statistics. *PLoS Genet* **10**:e1004383 <https://doi.org/10.1371/journal.pgen.1004383> | PubMed
- Noyvert B, Erzurumluoglu AM, Drichel D, Omland S, Andlauer TFM, Mueller S, Sennels L, Becker C, Kantorovich A, Bartholdy BA, et al. (2023) Structural variant multi-ancestry imputation panel. OpnMe, Genomic Lens. <https://opnme.com/genomiclens>
- Schloissnig S, Pani S, Ebler J, Hain C, Tsapalou V, Söylev A, Hüther P, Ashraf H, Prodanov T, Asparuhova M, et al. (2025) Structural variation in 1,019 diverse humans based on long-read sequencing. ENA. ID PRJEB89727 <https://www.ebi.ac.uk/ena/browser/view/PRJEB89727>

Peer reviews

Reviewer #1 (Public review):

Summary:

The authors sequenced 888 individuals from the 1000 Genomes Project using the Oxford Nanopore long-read sequencing method to achieve highly sensitive, genome-wide detection of structural variants (SVs) at the population level. They conducted solid benchmarking of SV calling and systematically characterized the identified SVs. While short-read sequencing methods, including those used in the 1000 Genomes Project, have been widely applied, they exhibit high accuracy in detecting single nucleotide variants (SNVs) and small insertions and deletions but have limited sensitivity for SV detection. This study significantly enhances SV detection capabilities, establishing it as a valuable resource for human genetic research. Furthermore, the authors constructed an SV imputation panel using the generated data and imputed SVs in 488,130 individuals from the UK Biobank. They then conducted a proof-of-

principle genome-wide association study (GWAS) analysis based on the imputed SVs and selected traits within the UK Biobank. Their findings demonstrate that incorporating SV-GWAS analysis provides additional insights beyond conventional GWAS frameworks focusing on SNVs, particularly in improving fine-mapping.

Strengths:

The authors constructed a high-sensitivity reference panel of genome-wide SVs at the population level, addressing a critical gap in the field of human genetics. This resource is expected to significantly advance research in human genetics. They demonstrated the imputation of SVs in individuals from the UK Biobank using this panel and conducted a proof-of-concept SV-based GWAS. Their findings highlight a novel and effective strategy for integrating SVs into GWAS, which will facilitate the analysis of human genetic data from the UK Biobank and other datasets. Their conclusions are supported by comprehensive analyses.

Weaknesses:

The authors have addressed many of my previous comments, and I appreciate their efforts. However, I still have two related concerns.

(1) Shortly after reviewing this manuscript last year, my laboratory obtained access to the UK Biobank (UKB) Tier 3 dataset for an unrelated project. In August 2025, I searched the UKB Research Analysis Platform (UKB-RAP) for the imputed structural variant (SV) dataset described in this manuscript but was unable to locate it. After contacting UKB, I was informed that they were developing the system for releasing the data. To the best of my knowledge, the dataset remains unavailable. A major contribution of this work is the generation of an imputed SV resource for approximately 500,000 UKB participants with extensive phenotypic information. If this resource is not accessible to the research community, even to authorized UKB users, the practical impact and utility of the study are substantially diminished.

(2) Given that the imputed SV dataset is currently unavailable, it becomes even more important for the authors to provide a detailed, ready-to-run SV imputation pipeline for UKB-RAP, even if the "data processing simply consisted of running standard bioinformatics tools with the parameters exactly as described in the manuscript". In particular, the pipeline should include practical information such as computational requirements (e.g., memory and storage), expected running time, and estimated cost. Anyone with experience using UKB-RAP will agree that reproducing large-scale analyses on the platform can be both technically complex and financially expensive. Such pipeline would greatly improve the reproducibility and accessibility of this work.

Because my initial assessment of the manuscript was generally positive, I do not wish to change my overall evaluation, summary, or assessment of its strengths. However, I would view the work even more favorably if either (i) the imputed SV dataset became publicly available to authorized UKB users, or (ii) the authors extended their SV-GWAS analyses to the full range of UKB phenotypes and released the resulting summary statistics, analogous to the Pan UKBB ("<https://pan.ukbb.broadinstitute.org/>") resource. Although this would require considerable additional effort and computational resources, it would substantially enhance the long-term value and impact of the study.

