

Benchmarking and Optimization of Methods for the Detection of Identity-By-Descent in High-Recombining *Plasmodium falciparum* Genomes

Reviewed Preprint

v2 • July 7, 2025

Revised by authors

Reviewed Preprint

v1 • January 6, 2025

Bing Guo, Shannon Takala-Harrison , Timothy D O'Connor 

Center for Vaccine Development and Global Health, University of Maryland School of Medicine, Baltimore, United States • Institute for Genome Sciences, University of Maryland School of Medicine, Baltimore, United States

 https://en.wikipedia.org/wiki/Open_access

 Copyright information

eLife Assessment

This **important** study presents an evaluation of several tools used for detecting Identity-By-Descent (IBD) segments in highly recombining genomes, using simulated data to replicate the high recombination and low marker density of *Plasmodium falciparum*, the parasite responsible for malaria. The evidence presented by the authors is **convincing** demonstrating that users should be cautious calling IBD when SNP density is low and recombination rate is high. This study will be of interest to scientists working in the field of genome evolution and infectious diseases

<https://doi.org/10.7554/eLife.101924.2.sa3>

Abstract

Genomic surveillance is crucial for identifying at-risk populations for targeted malaria control and elimination. Identity-by-descent (IBD) is increasingly being used in *Plasmodium* population genomics to estimate genetic relatedness, effective population size (N_e), population structure, and signals of positive selection. Despite its potential, a thorough evaluation of IBD segment detection tools for species with high recombination rates, such as *P. falciparum*, remains absent. Here, we perform comprehensive benchmarking of IBD callers – probabilistic (hmmIBD, isoRelate), identity-by-state-based (hap-IBD, phased IBD) and others (Refined IBD) – using population genetic simulations tailored for high recombination, and IBD quality metrics at both the IBD segment level and the IBD-based downstream inference level. Our results demonstrate that low marker density per genetic unit, related to high recombination relative to mutation, significantly compromises the accuracy of detected IBD segments. In genomes with high recombination rates resembling *P. falciparum*, most IBD callers exhibit high false negative rates for shorter IBD segments, which can be partially mitigated through optimization of IBD caller parameters, especially those related to marker density. Notably, IBD detected with optimized parameters allows for more accurate capture of selection signals and population structure; IBD-based N_e inference is very sensitive to IBD detection errors, with IBD called from hmmIBD uniquely providing less biased estimates of

N_e in this context. Validation with empirical data from the MalariaGEN *Pf*7 database, representing different transmission settings, corroborates these findings. We conclude that context-specific evaluation and parameter optimization are essential for accurate IBD detection in high-recombining species and recommend hmmIBD for *Plasmodium* species, especially for quality-sensitive analyses, such as estimation of N_e . Our optimization and high-level benchmarking methods not only improve IBD segment detection in high-recombining genomes but also enhance overall genomic analysis, paving the way for more accurate genomic surveillance and targeted intervention strategies for malaria.

Introduction

Malaria is a mosquito-borne disease caused by *Plasmodium* parasites. It poses a significant public health challenge globally, with an estimated 263 million clinical cases and 597,000 deaths occurring in 2023 [1]. Despite intensive efforts toward malaria control and elimination, malaria reduction has slowed or plateaued in recent years, due to multiple factors including antimalarial drug resistance and lack of highly efficacious vaccines and detailed, timely surveillance. Advances in sequencing technologies and the scale of resequencing studies now allow for parasite genomic surveillance, which can provide insights into the efficacy of malaria interventions and guide the design of targeted elimination strategies in different transmission settings [2–4].

Identity-by-descent (IBD) is an essential tool in population genomics that has been used to estimate genetic relatedness [5–8], positive selection [6, 9–11], effective population size (N_e) [10, 12], fine-scale population structure [4, 10, 13], and migration patterns [4, 14]. IBD-based analyses of parasite genomic data have applications in diverse epidemiological settings, enhancing malaria surveillance and control by providing crucial information to researchers and policymakers [3, 15, 16]. In high transmission settings, a rapid decrease in genetic diversity and effective population size may indicate a successful malaria intervention [12]. In intermediate transmission settings, IBD-based analysis of parasite population structure and migration provides valuable insights into sources and sinks of transmission for planning targeted elimination strategies [4, 6, 10]. In low transmission settings, IBD-based estimates of pedigree relationship analysis between infections can help differentiate local transmission from importation and causes of recurrent infection [17, 18]. Additionally, IBD-based detection of positive selection can assist in identifying and monitoring the emergence and spread of antimalarial drug resistance [6, 10, 19].

However, the reliability of the IBD-based analysis is highly dependent on the accuracy of the detected IBD segments. Insufficient density of genetic markers, on a local or genome-wide scale, probably contributes to high error rates in the identification of IBD segments [11, 20, 21] and reduced accuracy of IBD-based estimates of population demography [5, 22]. Many IBD detection methods have been designed for human genomes, where the demographic history and evolutionary parameters, including the recombination rate, differ considerably from *Plasmodium falciparum* (*Pf*). The effective population size of humans has increased rapidly in recent history [23], while that of *Pf* is decreasing, particularly in regions such as Southeast Asia and South America [1, 24], due to enhanced malaria elimination efforts. More importantly, *Pf* genomes recombine about 70 times more frequently per unit of physical distance [19, 25–28] than the human genome [29], while sharing a similar mutation rate [2, 15, 30–34] as human genomes [35] on the order of 10^{-8} per base pair per generation. The decreasing population size [10] and the high recombination rate in *Pf* result in a reduced number of common variants, such as single nucleotide polymorphisms (SNPs), per unit of genetic distance. Large human whole-genome sequencing data sets typically provide millions of common biallelic SNP variants [36], while *Pf* data sets only have tens of thousands [37]. Given that the human genome is about twice as large as *Pf* in genetic units, the per-centimorgan (cM) SNP density in *Pf* can be two orders of magnitude lower than in humans, which may not provide sufficient

information for detecting IBD segments. Thus, it is critical to understand whether IBD detection methods can still generate accurate IBD segments under low SNP density conditions, considering the specific evolutionary parameters of the *Pf* genome.

Evaluating the quality of the detected IBD segments requires benchmarking with the known ground truth through simulation studies [20, 21, 38, 39]. The performance of IBD detection tools developed for use in the human context is typically measured using simulated genomes reflecting demographic and evolutionary parameters of human genomes [20, 21, 38–40], which likely do not apply directly to *Pf*. For tools explicitly designed for malaria parasites, such as isoRelate and hmmIBD, the evaluation of the quality of IBD was based on parent-offspring [8] or pedigree-based simulations (up to 25 generations) [6] that focused primarily on close relatives, which more likely mirrors low malaria transmission settings than high transmission settings. Furthermore, benchmarking methods and definitions of IBD accuracy are inconsistent across studies [6, 8, 20, 21], making the results of the quality evaluation of IBD difficult to compare. Considering the limitations of existing evaluations of IBD detection methods for *Pf* genomes, a unified benchmarking framework specifically designed for high recombining *Pf* genomes from low and high-transmission settings is needed. Such a framework will assist researchers in comparing and prioritizing different IBD detection methods for intended downstream analysis.

In the present study, we developed a unified IBD benchmarking framework that reflects the demographic and evolutionary parameters of *Pf* (Fig 1 and S1 Data). We evaluated how different recombination rates and marker densities affect the quality of detected IBD segments, performed IBD caller-specific parameter optimization, and benchmarked different IBD detection methods with their optimized parameters at both the IBD segment and downstream inference levels. Furthermore, we validated our findings from simulation analysis with empirical data sets constructed from subsets of samples from the publicly available whole-genome sequencing database MalariaGEN *Pf* 7. Our findings indicate that a high recombination rate (given the same mutation rate) is associated with a low SNP density (per genetic unit), which substantially affects the accuracy of the detected IBD segments. To obtain optimal results, we generally recommend using hmmIBD when genetic data from phased haploid genomes are available. If human-oriented IBD callers are used, we recommend optimizing the parameters prior to applying to *Pf* genomes.

Results

Low SNP density due to high recombination rate affects the accuracy of IBD calls

Accurate estimation of IBD segments often requires dense genetic markers to capture ancestral relationships left by recent recombination events [20, 38, 47–50]. However, high-recombining species like *Pf* have very low marker densities per genetic unit, which may significantly affect the accuracy of inferred IBD segments.

To assess how recombination rates affect marker density per genetic unit and the detection of IBD segments, we simulated genomes with varying recombination rates but a fixed mutation rate. Under a single-population demographic model (see Methods), we found that the density of common biallelic SNPs (minor allele frequency ≥ 0.01) per cM, in selectively neutral scenarios, is inversely correlated with recombination rates ranging from 3×10^{-9} to 10^{-6} per base pair per generation (Fig 2a). For instance, the SNP density of *Pf*-like genomes is 25 SNPs per cM, which is approximately 1/67 of that of the human-like genomes (1,660 SNPs per cM). We further assessed how low marker densities associated with high recombination rates affect the accuracy of detected IBD segments, calculating two metrics (via *ishare/ibdutils*, see Code availability), including the false negative rate (FN), which represents the proportion of a true IBD segment (obtained via tskibid

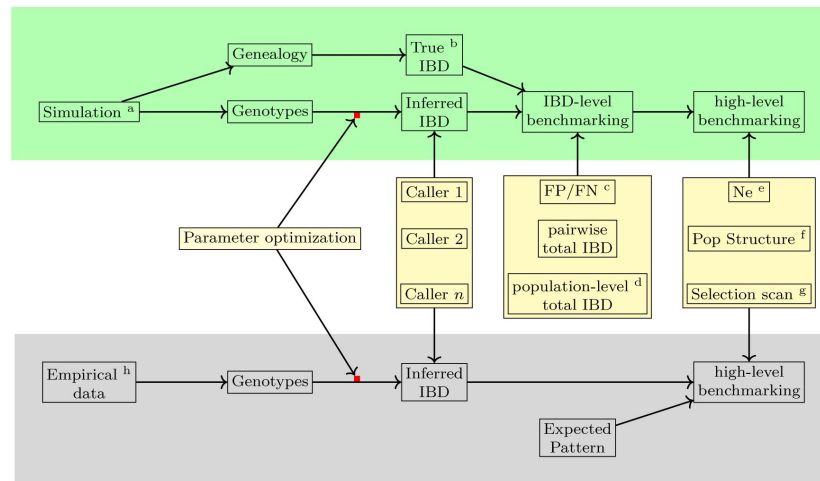


Fig 1.