Finally, I would like to emphasize that these comments are not intended to create unnecessary difficulties for the authors or the editors. Rather, I believe this highlights a broader issue in the use of this kind of large public datasets: reviewers cannot independently verify key results, and readers cannot readily build upon the work if the underlying resources are inaccessible, even after obtaining authorized access to the original dataset. I hope the authors, together with the eLife editors and UK Biobank where appropriate, can help facilitate the timely release of this valuable resource.

<https://doi.org/10.7554/eLife.106115.2.sa1>

Author response:

The following is the authors' response to the original reviews.

Our revision includes:

- (1) The generated data from long-read whole-genome sequencing of 1000 Genomes Project samples, including FASTQ files, SV calls, and the imputation panel, are now openly available via ENA and OpnMe. The imputed structural variant data have been submitted to UK Biobank for release through the UK Biobank Research Analysis Platform, subject to UK Biobank release procedures. SV-WAS summary statistics have been made available via OpnMe.
- (2) Clarification of analyses and methods, addition of two new Supplementary Figures, and correction of minor issues throughout the manuscript.
- (3) A significantly expanded Discussion to address the reviewers' comments and better contextualise our methods and results.

eLife Assessment

This fundamental work significantly enhances our understanding of how structural variants influence human phenotypes. The conclusion is convincingly supported by rigorous analyses of long-read sequencing data. If the raw data are made publicly available, these high-quality datasets and findings will further advance our knowledge of genetic variation in the human population.

We thank the editors for this positive assessment of our work. The raw long-read sequencing data (FASTQ files) can now be accessed through the European Nucleotide Archive (ENA) under accession number PRJEB89727, as part of a larger collection of 1019 sequenced probands from the 1000 Genomes Project (<https://www.ebi.ac.uk/ena/browser/view/PRJEB89727>). The generated imputation panel and the structural variant calls, based on the 888 probands used in the present manuscript, remain freely available for download at <https://opnme.com/genomiclens>. We have now added the summary statistics of 32 SV-wide association studies to the same resource. In addition, we have submitted the imputed SV genotypes for UK Biobank participants to the UK Biobank; once processed by UK Biobank, these genotypes will be released via the UK Biobank Research Analysis Platform (RAP).

Public Reviews:

Reviewer #1 (Public review):

Summary:

The authors sequenced 888 individuals from the 1000 Genomes Project using the Oxford Nanopore long-read sequencing method to achieve highly sensitive, genome-wide detection of structural variants (SVs) at the population level. They conducted solid benchmarking of SV calling and systematically characterized the identified SVs. While short-read sequencing methods, including those used in the 1000 Genomes Project, have been widely applied, they exhibit high accuracy in detecting single nucleotide variants (SNVs) and small insertions and deletions but have limited sensitivity for SV detection. This study significantly enhances SV detection capabilities, establishing it as a valuable resource for human genetic research. Furthermore, the authors constructed an SV imputation panel using the generated data and imputed SVs in 488,130 individuals from the UK Biobank. They then conducted a proof-of-principle genome-wide association study

(GWAS) analysis based on the imputed SVs and selected traits within the UK Biobank. Their findings demonstrate that incorporating SV-GWAS analysis provides additional insights beyond conventional GWAS frameworks focusing on SNVs, particularly in improving fine mapping.

The authors constructed a high-sensitivity reference panel of genome-wide SVs at the population level, addressing a critical gap in the field of human genetics. This resource is expected to significantly advance research in human genetics. They demonstrated the imputation of SVs in individuals from the UK Biobank using this panel and conducted a proof-of-concept SV-based GWAS. Their findings highlight a novel and effective strategy for integrating SVs into GWAS, which will facilitate the analysis of human genetic data from the UK Biobank and other datasets. Their conclusions are supported by comprehensive analyses.

We thank the reviewer for highlighting the value of our SV imputation reference panel.

Weaknesses:

(1) Although the authors employ state-of-the-art analytical approaches for the identification of SVs, the overall accuracy remains suboptimal, as indicated by an F1 score of 74.0%, particularly in tandem repeat regions. To enhance accuracy, it would be beneficial to explore alternative SV detection methods or develop novel approaches. Given the value of the reference panel and the fact that improved SV accuracy would lead to more precise SV imputation and GWAS results, investing effort in methodological refinement is highly encouraged.

Accurate SV calling remains an active area of research and is beyond the scope of the present study. Tandem repeat regions are particularly challenging for standardised SV detection. We believe that achieving a benchmark for NA12878 of $F1 = 74\%$ on a genome-wide level and, notably, $F1 = 91\%$ when excluding longer tandem repeats, represents strong performance. This result is especially convincing when considering that our benchmarking compared the SV calls to data generated using a different sequencing technology and processed using different bioinformatics pipelines.

(2) From the Methods section, it appears that the authors employed Beagle for both imputation and the UK Biobank imputation.

(a) It would be better to explicitly clarify this in the Results section and provide a detailed description of the corresponding procedures and parameters in the Methods section for both analyses, as this represents a key aspect of the study.

We thank the reviewer for these suggestions. Accordingly, we added the clarification to the Methods section that the leave-one-out imputation used exactly the same pipeline and settings as the UK Biobank imputation (page 14, section “Leave-one-out imputation performance”):

“We excluded one individual from the panel and imputed SVs for this individual using the panel of the remaining 887 samples, applying exactly the same pipeline and settings as those later used for SV imputation into UK Biobank (see below).”

(b) Additionally, Beagle is not specifically designed for SV imputation, the imputation quality of SVs is generally lower than that of SNVs. Exploring strategies to improve SV imputation, such as developing a novel method with reference panel data, may enhance performance.

As stated in our manuscript (page 4), we believe that, in our study, the imputation quality of SVs is lower than of that of SNVs primarily because of a) the greater difficulty of SV calling

compared to SNV genotype calling and b) the heterogeneity in SV representation across samples. Once SVs are encoded as bi-allelic markers in the reference panel, they can be imputed using the same LD/haplotype-based framework as any other variants. Accordingly, improving SV imputation is likely to benefit most from more accurate upstream SV calling and genotyping (e.g., through more robust multi-sample calling and harmonised variant representations) and not so much from improved or SV-specific imputation methods. While improved imputation is an important research direction, it is beyond the scope of the present manuscript.

(c) It is also important to assess how this reduced imputation quality may influence GWAS results. For instance, it would be useful to examine whether associated SVs exhibit higher imputation quality and whether SVs with lower quality are less likely to achieve significant association signals. In addition, the lower imputation quality observed for INV, DUP, and BND variants (Figure 3) may be due to their greater lengths (Figure 2). It is better to investigate the relationship between SV length and imputation quality.

We agree that imputation quality can influence GWAS results. For example, for the FEV1/FVC phenotype, SVs with $\text{INFO} > 0.9$ are almost twice as likely to reach genome-wide significance ($p < 5e-8$) compared with SVs with $0.7 < \text{INFO} < 0.9$ (odds ratio 1.95; Fisher's exact test p-value $2.5e-5$). This is consistent with the intuitive notion (applicable to any variants, not only to SVs) that greater uncertainty in the imputed genotypes dilutes association signals and therefore reduces power. For a detailed discussion of the relationship between allele frequency, imputation accuracy, and GWAS association results, see Zhang et al., Human Molecular Genetics 31(1):146–155 (2022), <https://doi.org/10.1093/hmg/ddab203>

We have now investigated the relationship between SV length and imputation quality (the new Supplementary Figure 6). The results suggest that the observed association between imputation quality and SV size is primarily driven by the SV-size-dependent minor allele frequency in the imputation panel.

(3) All examples presented in the manuscript focus on SVs that overlap with genes. It may also be valuable to investigate SVs that do not overlap with genes but intersect with enhancer regions. SVs can contribute to disease by altering regulatory elements, such as enhancers, which play a crucial role in gene expression. Including such analyses would further demonstrate the utility of SV-GWAS and provide deeper insights into the functional impact of SVs.

We agree with the reviewer that examining SVs intersecting with enhancer regions could be an interesting direction for future studies, as it would provide additional insights into regulatory mechanisms and disease associations. However, in the present proof-of-principle study, we prefer focusing on SVs overlapping with genes and have now highlighted this in additional detail in the revised manuscript (Discussion, page 7):

“In the present proof-of-principle study, we focused on SVs overlapping with the coding sequence of genes. In future applications of our SV imputation panel, more refined gene mapping approaches could be employed, e.g., including SVs overlapping enhancer regions or epigenetic marks. Such an enhanced mapping would increase the number of identified associated genes and thus provide additional insights into regulatory mechanisms and disease biology.”