Overview of methods used in benchmarking IBD detection methods.

Benchmarking and optimization of IBD callers for *Pf* include simulation analyses (top, green shading) and empirical data-based validation analyses (bottom, gray shading). (1) For the simulation study, the genealogical trees and phased genotype data are generated via the combination of forward (SLiM [41, 42]) and coalescent (msprime [43]) simulations (indicated by the superscript a in the diagram). True IBD is obtained from a simulated genealogy tree via tsKibd [10] (b) and inferred IBD from the phased genotype using different IBD callers, including hap-IBD [20], hmmIBD [8], isoRelate [6], Refined IBD [44], and phased IBD [21]. IBD benchmarking is performed at two levels. The first is at the IBD segment level. The metrics include false positive and false negative rates (c), population-level total IBD per length bin [22] (d), and total IBD per isolate pair. The second is at the level of IBD-based downstream analyses, including the effective population size (N_e) by IBDNe [22] (e), community membership through the InfoMap algorithm [45, 46] (f), and selection signals by statistics X_{iHS} [10] (g). As the default parameters of different IBD callers, particularly those developed for human data, may not be ideal for *Pf* genomes, we performed grid searches for key parameters for each IBD caller so that the comparison was based on the best performance of each caller (see S1 Data for detailed information on simulation and IBD calling parameters and used values). (2) For validation in empirical data where true IBD is not available, we obtained IBD-based estimates using IBD from different callers and assessed which version of IBD can generate the expected patterns. The empirical data sets are subsampled from the MalariaGEN *Pf* 7 database [37] (h).

[10]) not covered by inferred IBD segments of the same genome pairs, and the false positive rate (FP), which indicates the fraction of an inferred segment not covered by true segments of the same genome pairs (see

Methods for detailed definitions, and Fig 1 method overview). Our analysis showed that as the recombination rate increases, both the genome-wide FN and FP (Fig 2b) increase for IBD inferred from hmmIBD. The patterns vary in the other four IBD detection methods, and all suffer elevated FNs and/or FPs as the recombination rate increases (Fig 2b and S1 Fig), with the exception of isoRelate, which has better IBD quality with lower marker densities. The results suggest that low SNP density per genetic unit can dramatically affect the reliability of detected IBD segments.

Varying quality of IBD inferred from simulated *Pf* genomes via different IBD callers

Multiple IBD callers have been used for *Pf*, including *Pf*-oriented, Hidden Markov Model-based methods, such as hmmIBD [8] and isoRelate [6], and those originally designed for human genomes, such as Refined IBD [12] and Beagle (version 4.1) [4]. We analyzed hap-IBD [20] and phased IBD [21] in addition to hmmIBD, isoRelate, and Refined IBD, since the former represents two recent key advancements in the development of IBD detection methods that scale well to large sample sizes and genome sizes.

To evaluate the applicability and accuracy of these IBD detection methods in analyzing *Pf* genomes, we performed benchmarking analyses in simulated genomes (Fig 1, top panel), mimicking the high recombination rate and the decreasing population size of *Pf* populations. Our analyses include three sets of comparisons: (1) baseline benchmarking, where we mainly used the default parameter values for each IBD caller and compared the performance in *Pf* genomes at the level of an IBD segment and their simple statistics; (2) post-optimization benchmarking, where we used parameter values optimized specifically for each IBD caller so that the comparisons are based on the optimal performance of each method; (3) human-like genome benchmarking, where detected IBD segments are expected to have low error rates for human-oriented IBD callers and thus were used as an internal control to validate our benchmarking pipeline (see S1 Table for the IBD caller parameters and Methods for demographic models).

Baseline benchmarking analysis shows that all callers except hmmIBD suffer from high FN rates, especially for short IBD segments (genome-wide FN/FP rates reflect shorter IBD segments as most segments are short) (Fig 3). Similarly, genetic relatedness metrics based on pairwise total IBD are largely underestimated for most callers (S2 Fig). We found that by default hmmIBD has relatively low FN/FP error rates (Fig 3) and is less biased for relatedness estimates (S2 Fig). In contrast, isoRelate and human-oriented callers have high FN rates and varying FP rates (Fig 3). Thus, both *Pf*- and human-oriented IBD callers can suffer high error rates when detecting IBD from genomes with a high recombination rate and a low marker density, with the extent depending on underlying assumptions and methodologies.

IBD-caller-specific parameter optimization for *Pf* improves IBD accuracy

Since these IBD callers are optimized for different species or genotype data sets, we hypothesized that optimization of IBD caller parameters under a unified framework designed for *Pf* genomes can improve the performance of these callers in analyzing malaria parasite data. As searching the entire IBD caller parameter space is inefficient, our optimization focused mainly on parameters potentially affected by or needing adjustment due to differences in marker density between the high-recombining species (e.g., *Pf*) and the lower-recombining species (e.g., humans). For IBD

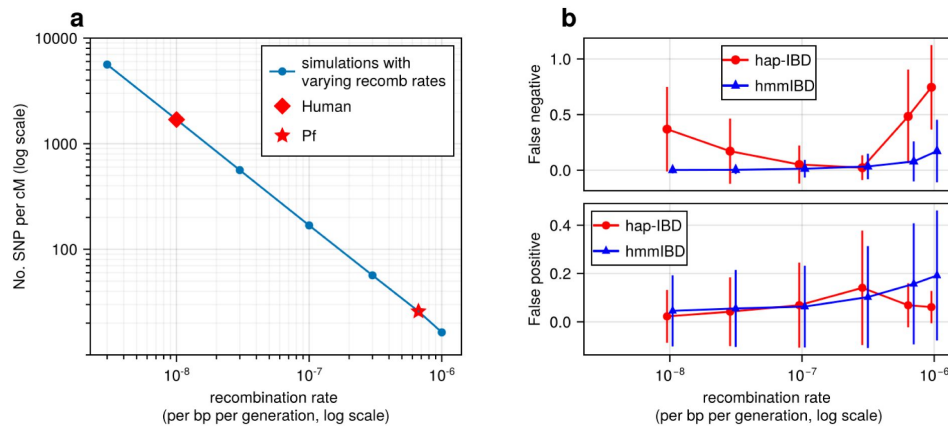


Fig 2.

High recombination rates reduce genetic marker density and affect the quality of detected IBD segments.

a, The number of common single nucleotide polymorphisms (SNPs) (minor allele frequency ≥ 0.01) per genetic unit (centimorgan, cM) in simulated genomes with different recombination rates. In these simulations (blue line), the mutation rates are fixed; the recombination rates vary widely to include the rate for both humans (red diamond) and *Pf* (red star). For each recombination rate, $n = 4$ independent simulations of chromosomes were performed. Data were plotted in the form of mean (marker) $\pm$ standard deviation (vertical lines, which are difficult to visualize given low variation among chromosomes).

b, Accuracy of IBD segments detected from genomes simulated with different recombination rates. The accuracy of IBD segments is measured by the false negative rates (top panel) and false positive rates (bottom panel). The plotted error rates represent the genome-wide rates (as defined in Methods) of IBD segments identified with default IBD caller parameters unless stated otherwise in S1 Table. These rates are based on one representative set from $n = 3$ simulation sets. The plotted vertical lines indicate standard deviations of error rates across all genome pairs in the representative simulated set. Both the vertical lines and markers are horizontally staggered for clarity. Only error rates for two IBD detection methods, hmmIBD and hap-IBD, are included in (b) for simplicity. The error rates for all 5 IBD callers are provided in S1 Fig. For both (a) and (b), the genomes were simulated under the single-population model (see Methods). Note that log scales are used for the y axis in (a) and the x axis in (b).

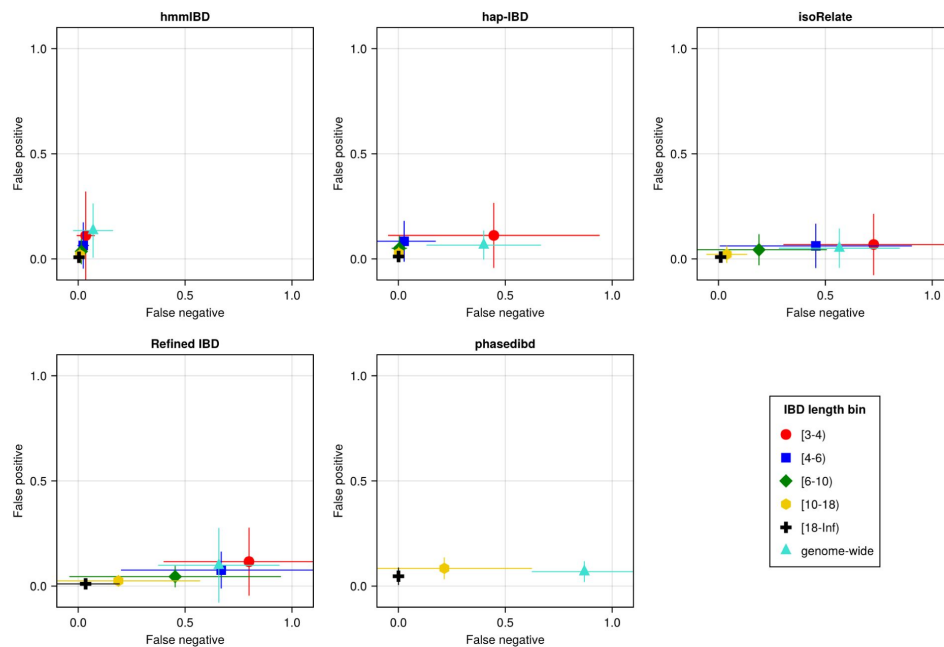


Fig 3.

Accuracy of IBD segments detected from *Pf* genomes varies across IBD callers.

IBD segments were inferred from genomes simulated under the single-population model with a shrinking population size and a recombination rate compatible with *Pf*. The accuracy of IBD was evaluated using the calculated false positive rate (y axis) and false negative rate (x axis). The rates were calculated for different length bins in centimorgans, including [3-4), [4-6), [6-10), [10-18), [18, inf) centimorgans and at the genome-wide level (defined by overlapping analysis between true IBD segments and inferred IBD segments from each genome pair). These rates are based on one representative set from $n = 3$ simulation sets. The plotted vertical and horizontal lines represent the standard deviations of error rates (horizontal for false negatives and vertical for false positives) calculated across all relevant segments for length-bin specific rates or across all genome pairs for genome-wide rates in the representative simulated set. The titles of the subplots indicate the IBD callers analyzed. The results of the simulations under the multiple-population model and the UK human demographic model are provided as S2 Data.

callers that do not explicitly have marker density-related parameters, such as hmmIBD, we explored other parameters that likely affect IBD quality. We performed grid searches to find parameter values that generate inferred IBD with low and balanced error rates (see S2 Table for parameters explored and their corresponding values, and S2 Data for detailed results).

We found that most IBD callers have a key parameter that can dramatically affect their FN/FP rates (S2 Table). For example, the FN rates of IBD called from hap-IBD change substantially when the min-marker parameter varies (see S2 Data). With a value of 70, the FN rate for short IBD segments dramatically decreases such that the FN and FP error rates become more balanced (**Fig 4a** [↗](#)). Consistently, genetic relatedness estimates change from being highly underestimated before parameter optimization (**Fig 4b** [↗](#), left column) to being more balanced after optimization (**Fig 4b** [↗](#), right column). Similar improvements were observed for Refined IBD and phased IBD (S3 Fig). In contrast, the quality metrics remained largely unchanged during parameter optimization attempts for hmmIBD and isoRelate, with hmmIBD being more accurate and unbiased and isoRelate suffering from high false negative rates and underestimated relatedness (S3 Fig).

The human-oriented IBD callers, when not optimized for *Pf*, underperformed hmmIBD, especially for Refined IBD and phased IBD (S3a Fig). To exclude potential problems in our benchmarking pipeline, we simulated genomes with recombination rates and demographic history consistent with the human population in the UK (S4a Fig, left column, and S4b Fig, left column; also see Methods). These callers, evaluated *without* *Pf*-optimized parameters, indeed perform much better on human genomes, showing consistently lower FN/FP error rates (S4a Fig, left column) and less biased total IBD-based relatedness estimates (S4b Fig, left column), compared to *Pf*-like genomes (S3a Fig and S3b Fig, left columns). The results support the robustness of our benchmarking pipeline (S4a Fig and S4b Fig, left columns) and demonstrate challenges in applying human-oriented IBD callers to *Pf*. Furthermore, we found that IBD caller performance varies with demographic configurations (the single-population model in S3a Fig and S3b Fig, left columns, versus UK human model in S4a Fig and S4b Fig, right columns) even with the *same* (*Pf*) recombination/mutation rates and default IBD caller parameters, suggesting that the optimization is demography-dependent (see details in S2 Data).

Post-optimization benchmarking via downstream inferences

IBD-based downstream analyses, such as estimation of N_e , selection signals, and population structure, are key applications of IBD segments in population genetics, which often rely on the high quality of input IBD segments. With optimized IBD caller-specific parameters tailored for *Pf*-like genomes, we can expect IBD callers to perform at their best, which allows high-level benchmarking by comparing IBD-based downstream estimates.

For IBD-based selection detection, we simulated positive selection on each of the 14 chromosomes via the single-population model, and identified IBD peaks with different callers (see Methods). These peaks, inferred as regions under selection, were considered true signals if they contained the selected site from simulations. We found that most callers can capture the majority of the simulated signals except Refined IBD, which is less sensitive and only detects 2-3 out of 14 selected loci (S5a Fig); isoRelate shows an increased level of false positives or low signal-to-noise ratios, evident in IBD coverage curves (S5a Fig).

For IBD-based population structure inference, we simulated *Pf*-like genomes under a selectively neutral condition using the multiple-population demographic model, and performed IBD network community detection via InfoMap [[45](#) [↗](#), [46](#) [↗](#)] to define subpopulations (see Methods). Similarly to IBD-based selection signal detection, we found that IBD inferred from most callers can accurately recapitulate the simulated population structure, comparable to true IBD (S5b Fig). The exception is that isoRelate tends to generate many smaller groups, which is likely due to high FN

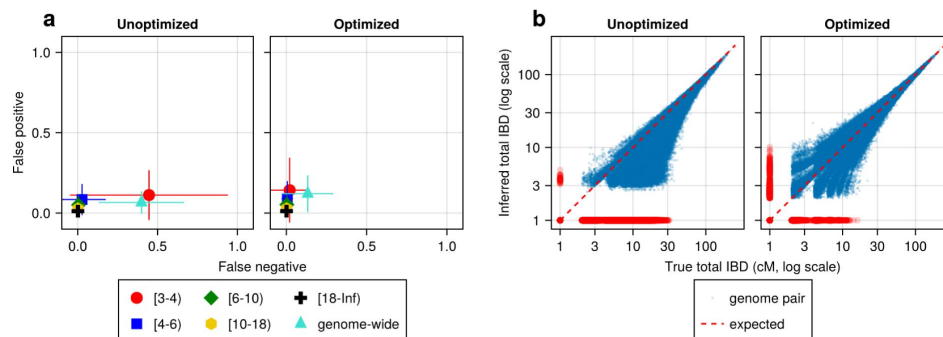


Fig 4.

IBD caller-specific parameter optimization can improve the quality of IBD segments inferred from simulated *Pf* genomes (using hap-IBD as an example).

a, Quality of detected IBD measured by false positive and false negative rates before (left column) and after (right column) hap-IBD-specific parameter optimization. As indicated in the axis legend, the error rates were calculated for different length ranges (in centimorgans), including [3-4], [4-6], [6-10], [10-18], [18, inf) and at the genome-wide level. These rates are based on one representative set from $n = 3$ simulation sets. The plotted vertical and horizontal lines represent the standard deviations of error rates (horizontal for false negatives and vertical for false positives) calculated across all relevant segments for length-bin specific rates or across all genome pairs for genome-wide rates in the representative simulated set. b, Quality of detected IBD measured by total genome pairwise IBD, an estimate of genetic relatedness, before (left column) and after (right column) hap-IBD parameter optimization. Data of $n = 1,000$ haploid genomes from a representative simulation set was plotted with each subplot for hap-IBD before and after parameter optimization as indicated in the titles. Each dot represents a pair of genomes with the coordinates x and y being true and inferred total IBD. In (b), the blue dots are the pairs with nonzero true and inferred total IBD while red dots are pairs with either true total IBD or inferred total IBD being 0; zero-valued total IBD was replaced with 1.0 cM for visualization purposes. The red dotted line of $y = x$ indicates the expected pattern, that is, true total IBD equal to inferred total IBD if the inferred IBD was 100% accurate. Note that log scales are used in both the x and y axes in (b).

rates for short IBD segments, missing connectivity among distantly related genomes, and only showing closely related subgroups of small sizes. As indicated by the low adjusted Rand indices, there was little agreement between the community labels inferred by isoRelate and the true labels.

For IBD-based N_e inference, we simulated neutral *Pf* genomes using the single-population model and estimated N_e from detected IBD via IBDNe [22]. We found most of the compared IBD callers suffer from wild oscillations, which have previously been observed [22], and deviate significantly from the truth for older generations (Fig 5), consistent with the general pattern of high error rates in shorter IBD length bins for these IBD callers (S3a Fig, right column). Meanwhile, the IBD inferred by hmmIBD generated highly accurate estimates comparable to true IBD (Fig 5). We explored the mechanisms underlying bias in N_e estimates and found that the strong bias is likely due to the underestimation of population-level total IBD for short IBD segments, which is most obvious for hap-IBD, isoRelate, and phased IBD, followed by Refined IBD (S6 Fig). For hmmIBD, both the estimates of N_e (Fig 5) and population total IBD are relatively unbiased (S6 Fig, column 2), consistent with its relatively low and balanced FP/FN rates (Fig 3, leftmost panel). These results suggest that IBD-based N_e estimates are highly sensitive to the quality of input IBD segments, and that hmmIBD is more accurate for this analysis.

Given that N_e estimation in *Pf* is very sensitive to the quality of detected IBD, we explored whether excluding short, error-prone IBD segments (< 4 cM) could improve N_e estimates for callers other than hmmIBD. The exclusion results in reduced oscillation of the trajectory for some callers, like hap-IBD, but wide confidence intervals or underestimation in older generations, in both simulated and empirical data (S7 Fig). We then explored reasons underlying the recent oscillation or drop (around 20 generations ago) commonly observed in the estimated N_e trajectories (S7a Fig, second row and S7b Fig, both rows) [12, 40, 51]. We hypothesized that this oscillation is partially due to IBD segments with TMRCA < 1.5 generations ago being included in the IBDNe input [22]. We found that removing these segments can greatly mitigate this problem (S7a Fig), especially for hmmIBD and Refined IBD. The findings suggest caution when interpreting (1) a recent drop in an estimated N_e trajectory in empirical data sets where TMRCA-based filtering is less practical, and (2) extremely large estimates in older generations stemming from high error rates for short IBD segments.