(4) The data availability link currently provides only a VCF file ("sniffles2_joint_sv_calls.vcf.gz") containing the identified SVs.

(a) It would be beneficial for the authors to make all raw sequencing data (FASTQ files) and key processed datasets (such as alignment results and merged SV and SNV files) available. Providing these resources would enable other researchers to develop improved SV detection and imputation methods or conduct further genetic analyses.

Thank you for emphasising the importance of data sharing, which we agree with.

The Data Availability section of the manuscript already includes a link to <https://opnme.com/genomiclens>, where we made both the SV calls and the full and reduced SV imputation panel files freely available. We have now added SV summary statistics from 32 SVwide association studies to the same resource. Based on the reviewer's request, we now also reference the ENA repository project PRJEB89727 (<https://www.ebi.ac.uk/ena/browser/view/PRJEB89727>), where the raw FASTQ files are available for download, in the manuscript.

We have appended the Data Availability statement on page 22 of the revised manuscript as follows:

“Raw SV calls, the long-read sequencing-based SV imputation panel, and the SV summary statistics from 32 SV-wide association studies are available through the OpnMe initiative of Boehringer Ingelheim GmbH (<https://opnme.com/genomiclens>). The raw long-read sequencing data (FASTQ files) for the 1000 Genomes Project samples included in this study are accessible via the European Nucleotide Archive under accession number PRJEB89727 (<https://www.ebi.ac.uk/ena/browser/view/PRJEB89727>). The dataset analysed here constitutes a subset of this broader collection.”

(b) Furthermore, establishing a dedicated website for data access, along with a genome browser for SV visualization, could significantly enhance the impact and accessibility of the study. Additionally, all code, particularly the SV imputation pipeline accompanied by a detailed tutorial, should be deposited in a public repository such as GitHub. This would support researchers in imputing SVs and conducting SV-GWAS on their own datasets.

The Methods section provides a full and detailed description of the imputation pipeline and parameters in the section “Preprocessing and imputation of SVs into UK Biobank” on page 14 of the revised manuscript. Our data processing simply consisted of running standard bioinformatics tools with the parameters exactly as described in the manuscript.

Reviewer #1 (Recommendations for the authors):

(1) In the Results section, Figure 3b is mentioned before 3a, and it is better to switch them in the Figure.

Thank you for highlighting this fact. We acknowledge that typically the sequence of sections matches exactly between text and figures. However, in this specific case, we would prefer to deviate from the norm: In our opinion, Figure 3 is easier to interpret in its current sequence. At the same time, the text flows more logically in its current sequence, describing 3b before 3a. We would therefore prefer to stick to the current order, even if it means that Fig. 3b is described before 3a in the text.

(2) Page 10, "Figure 1e" -> "Figure 2e".

Thank you, we corrected this issue.

(3) Page 14, "Leave-one out" -> "Leave-one-out".

Thank you, we corrected this mistake.

(4) It is better not to use abbreviations in the subheadings, especially "UKB" (page 3).

Thank you, we changed the acronym 'UKB' to 'UK Biobank' in all subheadings.

Reviewer #2 (Public review):

Summary:

The authors aimed to develop a novel and efficient method for SV detection, utilizing data from the 1000 Genomes Project (1KGP) for modeling and calibration. This method was subsequently validated using UK population data and applied to identify structural variants associated with specific disease phenotypes.

Strengths:

Third-generation single-molecule sequencing data offers several advantages over traditional high-throughput sequencing methods, particularly due to its long-read lengths, which provide valuable insights into significant forms of genomic variation. The authors have developed an efficient method for detecting structural variations and optimizing the utilization of genomic data. We hope that this method will continue to be refined, enabling researchers to more effectively leverage long-read data, high-throughput data, or even a synergistic combination of both.

Weaknesses:

Although this research contributes to our ability to more effectively utilize long-length and high-throughput data, there are some key issues that need to be addressed in terms of analyzing the specific results as well as writing the article.

Reviewer #2 (Recommendations for the authors):

(1) How to discuss the lower detection rate of structural variations (SVs) in East Asian populations, it is worth considering whether the authors' training dataset, which may have been based on raw data with insufficient representation of East Asian individuals, could have introduced a bias favoring other populations. This potential bias might arise from the relatively limited data available for Asian ancestry. Alternatively, the observed differences could also be influenced by the role of natural selection, which may have shaped the genomic landscape of East Asian populations in distinct ways. Further investigation is needed to clarify these possibilities.