To confirm that parameter optimization improves downstream inferences, we compared postoptimization results (S5 Fig and Fig 5) with pre-optimization estimates (S8 Fig and S9 Fig). We found that parameter optimization increased the accuracy of selection detection (in isoRelate, Refined IBD, and phased IBD), improved population structure inference (phased IBD and hap-IBD), and reduced oscillation on the N_e trajectory (Refined IBD).

Validation in empirical data sets

We further validated the findings from simulation analysis in empirical data sets, by constructing “single” or “multiple” population data sets based on the MalariaGEN *Pf* 7 data [37] (see Methods for details). As true IBD segments are not available here, we focused on high-level benchmarking by evaluating whether IBD-based downstream estimates are consistent with expected patterns, including N_e estimation and selection signal detection with “single” population data sets and InfoMap population structure inference with the “multiple” population data set (S3-S5 Tables).

With optimized parameters, all callers, except Refined IBD, capture most known selection signals from the Southeast Asia data set. These signals include selective sweeps associated with antimalarial drug resistance and sexual commitment, such as dihydrofolate reductase (*dhfr*) [52], multidrug resistance protein 1 (*pfdmr1*) [53], amino acid transporter 1 (*pfaat1*) [19], chloroquine resistance transporter (*pfcr1*) [54], dihydropteroate synthase (*dhps*) [55], Apicomplexan-specific ApicAP2-g (*ap2-g*) [56] and apicoplast ribosomal protein S10 (*arps10*) [57] (Fig 6a). hmmIBD detects more peaks but suffers from noise, likely due to the relatively high FP to FN ratios for short (<4 cM) IBD segments (Fig 2 and Fig 3).

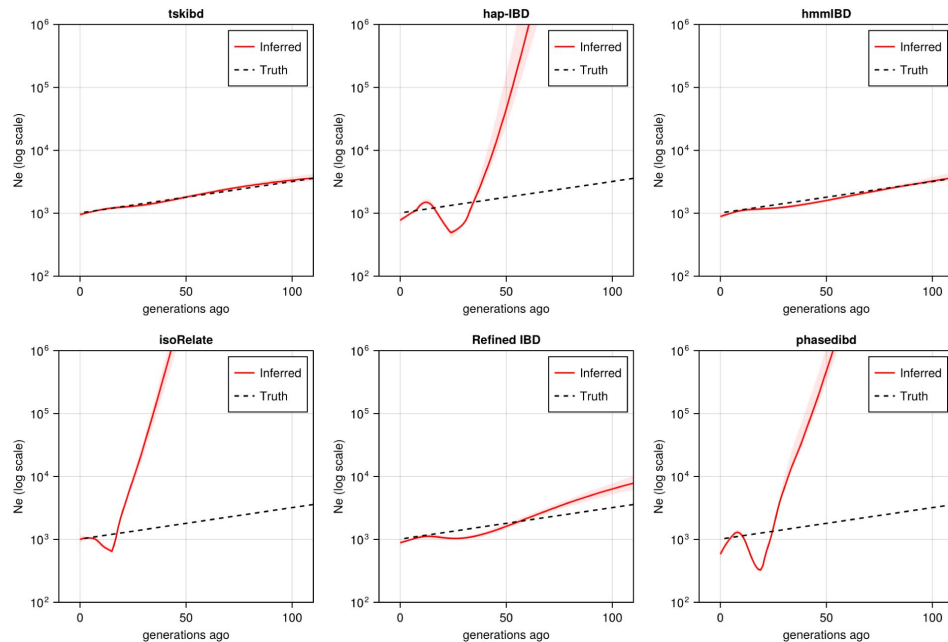


Fig 5.

Post-optimization benchmarking of different IBD callers by comparing downstream estimates

N_e . With parameters optimized for each IBD caller, the performance of IBD callers was evaluated by comparing the N_e trajectory for the recent 100 generations estimated via IBDNe based on true (black dashed line) IBD versus inferred IBD (red solid line). True IBD was calculated from simulated genealogical trees via tskibd; inferred IBD includes those inferred from hap-IBD, hmmIBD, isoRelate, Refined IBD, and phased IBD, with their N_e estimates shown from top left to bottom right. The shaded areas surrounding the red lines indicate 95% confidence intervals as determined by IBDNe. The plots show results from one representative set of $n = 3$ replicated simulation sets. See S9 Fig for pre-optimization results. Note that log scales are used on the y axes.

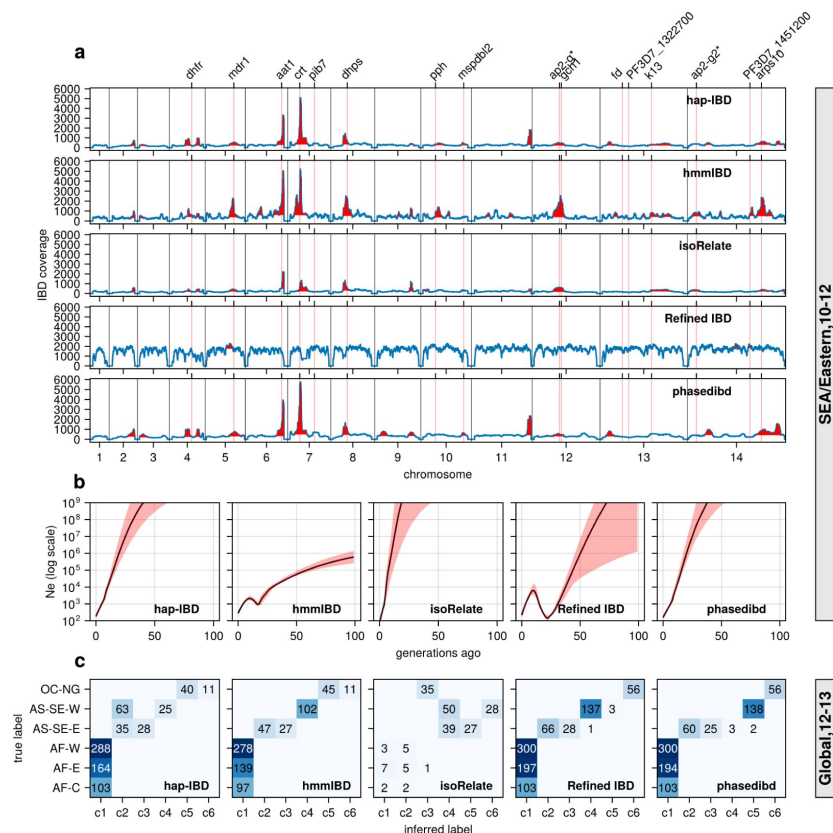


Fig 6.

Validation of the performance of IBD callers in empirical data sets by comparing IBD-based downstream analyses.

a, IBD coverage and detected selection signals in the SEA data set using different IBD callers (rows 1 to 5). Annotations and corresponding vertical dotted lines at the top indicate the center of known and putative drug resistance genes and genes related to sexual commitment; red shading indicates regions that are inferred to be under positive selection (see Methods for definitions). b, N_e estimates of the SEA data set based on IBD inferred from different callers. Line plots are point estimates; the shaded areas around the line plots indicate confidence intervals based on bootstrapping (generated by IBDNe). c, Inference of the population structure of the structured data set by the InfoMap community detection algorithm using the IBD inferred from different IBD callers. The rows of the heatmap are geographic regions of isolates, and the columns are the largest, inferred communities, labeled as c1 to c6. The heatmap color represents the number of isolates in each block with the given row and column labels. The columns are rearranged so that the diagonal blocks tend to have the largest values per row for better visualization. Note that log scales are used in y axes in (b).

Similar to the simulation analysis, IBD detected using most callers resulted in N_e estimates with unrealistic oscillations for the Southeast Asia data set, including extremely large estimates in the more distant past (> 20 generations ago) (**Fig 6b**). The problems are much less severe with the estimates from hmmIBD, which mirrored the expected reduction in malaria in this region due to the intense efforts to eliminate malaria in recent decades [1].

InfoMap-based community detection reveals expected continental population structure across most callers: African *Pf* parasites are less structured, and Southeast Asian parasites are more structured and distinct from Oceanian parasites (**Fig 6c**), consistent with previous non-IBD-based methods [37]. isoRelate IBD estimates generate many small, close groups, likely due to high false negative rates for short IBD segments, especially in high transmission settings like Africa where parasites have low relatedness and mainly share short IBD segments (**Fig 6c**).

To further confirm the improvement of IBD-based estimates through parameter optimization, we performed analyses with IBD detected with unoptimized parameters (see S1 Table). We found that the height and number of IBD peaks for hap-IBD and phased IBD decreased significantly (S10a Fig versus **Fig 6a**). Pre-optimization N_e estimates show more extreme oscillations, especially for human-oriented IBD callers (S10b Fig). Pre-optimization IBD estimates from hap-IBD and phased IBD fail to reveal the expected population structure, particularly in African parasite populations (S10c Fig). These differences underscore the importance of parameter optimization for *Pf*, especially for IBD callers not validated for *Pf*.