Thank you for raising this important point. Although an interesting research direction, a detailed investigation of the factors affecting SV detection rates is beyond the scope of the present study. However, we do not think that the lower detection rate in East Asians is due to an underrepresentation of Asian ancestry in our dataset. To explain this to all readers, we have added the following explanation to page 7 of the Discussion:

"In this context, we observed that the number of SVs detected per individual differed between superpopulations. We identified the highest average number of SVs in individuals of African descent and a slightly lower average in East Asians, compared to the other superpopulations. While we included a higher number of African ancestry individuals, the number of East Asian individuals included in our reference panel was comparable to the number of individuals from other, non-African ancestries. In fact, it was even larger than the number of European ancestry individuals (AFR n=241, SAS n=171; EAS n=168; EUR n=164; AMR n=144). Therefore, we do not expect a major bias from underrepresentation of any superpopulation in the training dataset. It is well established that African ancestry is more diverse than is the case for other superpopulations [32, 33] and previous studies indicate that

East Asian populations tend to exhibit slightly lower genetic diversity compared to European populations [34], which is consistent with the lower observed SV counts per genome.”

(2) *The authors did not present the results of the detection of CNV.*

Copy number variations (CNVs) are considered a subclass of structural variants. In our analysis, we detected deletions and duplications, which represent the most common forms of CNVs. However, we did not specifically investigate high copy-number SVs, as these are often larger than what can be reliably detected using long-read sequencing. Large-scale CNVs are typically identified in biobank studies through analysis of intensity data from genotyping microarrays using tools like PennCNV, and there is extensive literature supporting the use of this microarray approach in UK Biobank and other genotyped cohorts, see for example Aguirre et al.: Phenomewide Burden of Copy-Number Variation in the UK Biobank. *Am J Hum Genet.* 2019, Aug 1;105(2):373-383. doi:10.1016/j.ajhg.2019.07.001.

(3) *Multiple testing correction is essential for ensuring the validity of large-scale structural variation (SV) association analyses. It is strongly recommended that the statistical methods and correction strategies employed, such as Bonferroni correction or false discovery rate (FDR) control, be explicitly detailed to enhance the transparency and reliability of the findings.*

For genome-wide SV association analyses, we applied the commonly used genome-wide significance threshold of 5×10^{-8} , which is standard in genome-wide studies. Given that these were exploratory proof-of-principle analyses illustrating use cases for SV analyses, we decided not to correct on top of that for multiple testing for the number of traits (32) tested. For the pQTL analyses, we further adjusted this threshold using a Bonferroni-type correction based on the number of proteins tested (1,463), to account for the increased number of multiple comparisons.

We have now added a more detailed description of this multiple testing procedure to the Methods subsection “SV-wide association studies in UK Biobank” on page 16 of the revised manuscript:

“In the exploratory SV-WAS, we used the standard threshold for genome-wide significance of $p < 5 \times 10^{-8}$. For the pQTL analyses, we applied Bonferroni correction for multiple testing on top of that genome-wide threshold, correcting for the number of tested protein levels ($n=1463$): $p < 5 \times 10^{-8} / 1463 = 3.4 \times 10^{-11}$.”

(4) *The study primarily relied on data from the 1000 Genomes Project (1KGP) and the UK Biobank; however, the UK Biobank cohort is predominantly composed of individuals of European ancestry, which may restrict the generalizability of the research findings to other populations.*

Our reference panel was constructed to cover multiple ancestries, enabling imputation for diverse populations. Thus, our imputation panel can be applied to biobanks around the world and is freely available for this purpose. As a proof of principle, we have demonstrated the feasibility and performance of SV imputation in UK Biobank as an example of a broadly accessible cohort. We are looking forward to biobanks from diverse ancestries downloading our imputation panel and applying it to their populations.

(5) *Although the study employed long-read sequencing technology, the validation of structural variation (SV) detection accuracy predominantly relied on internal data, such as 'leave-one-out' validation. To further strengthen the reliability of the SV detection methods, it is recommended to incorporate additional external independent datasets for validation.*

The leave-one-out procedure in our study was used to validate the imputation performance, not the accuracy of SV detection. To assess SV calling accuracy, we performed extensive benchmarking against external SV call datasets derived from PacBio long-read sequencing and Illumina short-read sequencing. These details are provided under the subheading ‘Structural variant calling and benchmarking’ in the Results section on page 2 of the manuscript.