Computational efficiency comparison

With the decrease in whole-genome sequencing cost and the increase in sample availability, it is important to prioritize IBD callers that scale well for large sample sizes, such as MalariaGEN *Pf*7 ($n = 20,864$) [37]. We compared the IBD inference time and maximum memory usage for different IBD callers with or without parallelization. When using a single thread, probabilistic inference algorithms like isoRelate, Refined IBD, and hmmIBD are about two orders of magnitude slower than those based on identity-by-state-based or positional Burrows-Wheeler transform (PBWT) based algorithms, such as hap-IBD and phased IBD (**Fig 7a**). Maximum memory consumption is highest in Refined IBD, with hmmIBD being around ten times more efficient (S11a Fig). With multithreading, the patterns are similar to single-thread comparison as most allow parallelization (**Fig 7b** and S11b Fig). The exception is hmmIBD, as it currently only supports single-thread computation. Despite the computational efficiency of hmmIBD compared to isoRelate and its high accuracy in detecting IBD segments from *Pf* genomes, it remains significantly slower than IBS-based methods, highlighting the need for further enhancements for large data sets like MalariaGEN *Pf*7.

Discussion

In this work, we evaluated the reliability of IBD detection methods in high-recombining genomes of malaria parasites. Our findings indicate that low marker density per genetic distance significantly affects IBD detection accuracy. Optimizing parameters of IBD detection methods for *Plasmodium* genomes enhances the accuracy of detected IBD segments, thereby improving subsequent downstream analyses such as the inference of positive selection signals and population structures. However, variations in performance among IBD detection methods remain substantial, particularly in analyses sensitive to IBD quality, such as the estimation of effective population size (N_e), even after parameter optimization. We generally recommend hmmIBD for IBD estimation in *Plasmodium* genomes when phased haploid genomes are available. We further emphasize the necessity of performance evaluation and parameter optimization for IBD callers prior to application in untested species or scenarios.

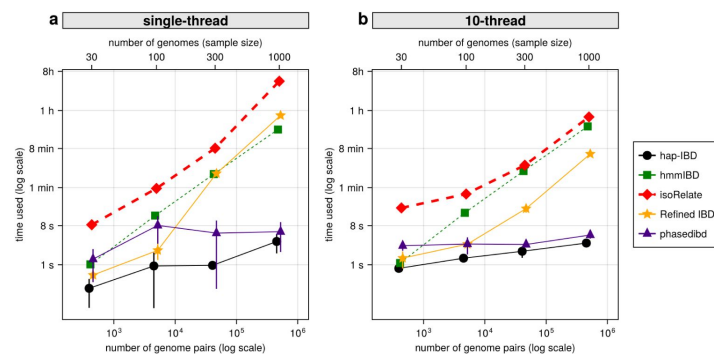


Fig 7.

Comparison of computational runtime for IBD calling process for different callers.

a, Runtime for different IBD callers to detect IBD from genomes of different sample sizes in single-thread mode. The comparison is based on *Pf* genomes of size of 100 cM simulated under the single-population model. The x-axis tick labels include the number of pairs of genomes simulated and analyzed (below the plot, reflecting the number of computation units) and the number of haploid genomes (above the plot, representing sample size) analyzed. The line styles and markers for different callers/tools are provided in the legend box on the far right of the figure, which is shared across the two subplots. Values on the y axis represent means (markers) and standard deviations (vertical error lines) from $n = 3$ sets of independent simulations. Note that at each sample size, the x values for different IBD callers are slightly staggered to prevent the error lines from overlapping; some of the error lines are hard to visualize as they are relatively small. b, Runtime in multithreading mode. (b) is organized similarly to (a), except that the IBD calling processes were run in multithreading mode with 10 CPU threads. Note that log scales are used in y axes and bottom x axes. Also, see S11 Fig for the maximum memory usage for different callers.

Comparing the performance of multiple IBD detection methods requires a unified framework, which should include a uniform definition of accuracy and a simulated ground truth that mimics *Pf* genomes. Our benchmarking framework incorporates several novel features. First, it utilizes a consistent definition of IBD length-specific accuracy based on the overlap of IBD segment lengths, closely aligned with metrics used in hap-IBD [20], phased IBD [21], and Refined IBD [44], as detailed in our Methods section. The approach differs from the original evaluation of hmmIBD, which assesses the accuracy of IBD based on the fraction of SNPs that share IBD, which could overlook the precise accuracy of the length of the detected IBD segments [8]. In particular, the original study of isoRelate defined the accuracy (true positive rate) using a less stringent overlap-by-count criterion where a segment is counted as accurate when at least 50% of a detected IBD segment is overlapped by true segments [6]. Second, we generated *Pf*-like genomes via population genetic simulation that reflect a realistic distribution of IBD segment lengths. This contrasts with previous studies, where methods such as hap-IBD and Refined IBD used human-like data from population genetic simulations [20, 44], whereas hmmIBD, isoRelate, and phased IBD relied on simulations based on artificial recombination or pedigree models [6, 8]. These models (non-population-based genetic simulation) often produce long shared IBD segments typical of close relatives, failing to capture the IBD length distribution in population samples predominantly comprising distant relatives. Third, our benchmarking extends beyond the segment-level evaluation of IBD callers and includes downstream inferences of population structure and effective population size, providing a more thorough assessment of their application in real-world analyses.

Of note, our benchmarking highlights the high false negative rate of short segments detected by isoRelate, despite it being developed for malaria parasites. The relatively poor performance of isoRelate compared to hmmIBD is likely due to differences in the underlying HMM models, where the isoRelate model assumes unphased data even when phase is provided in our benchmarking analysis. Given that isoRelate can be applied directly to unphased genotype data, it might outperform hmmIBD when genotype phasing, required by hmmIBD, is error-prone. However, this possibility was not examined in this study and warrants further investigation beyond the current scope.

The density of markers per genetic unit plays a crucial role in IBD detection. IBS-based methods, such as hap-IBD and phased IBD, first identify long IBS segments (≥ 2 centimorgan) as candidate IBD segments, and subsequently merge short ones separated by small gaps [20], allow a certain number of discordant markers to account for phasing errors [21], or eliminate false positive segments by removing candidate segments supported by only a small number of markers [20]. Similarly, Refined IBD, which combines an IBS-based method with an HMM probabilistic model, uses a LOD score to decide whether a candidate IBD segment should be rejected or accepted [44]. In these studies, default values for marker density-related parameters were shown to be effective for human genomes but have not been evaluated in high-recombining genomes. We evaluated how different levels of per-genetic-unit marker density affect the detection of IBD segments by varying recombination rates in simulations. With simulated *Pf*-like genomes of low marker density and high recombination rate, we found that human-oriented IBD callers suffer high false negative rates.

One potential explanation is that the thresholds optimized for human data are too stringent for *Pf*, causing excessive rejection of candidate segments. The effect of marker density on IBD detection is further confirmed by our findings that adjusting the values of marker density-related parameters using a grid-search approach could significantly reduce IBD error rates and generate more accurate IBD-based downstream estimates. Even though IBD accuracy can be improved by parameter optimization, we found that error rates of the detected IBD segments are still higher in *Pf* genomes than those of human genomes even after IBD caller parameters are optimized. There are several possible explanations for the high error rates of IBD segments detected from low marker density data: (1) Detected IBD segments can only start and end at the genotype marker site,

bypassing any non-genotyped sites; (2) A lower marker density is linked with greater uncertainty in the distribution of IBD endpoints [11]; (3) Ancestral relationships, including IBD, are too difficult to be reliably inferred given limited mutational information [49, 50, 58, 59].

In our benchmark analysis, we used only common biallelic SNPs as markers for inferring IBD, excluding rare variants and indels. The use of this additional information can potentially provide denser genotype information, thus enhancing our understanding of the population's ancestral relationships, a key aspect on which the inference of the IBD segment depends. For instance, large-scale whole-genome sequencing studies reveal that rare variants account for the majority of all segregating sites [36, 60], which contain crucial information for deciphering recent evolutionary history. However, rare variants are typically not utilized for two main reasons: (1) rare genotypes are very sparsely distributed across many sites and are less informative per site [8, 11]; (2) rare genotype calls are more prone to genotyping or phasing errors. As a result, including rare variation may cause reduced accuracy in detected IBD segments due to genotyping/phasing errors and increase IBD inference time due to marker density. Although our simulation analysis indicates that including rare variants is of little effect or detrimental for IBD detection (S2 Data), the findings could be skewed due to the absence of genotyping errors, the small sample size simulated, and the limited number of cut-off values tested. Further investigation in the context of large sample sizes and varying levels of genotype errors is needed to inform the usability of rare variants for different IBD detection algorithms. Indels are another significant source of underutilized genetic variations in the *Pf* genome, with abundance on par with that of biallelic SNPs in the MalariaGEN *Pf*7 data set [37]. These indels are, in part, the result of microplasticity related to the high AT content (up to 90% in non-coding regions) in *Pf* genomes [31]. Using these variants could also increase the marker density to infer IBD segments. However, additional research is necessary to determine whether the inclusion of these variants can reduce uncertainty in the inference of IBD for *Pf* or introduce more bias due to challenges such as sequence-read mapping.

While demonstrating high accuracy in IBD detection and downstream analysis, hmmIBD tends to be slower than IBS-based methods like hap-IBD. The trade-off between speed and accuracy may limit hmmIBD's suitability for large sample-size data sets. Our ongoing efforts are directed towards enhancing the computational efficiency and functionalities of the model used by hmmIBD in an adapted tool, thereby extending its applicability to large-scale data sets of relatively small and high-recombining genomes, such as *Plasmodium* genomes.

Although a substantial portion of this study concentrated on *Pf*, the main findings and methodologies may be relevant to high-recombining species beyond *Pf*. For instance, in regions with intermediate and low malaria transmission, the incidence of *Pf* has markedly decreased, allowing other species, such as *Plasmodium vivax*, to become predominant [1]. Clinical malaria caused by simian *Plasmodium* species, e.g., *Plasmodium knowlesi*, has also increased in some geographic areas where human *Plasmodium* species have declined [61]. We expect that the performance of IBD callers will be similar in other *Plasmodium* species, given their likely comparable high recombination rates [62, 63]; however, the generalization would need further exploration as part of future work, considering variations in evolutionary histories and parasite biology [64–66]. Beyond *Plasmodium* parasites, there are many other high-recombining organisms [67] such as Apicomplexan species like *Theileria* [68], insects like *Apis mellifera* (honeybee) [69, 70], and fungi like *Saccharomyces cerevisiae* (Baker's yeast) [71, 72]. For these species, our optimized parameters may not be directly applicable, but the benchmarking framework established in this study can be utilized to prioritize and optimize IBD detection methods in a context-specific manner.

While we have conducted numerous simulation analyses complemented by carefully designed validation studies, our work is subject to a few caveats: (1) Our simulations did not explicitly incorporate inbreeding within the complicated life cycle of the parasite [73], except for the

increased inbreeding potential due to the reduced population size in the single-population model. Inbreeding can be pervasive, especially in low transmission settings, leading to a change in the length distribution of IBD toward longer segments and a potentially reduced marker density. (2) Our optimization is based on simple accuracy metrics and only focuses on a subset of parameters to allow faster iteration over different values. Investigating a larger parameter space, including genotyping error rate and higher-level accuracy metrics, may generate different optimal values and further improve IBD-based downstream estimates. (3) We assume a constant recombination rate and mutation rate as static genomic/population parameters, rather than traits capable of evolving over time. If this assumption proves to be inaccurate, such as with recombination rates that vary between individuals and populations [74], a more complex benchmarking framework will be required.

In conclusion, we evaluated the performance of existing IBD segment detection methods for analyzing genomic data of the malaria parasite *Pf*, which is characterized by high recombination rates and low marker densities. Our findings underscore that a high recombination rate, relative to the mutation rate, can compromise the accuracy of detected IBD segments when using methods originally calibrated for the human genome, characterized by a significantly lower recombination rate and a higher marker density (per genetic unit). The accuracy of IBD detection can be improved by parameter optimization via grid search techniques. We advocate for a context-specific evaluation of IBD detection methods when applying them to untested species. Specifically for *Pf*, our research indicates that hidden Markov model-based probabilistic methods, such as *hmmIBD*, produce less biased IBD estimates, leading to more accurate downstream inferences. This is especially important for analyses that heavily rely on the accuracy of detected IBD segments, such as N_e inference. These findings will improve the accuracy of IBD detection and downstream analysis, providing more robust estimates of malaria transmission patterns, essential to effective malaria control and elimination efforts.

Methods

Simulation overview

We used population genetic simulations to allow the generation of (1) ground truth, including true IBD segments, true sites under positive selection, true trajectory of population size, and true subpopulation assignments (population structure), and (2) inferred patterns, including IBD inferred from phased genotype data via different IBD callers and IBD-based downstream inferences of N_e , positive selection, and population structure. By comparing inferred patterns with ground truth, we calculated metrics at the IBD segment level and the IBD-based downstream estimate level (high level) for benchmarking and optimizing various IBD detection methods for *Pf* genomes. As described in our accompanying work [10], we combined the flexible forward simulator SLiM [41, 42] and the efficient coalescent simulator *msprime* [43] to simulate genomes similar to *Pf*, reflecting the high recombination rate, strong positive selection, and population size shrinkage due to malaria reduction (Fig 1). Detailed simulation parameter values used are provided in S1 Data. A detailed implementation of the simulations can be found in a dedicated GitHub repository (https://github.com/bguo068/bmibdcaller_simulations).

In these simulations, we assumed constant recombination rates over the genome, such as 6.67×10^{-7} per base pair per generation for *Pf* [19, 75] and 1.0×10^{-8} for humans [29], and a mutation rate of 1.0×10^{-8} for both *Pf* [15, 30] and humans unless otherwise specified. Parameters for modeling population size changes, population structure, and positive selection are detailed in the following section or our related publication [10].

Simulated demographic models

We used three different demographic models in the simulations, including the single-population model, the multiple-population model, and the UK European human population model [20].

Single- and multiple-population models have been described in our accompanying work [10]. The single-population model mimics malaria reduction in settings like Southeast Asia, with a population size decreasing from 10,000 to 1,000 over the last 200 generations. This model was used to benchmark IBD detection methods at the IBD segment level and the (high) level of downstream estimation, including selection signal detection and N_e estimation. The multiple-population model was mainly used to benchmark IBD calling methods via IBD network-based community detection. Implementation of the two models is provided in a GitHub repository (see Code availability).

The UK human demographic model, similar to the one used in [20], simulates a population bottleneck event from a constant size of 10,000 to 3,000 that occurred 5,000 generations ago, followed by growth at rates of 1.4% and 25% per generation beginning 300 and 10 generations ago, respectively. We simulated 14 chromosomes with a size of 60 cM each for a smaller genome size to reduce simulation time. This demographic model serves as a control to detect IBD segments with human-oriented callers, which can help validate our IBD accuracy evaluation pipeline. We also used this model to test whether demographic models impact the performance of IBD callers by replacing the human recombination rate with that of *Pf*.

To separate the effects of positive selection from those of demographic models and recombination rates, we mostly simulated neutral genomes by setting the selection coefficient s to 0 for the above models, except in the case where we needed to benchmark IBD callers for detecting positive selection signals.

Positive selection simulation

To evaluate the performance of IBD callers via IBD-based selection signal detection, we simulated positive selection within the single-population model with selection coefficients s of 0.2 with a single origin starting 80 generations ago. Fourteen chromosomes with a size of 100 cM were simulated independently, each with a selected site at 33.3 centimorgans from their left ends. For selection simulation, we conditioned on the establishment of the selective sweeps; that is, the allele under selection should not be lost in the present-day generation. If lost, the simulation was rerun for a maximum of 100 times until the selective sweep was established.

IBD calling and default parameters

We generated true IBD from simulated genealogical trees in the tree sequences using the tsKibd algorithm [10]. Briefly, we sampled local/marginal trees along a chromosome and tracked changes in the most recent common ancestor (MRCA) for each pair of sample nodes. If the MRCA changes, the shared ancestral segment breaks. We report a shared ancestral segment as an IBD segment if its length exceeds a threshold, such as 2 centimorgans.

For inferred IBD, we used phased genotype data as input. Only biallelic sites with a minor allele frequency no less than 0.01 were included, unless otherwise stated. When needed, a genetic map is generated based on the constant recombination rate as specified in the corresponding simulations. For IBD detection methods designed for diploids, we converted each haploid to a pseudo-homozygous diploid. The resulting IBD segments of each pair of pseudo-homozygous diploids (A1/A2 and B1/B2) have redundant information due to the 100% runs of homozygosity. We only keep one pair, A1-B1, and remove other combinations.

As each IBD detection method provides multiple tunable parameters, we detailed values used in S1 Table for both default and optimized scenarios. For the default scenarios, the parameters mostly follow the original documentation. Exceptions are parameters that need consistency across different IBD callers for benchmarking, including the minimum IBD length and the minimum minor allele frequency. The process of obtaining optimized parameter values is described below.

Benchmarking metrics at the IBD segment level

Several metrics were calculated to benchmark IBD methods at the IBD segment level or using their simple aggregates, including the false negative rate and false positive rate, pairwise total IBD, and population-level total IBD per length bin.

False positive rate and the false negative rate were obtained following the definition used in the work of Zhou *et al.* [20]. Rates were first calculated per segment via segment-overlapping analysis; then, they were averaged over all segments of the same length bin. The false positive rate per segment is defined as the proportion of a given IBD segment of some genome pair from the inferred set (for example, IBD called via *hmmIBD*) that is not covered by any IBD segment from the truth set (generated by *tskibd*), for the same genome pair. Similarly, the false negative rate per segment is defined as the proportion of a given IBD segment of some genome pair from the truth set that is not covered by any IBD segment from the inferred set. The average false positive rate is calculated as the average per-segment false positive rates for all inferred IBD segments the length of which falls in a certain range (length bin); the average false negative rate is calculated as the average per-segment false negative rates for all true segments of a length bin. The following length bins were used: [3-4), [4-6), [6-10), [10-18), [18, inf) cM, similar to Zhou *et al.*'s method.

Genome-wide FP/FN rates per pair and their averages across all genome pairs were calculated to capture genome-wide bias. The per-pair genome-wide false positive rate is the ratio of two sums:

(1) the numerator sum is the total length of parts of all inferred IBD segments of a certain genome pair that are not covered by any true IBD segments of the same genome pair; (2) the denominator sum is the total length of all inferred IBD segments of a certain genome pair. The per-pair genomewide false negative rate is defined in a similar way as the percentage of pairwise total true IBD that is not overlapped by any inferred segments of the same pair. We then obtain the aggregate metrics by averaging these rates for all genome pairs.

Pairwise total IBD from truth *versus* inferred set was calculated as it's a useful metric to estimate genetic relatedness and build IBD sharing networks. It was calculated as the sum of the lengths of all inferred or true IBD segments of each genome pair.

Given that the N_e estimator *IBDNe* internally utilizes quantities of population-level total IBD of different length bins, these quantities were calculated here to better examine IBD accuracy for N_e inference. We defined non-overlapping length bins of 0.05 centimorgan width that cover all possible lengths. For each length bin, a population total IBD was defined as the sum of the length of all IBD segments with segment lengths falling into this bin from any genome pair.

To expedite IBD segment-level analysis and alleviate computational burdens, we developed an open-source tool, *ishare/ibdutils* (available at <https://github.com/bguo068/ishare>), which harnesses algorithms such as interval trees to efficiently calculate metrics like those described above.