(6) Some of the SVs mentioned in the study overlap with disease association loci in the GWAS Catalog, but functional annotation and exploration of the biological mechanisms of these SVs are more limited. It is suggested that LD can be added to analyse whether there are SNP that are highly linked to them to further explore their functions.

We thank the reviewer for this suggestion. We have actually conducted an analysis addressing exactly this question: We performed conditional association analyses of the SV signals with nearby short variants (SNPs and InDels) at the SV locus. Such a conditional analysis addresses whether the observed SV association is influenced by LD-correlated SNPs or not. The results of this analysis are reported in Supplementary Tables 16 and 17. These tables include both the conditional analysis results and the LD between each SV and the variant at the locus with the second-highest evidence for an association.

Researchers interested in exploring the biological significance of the SV-WAS results in more detail can now download the full SV-WAS summary statistics from <https://openme.com/genomiclens> .

(7) The discussion section could be further expanded to explore the role of SV in complex diseases and its potential application in precision medicine. For example, it could discuss how SV information can be integrated into existing GWAS frameworks to enhance the accuracy of disease risk prediction.

Thank you for the suggestion, we have now added the following sentences to the discussion (page 7/8):

“Structural variants can influence complex disease biology through either the disruption of coding sequence or an altered regulation of gene expression. Such effects may not be well captured by short variants alone. Incorporating SVs into GWAS and follow-up analyses would thus provide more accurate disease risk prediction, uncover underlying pathomechanisms by highlighting actionable pathways and targets, and support precision medicine by providing biomarkers for patient stratification.”

(8) The geographic labeling of certain samples in Figure 2 appears to contain inaccuracies. For instance, the CDX sample, which represents the Dai population from Xishuangbanna in China's Yunnan Province, is currently mislabeled as originating from China's Inner Mongolia. This discrepancy should be corrected to ensure the accuracy of the data representation.

We apologise for the misunderstanding. The geographic map in Figure 2a serves as an illustrative mapping of the samples to countries. It is intended to provide readers with an overview of population coverage, rather than to indicate the precise geographic origins of individual populations. The populations CDX, CHB, and CHS are displayed within the outline of China in alphabetical order, without any intention to indicate their exact geographic origin. We changed the respective figure caption to make this clear (page 19 of the revised manuscript):

“Map of the 888 samples from the 1000 Genomes project, mapping the samples to countries and not indicating detailed geographical origins of populations.”

Reviewer #3 (Public review):

This study successfully identified genetic loci associated with various traits by generating large-scale long-read sequencing data from a diverse set of samples. This study is significant because it not only produces large-scale long-read genome sequencing data but also demonstrates its application in actual genetics research. Given its potential utility in various fields, this study is expected to make a valuable contribution to the academic community and to this journal. However, there are several critical aspects that could be improved. Below are specific comments for consideration.

Strengths:

Producing high-quality, large-scale variant datasets and imputation datasets

Weaknesses:

(1) Data availability

Currently, it appears that only the Genomic Lens SV Panel is available on the webpage described in the Data Availability section. It is unclear whether the authors intend to release the raw sequencing data. Since the study utilized samples from the 1000 Genomes Project, there should be no restriction on making the data publicly accessible. Given this, would the authors consider making the raw sequencing reads publicly available? If so, NCBI SRA or EBI ENA would be the most appropriate repositories for data deposition. I strongly encourage the authors to consider public data release. Additionally, accessing the Genomic Lens SV Panel data does not seem straightforward. The manuscript should provide a more detailed description of how researchers can access and utilize these data. In my opinion, the best approach would be to upload the variant data (VCF files) to a public database such as the European Variation Archive (EVA) hosted by EBI.

I strongly request that the authors publicly deposit the variant data. At a minimum:

(a) The joint genotype data for all 888 samples from the 1000 Genomes Project must be publicly available.

Thank you for emphasising the importance of data sharing, which we agree with.

The Data Availability section of the manuscript already includes a link to <https://opnme.com/genomiclens>, where we make both the SV calls and the full and reduced SV imputation panels (provided as multi-sample VCF files) freely available. Based on the reviewer's request, we now also reference the ENA repository project PRJEB89727 (<https://www.ebi.ac.uk/ena/browser/view/PRJEB89727>), where the raw FASTQ files are available for download.