IBD caller parameter optimization

We optimized key IBD caller-specific parameters by iterating each parameter over the list of discrete values, or two or more parameters over a grid of discrete values. Many optimized parameters were related to marker density for IBD callers *hap-IBD*, *phased IBD*, *Refined IBD*, and

isoRelate. Other parameters were searched, if they potentially have a great impact on the quality of the detected IBD. The optimal values for explored parameters are determined by the length-bin-specific or genome-wide error rates (FN and FP) for detected IBD segments as defined above. The parameter value or combination of values that generates lower and generally balanced error rates was selected as optimal values. The parameters searched, the value lists explored, and the optimal values selected are summarized in S1 Table and S2 Table. When the optimized values vary across different demographic models, we used the ones optimized for the single-population model for downstream analyses. We provided detailed simulation and IBD calling parameter values in S1 Data and heatmaps of error rates for all demographic models tested in S2 Data.

Benchmarking via IBD-based downstream analyses

At a higher level, we benchmarked different IBD callers by comparing downstream estimates based on true IBD sets *versus* inferred IBD sets. These IBD-based estimates include positive selection scans, N_e estimates, and population structure inference via the community detection algorithm InfoMap.

We scanned for positive selection signals using the IBD-based thresholding method followed by validation with integrated haplotype score-based statistics X_{iHS} as previously described [10].

We inferred the trajectory of effective population size, N_e , for the last 100 generations using IBDNe. As this method uses IBD shared by diploid individuals as input, we converted each haploid genome to pseudo-homozygous diploid individuals. We inferred the trajectory N_e using most of the default parameters except for setting the minregion parameter to 10 cM to allow the inclusion of short contigs in the analysis. The final estimates are scaled by 0.25 to compensate for the haploid-to-diploid conversion. By default, we used the value of 2 cM for the mincm parameter to only include IBD segments ≥ 2 cM for inferring the N_e trajectory; as indicated in the Results section, we also set mincm to 4 to test whether excluding short IBD segments can improve the accuracy of N_e estimates. As the IBDNe algorithm does not work well when IBD segments shared by close relatives are included in the input, we followed the procedure described in the original work by Browning *et al.* [22]. For simulated data, we utilized the TMRCA information of the true IBD segments to filter the IBD segments before calling IBDNe. For true IBD segments (generated by tskibd), we excluded IBD segments with TMRCA < 1.5 ; for inferred IBD segments (called by hap-IBD, hmmIBD, isoRelate, Refined IBD, phased IBD), we removed (inferred) segments that overlap with any true IBD segment with TMRCA < 1.5 shared by the same genome pair. For empirical data where true IBD and TMRCA are not available, we pruned highly related isolates by iteratively removing the genome that has the highest number of close relatives defined by pairwise total IBD > 0.5 of genome size until no close relatives are present in the remaining subgroup [10].

For population structure inference, we first built the pairwise total IBD matrix, each element being the total IBD for a pair of genomes. The matrix was then squared and used as a weighted adjacency matrix to construct an IBD-sharing network. We then ran the InfoMap algorithm to infer community membership. Genomes assigned the same membership were inferred to be of the same subpopulation. We calculated an adjusted Rand index using the igraph-python package [46] to analyze the agreement between true population labels and inferred community labels. For empirical data, we excluded IBD segments shorter than 4 cM when calculating the total IBD matrix to help reduce noise due to false positives and set each element with a value < 5 to zero in the unsquared IBD matrix to decrease the density of the IBD matrix.

Before all high-level benchmarking analyses, we pruned highly related samples as mentioned above. As IBD-based estimates can be biased by strong positive selection in empirical data, we removed high IBD-sharing peaks with peak impact index > 0.01 as previously described [10] before any downstream analyses.

Processing empirical data sets

We constructed empirical data sets for validation using genotype data from whole-genome sequencing samples from the MalariaGEN *Pf7* database [37]. We used *malariagen_data* 7.13 to download high-quality monoclonal samples that pass quality control. Monoclonal samples were determined by $F_{ws} > 0.95$ (F_{ws} table available from the MalariaGEN website). Quality control labels were extracted directly from the metadata (provided in the *malariagen_data* package).

We then generated haploid genomes (phased genotype data) using the dominant allele from genotype calls. The dominant allele of each genotype call was determined by the per-sample allele depth (AD) fields. For each sample and site, the allele supported by 90% of total reads (total AD values) in a genotype call with at least 5 total reads was used as the dominant allele. Genotype calls without dominant alleles were marked as missing; those with dominant alleles were replaced with a phased genotype homozygous for the dominant allele. The genotype data were further filtered by sample missingness and SNP minor allele frequency and missingness. The resulting genotype data had per-SNP and per-sample missingness < 0.1 and minor allele frequency ≥ 0.01 . Genotypes based on dominant alleles were further imputed without a panel using Beagle 5.1 [76]. These processing steps generated phased, imputed, pseudo-homozygous diploid genotype data, ready for IBD detection.

We constructed different data sets, including two “single” population data sets and a “structured, multiple” population data set, by subsampling the above haploid genomes according to sampling time and location. For each data set, we set the time window to 2-3 years to reduce the sample time heterogeneity, and then shifted the window within all possible sampling years and chose one that maximized the sample size. For the “single” population data sets, we further restricted the sample locations to a relatively small geographic region, such as eastern Southeast Asia, as the data set was used for N_e estimation, which assumes a homogeneous population. For the “multiple” population data set, we included samples from different continental or subcontinental regions using “Population” labels from the meta-information table provided with the MalariaGEN *Pf7* database.

To make the sample size of each “population” more balanced, we set a maximum number of samples of 300. Populations with samples larger than 300 are subsampled to a size of 300; populations with a size smaller than 100 were not included in the multiple-population data set. The details of the sampling location and time information were summarized in S3-S5 Tables.

Details about the preparation and analysis of the empirical data sets can be found at https://github.com/bgao068/bmibdcaller_empirical.

Measuring computational runtime and memory usage

The genomes were simulated with $N_0 = 1,000$ and $s = 0.0$ under the single-population model. The runtime and maximum memory usage were measured using the GNU time 1.7 utility. To allow a more appropriate comparison, we ensured: (1) the time used for data pre-and post-processing was excluded; (2) memory resource allocation was capped at 30 gigabytes per IBD call; (3) input genotype data included only common SNPs with minor allele frequency ≥ 0.01 ; (4) the minimum reported IBD segment length was set to 2.0 cM.

Replications and uncertainty of measures

Replications of simulations are conducted at two levels: (1) Replication of simulation sets: for each combination of simulation parameters, we performed $n = 3$ full sets of simulations of populations and sampled $n = 1,000$ haploid genomes per population. (2) Each of the 14 chromosomes in the

genomes of a population was simulated independently, which are replicates of each other (see S1 Data for detailed simulation parameters and replications at simulation set and chromosome levels).

Different analyses reported the uncertainty of measures at various levels of replications or units of observations (as mentioned in the corresponding figure legends and S2 Data). The choice was made based on the level of uncertainty most relevant to the measures. (1) For IBD accuracy-based IBD segment overlapping analysis, the mean $\pm$ standard deviation (SD) was calculated at the segment level for IBD segment false positive and false negative rates for each length bin, or at the genome-pair level for IBD genome-wide error rates. (2) For IBD-based genetic relatedness, the uncertainty is directly visualized in scatter plots at the genome-pair level. (3) For IBD-based selective signal scans, the mean $\pm$ SD of the number of true selection signals (peaks) and false selection signals were calculated at the simulation set level ($n = 3$ full simulation sets). (4) For IBD network community detection, the mean $\pm$ SD of the adjusted Rand index was reported at the simulation set level ($n = 3$). (5) For IBD-based N_e estimates, bootstrap confidence intervals were obtained directly from IBDNe using a single simulation set. (6) For the measure of computational efficiency and memory usage, the mean $\pm$ SD was calculated across chromosomes from the same simulation sets.

Given our large sample size of 1,000 haploid genomes, the uncertainty reported at the simulation set level is relatively small and can be measured with a limited number of replications. Additionally, full sets of simulation replications were computationally intensive. Therefore, we opted to run $n = 3$ full simulation sets when it was necessary to measure uncertainty at the simulation set level. For measures for which uncertainty was reported at the segment level or genome-pair levels, only results from a representative simulation set were reported if the results were consistent across $n = 3$ simulation sets.

Acknowledgements

This publication uses MalariaGEN data as described in “Pf7: an open dataset of *Plasmodium falciparum* genome variation in 20,000 worldwide samples” MalariaGEN et al, Wellcome Open Research 2023, 8:22 <https://doi.org/10.12688/wellcomeopenres.18681.1>. This work was supported by NIH 1R01AI145852 granted to ST-H and TDO by the U.S. National Institutes of Health.

Additional information

Code availability

Custom tools or scripts were provided in the following GitHub repositories: (1) `bmibdcaller_simulations`: a Nextflow pipeline to benchmark different IBD detection methods and optimize IBD caller-specific parameters by simulating *Plasmodium falciparum*-like genomes and using true IBD (https://github.com/bgao068/bmibdcaller_simulations, v0.1.0). (2) `bmibdcaller_empirical`: a Nextflow pipeline to benchmark IBD callers with empirical data by comparing IBD-based estimates with expected patterns (https://github.com/bgao068/bmibdcaller_empirical, v0.1.0). (3) `ishare/ibdutils`: a Rust crate and command-line tools designed to facilitate the analysis of rare-variant sharing and identity-by-descent (IBD) sharing (used here mainly for fast IBD segment over-lapping analysis) (<https://github.com/bgao068/ishare>, v0.1.11, command-line tool `ibdutils`).

Data availability

All empirical data used was publicly available from MalariaGEN Pf7 [37].

Supplementary figure and table legends

S1 Fig. High recombination rates affect the quality of detected IBD segments (extending Fig 2). a-b, Accuracy of detected IBD segments from genomes simulated with different recombination rates.