We have appended the Data Availability statement on page 22 of the revised manuscript as follows:

“Raw SV calls, the long-read sequencing-based SV imputation panel, and the SV summary statistics from 32 SV-wide association studies are available through the OpnMe initiative of Boehringer Ingelheim GmbH (<https://opnme.com/genomiclens>). The raw long-read sequencing data (FASTQ files) for the 1000 Genomes Project samples included in this study are accessible via the European Nucleotide Archive under accession number PRJEB89727 (<https://www.ebi.ac.uk/ena/browser/view/PRJEB89727>). The dataset analysed here constitutes a subset of this broader collection.”

(b) For the UK Biobank samples, at least allele frequency data should be disclosed.

Supplementary Table 5 [↗](#) includes the allele frequencies of the SVs imputed into UK Biobank.

(c) Since eLife has a well-established data-sharing policy, compliance with these guidelines is essential for publication in this journal.

By sharing the FASTQ files, the SV calls, the SV imputation panels, the SV summary statistics, and (once processed by UK Biobank) the genotypes of SVs imputed into UK Biobank, we are providing all SV data generated in our study.

(2) Long-read sequencing data quality

While the manuscript presents N50 read length and mean or median read base quality for each sample in a table, it would be highly beneficial to visualize these data in figures as well. A violin plot or similar visualization summarizing these distributions would significantly improve data presentation.

Notably, the base quality of ONT long-read sequencing data appears lower than expected. This may be attributed to the use of pore version 9.4.1, but the unexpectedly low base quality still warrants attention. It would be helpful to include a small figure within Figure 2 to illustrate this point. A visual representation of read length distribution and base quality distribution would strengthen the manuscript.

We thank the reviewer for this suggestion. We have now included two violin plots (the new Supplementary Figure 1 [↗](#)) to the revised manuscript, summarising a) the N50 read length per sequencing run and b) the median read quality per sequencing run. These plots provide a clearer visualisation of the underlying distributions. We do not consider the ONT base quality to be low. Importantly, structural variant detection is generally robust to modest variations of per-base quality. Therefore, we do not expect the observed base quality levels to significantly affect SV calling in this study.

(3) Variant detection precision, recall, and F1 score

This study focuses on insertions and deletions (indels) {greater than or equal to}50 bp, but it remains unclear how well variants <50 bp are detected. I am particularly interested in the precision, recall, and F1 score for variants between 5-49 bp.

While ONT base quality is relatively low, single-base variants are challenging to analyze, but variants {greater than or equal to}5 bp should still be detectable as their read accuracy is still approximately 90%, making analysis feasible. Given that Sniffles supports the detection of variants as small as 1 bp, I strongly encourage the authors to conduct an additional analysis.

A simple two-category classification (e.g., 5-49 bp and {greater than or equal to}50 bp) should suffice. Additionally, a comparative analysis with HiFi and short-read sequencing data would be highly valuable. If possible, I strongly recommend that all detected variants {greater than or equal to}5 bp be made publicly available as VCF files.

Because short InDels are available from high-coverage Illumina sequencing data generated for the same individuals (i.e., the data referred to as the NYGC dataset in our manuscript), we decided against calling such short variants from our lower coverage Oxford Nanopore data and thus concentrated our efforts on reliably calling longer variants covering at least 50 bp, consistent with the conventional definition of structural variants.

(4) Assembly-based methods

Given the low read accuracy and low sequencing depth in this dataset, it is understandable that genome assembly is challenging. However, the latest high-quality

human genome datasets-such as those produced by the Human Pangenome Reference Consortium (HPRC)demonstrate that assembly-based approaches provide significant advantages, particularly for resolving complex and long structural variants.

Since HPRC data also utilize 1000 Genomes Project samples, it would be highly informative to compare the accuracy of ONT sequencing in this study with HPRC's assembly-based genome data. The recent publication on 47 HPRC samples provides a valuable reference for such a comparison. Given its relevance, the authors should consider providing a comparative analysis with HPRC data.

The aim of the present study was to generate an SV reference panel that enables SV imputation for large biobanks. Detailed assessments of ONT sequencing quality in general and comparisons to other sequencing efforts and technologies are out of scope for the present manuscript. We invite the scientific community to use the FASTQ files provided at ENA for conducting such detailed assessments in follow-up studies.

<https://doi.org/10.7554/eLife.106115.2.sa0>