The accuracy of IBD segments is measured by the false negative rates (a) and the false positive rates (b). The plotted error rates reflect the genome-wide rates (as defined in Methods) of IBD segments called via default IBD caller parameters. These rates are based on one representative set from $n = 3$ simulation sets. The plotted vertical lines indicate standard deviations of error rates across all genome pairs in the representative simulated set. Both the vertical lines and markers are horizontally staggered for clarity.

Note, for both (a) and (b), the genomes were simulated under the single-population model (Methods). Data for low recombination rates is missing for isoRelate as genome sizes in base-pairs are too large for isoRelate to finish IBD calling in a timely manner. Note that x axes in both (a) and (b) use log scales.

S2 Fig. Pairwise total IBD of simulated *Pf*-like genomes from the different IBD callers compared with those based on true IBD (before parameter optimization). Quality of detected IBD as measured by genome-wide pairwise total IBD (total length of IBD segments ≥ 2 cM for each genome pair), an estimate of genetic relatedness. Data of $n = 1,000$ haploid genomes from a representative simulation set was plotted with each subplot for an IBD caller (as indicated in the titles). Each dot represents a genome pair with the x and y coordinates being true and the inferred total IBD. The blue dots are pairs with nonzero true and inferred total IBD, whereas red dots are pairs with the true total IBD or the inferred total IBD being 0; “0” total IBD was replaced with 1.0 cM for visualization. The red dotted line of $y = x$ indicates the expected pattern, i.e., true total IBD equal to inferred total IBD if the inferred IBD was 100% accurate. Note that log scales are used in all x and y axes in (b).

S3 Fig. IBD caller-specific parameter optimization can improve the quality of IBD segments inferred from simulated *Pf* genomes (extending Fig 4). a, Quality of detected IBD measured by false positive and false negative rates before (left column) and after (right column) IBD caller-specific parameter optimization. As indicated in the legend of the axes, the error rates were calculated for different IBD segment length ranges (in centimorgans), including [3-4], [4-6], [6-10], [10-18], [18, inf), and at the genome-wide level. These rates are based on one representative set from $n = 3$ simulation sets. The plotted vertical and horizontal lines represent the standard deviations of error rates (horizontal for false negatives and vertical for false positives) calculated across all relevant segments for length-bin specific rates or across all genome pairs for genome-wide rates in the representative simulated set. b, Quality of detected IBD measured by genome-wide pairwise total IBD, an estimate of genetic relatedness, before (left column) and after (right column) parameter optimization. Data of $n = 1,000$ haploid genomes from a representative simulation set was plotted with each subplot for each IBD caller (indicated in shaded vertical labels) before and after parameter optimization (indicated in the titles). Each dot represents a genome pair with the x and y coordinates being true and inferred total IBD. In (b), the blue dots are pairs with non-zero true and inferred total IBD. Whereas, red dots are pairs with either true total IBD or inferred total IBD being 0; “0” total IBD was replaced with 1.0 cM for visualization. The red dotted line of $y = x$ indicates the expected pattern, that is, the true total IBD equal to the inferred total IBD if the inferred IBD was 100% accurate. Note: This figure is similar to Fig 4 except that it includes the results for the five IBD callers analyzed. For both (a) and (b), each row represents results for a caller, as indicated by the vertical text label with a gray background on the left of the figure. Note that log scales are used in all x and y axes.

S4 Fig. Quality of IBD segments called via human-oriented callers from simulated genomes with UK human population demographic history but different recombination rates (human versus *Pf* rates). a, Quality of detected IBD as measured by false positive and false negative rates from genomes simulated with human (left axes) and *Pf* (right axes) recombination rates. As indicated in the legend of the axes (at the bottom of the panel), the error rates were calculated for different length ranges (centimorgan), including [3-4], [4-6], [6-10], [10-18], [18, inf) and at the genome-wide level. These rates are based on one representative set from $n = 3$ simulation sets. The plotted vertical and horizontal lines represent the standard deviations of error rates (horizontal for false negatives and vertical for false positives) calculated across all relevant segments for length-bin specific rates or across all genome pairs for genome-wide rates in the representative simulated set. b, Quality of detected IBD measured by genome-wide pairwise total IBD, an estimate of genetic relatedness, from genomes simulated with human (left column) and *Pf* (right column) recombination rates. Data of $n = 1,000$ haploid genomes from a representative simulation set using human-like versus *Pf*-like recombination rates (indicated in the titles) was plotted with each subplot for each IBD caller (indicated in shaded vertical labels). Each dot represents a genome pair with the x and y coordinates being true and inferred total IBD. Note: (1) In (b), the blue dots are pairs with non-zero true and inferred total IBD, whereas red dots are pairs with either true total IBD or inferred total IBD being 0; "0" total IBD was replaced with 1.0 cM for visualization. The red dotted line of $y = x$ indicates the expected pattern, i.e., true total IBD equal to inferred total IBD if the inferred IBD was 100% accurate. (2) Genomes were simulated under the demographic model mimicking the history of the UK human population (described in Zhou *et al.* 2020 [20]). Note that log scales are used in all x and y axes in (b).

S5 Fig. Post-optimization benchmarking of different IBD callers by comparing downstream estimates of selected loci and population structure (extending Figure 4). a, IBD coverage and estimates of selected loci based on IBD segments called after parameter optimization. Each subplot corresponds to one of the five benchmarked IBD callers. In each subplot, lines of different colors represent the IBD coverage over one of 14 chromosomes; text in the upper right corner indicates the observed (on the left of "/") and expected (on the right of "/") numbers of true peaks and false peaks. A peak is defined as true if the span of the peak contains a selected site determined from the simulation. b, Post-optimization inference of population structure via the InfoMap algorithm. In each subplot, the row labels indicate the true population labels set in the simulation, and the column labels indicate the detected community labels. C1 - C5 represent the largest five detected communities. The color within each block represents the number of samples assigned to the given true population label and the detected community label. Columns are rearranged for better visualization. Text in the upper right corner of each plot are adjusted rand indices, indicating the consistency between true population labels (simulation) and detected community labels (detected). For both (a) and (b), the numbers of true and false peaks and adjusted rand indices are presented in the form of mean $\pm$ standard deviation as determined from $n = 3$ independent simulation sets. See S8 Fig for pre-optimization results.

S6 Fig. Population-level total IBD per length bin. The population-level total length of IBD segments with a given length range is an intermediate quantity used by IBDNe internally to estimate the trajectory of N_e at different times in recent history. The population-level total IBD inferred by an IBD caller (as shown in the title of each subplot) is compared with that of true IBD. The log ratio of inferred to true population-level total IBD was plotted against length bins (represented by bin centers in centimorgans). Blue dots represent a length bin with non-zero population-level total IBD for both true and inferred sets.

Red dots indicate either the true or inferred population-level total IBD is zero for the corresponding length bin (x coordinate). A line of $y = 0$ is added as the expected line, meaning that the log ratios for all length bins would be 0 if there were no biases. The data shown here is based on $n = 1000$ haploid genomes for one representative set of $n = 3$ simulation replication sets. Note that log scales are used for all x axes.

S7 Fig. The effect of including and excluding short IBD segments (< 4cM) on N_e estimation.

a. N_e estimates for genomes simulated under the single-population model without selective sweep. In the upper row, IBD segments with TMRCA < 1.5 generations were excluded from the IBD input for N_e estimation; in the bottom row, they are kept. For true IBD segments (from tskibd), the TMRCA for each IBD segment was obtained directly from genealogical trees; for inferred IBD segments, any of them that overlapped with a segment of the truth set with TMRCA < 1.5 generations was determined to have TMRCA < 1.5 generations. The plots show results from one representative set of $n = 3$ replicated simulation sets. b. N_e estimates for genomes for different empirical data sets. Each row corresponds to a data set labeled on the right. The columns are estimates for different IBD callers. Solid red line: N_e estimates when short IBD segments are included; red dashed line: N_e estimates when short IBD segments are excluded; black dashed line: parameter population sizes used in simulations. For (a) and (b), log scales are used for all y axes, and shaded areas indicate the 95% bootstrap confidence intervals as determined by IBDNe.

S8 Fig. Pre-optimization benchmarking of different IBD callers by comparing downstream estimates of selected loci and population structure (for comparison with S5 Fig). a, IBD coverage and estimates of selected loci based on IBD segments called before parameter optimization. Each subplot corresponds to one of five benchmarked IBD callers. In each subplot, lines of different colors represent the IBD coverage over one of 14 chromosomes; text in the upper right corner indicates the observed (on the left of the “/”) and expected (on the right of “/”) numbers of true peaks and false peaks. A peak is defined as true if the span of the peak contains selected sites determined from simulation. b, Post-optimization inference of population structure via InfoMap algorithm. In each subplot, the row labels indicate the true population labels set in the simulation and the column labels indicate the detected community labels. C1 - C5 represent the largest five detected communities. The color within each block represents the number of samples assigned to the given true population label and the detected community label. Columns are rearranged for visualization purposes. Text in the upper right corner of each plot is adjusted rand indices, indicating the consistency between true population labels (simulation) and detected community labels (detected). For both (a) and (b), the numbers of true and false peaks and adjusted rand indices are presented in the form of mean $\pm$ standard deviation as determined from $n = 3$ independent simulation sets.

S9 Fig. Pre-optimization benchmarking of different IBD callers by comparing downstream estimates of N_e (for comparison with [Fig 5](#)). With mainly default values for IBD caller parameters, the performance of IBD callers was evaluated by comparing the N_e trajectory for the recent 100 generations estimated via IBDNe based on true (black dashed line) IBD versus inferred IBD (red solid line). True IBD was calculated from simulated genealogical trees via tskibd; inferred IBD includes those inferred from tskibd, hap-IBD, hmmIBD, isoRelate, Refined IBD and phased IBD, with their N_e estimates shown from left to right. Red shaded areas in each plot indicate the 95% confidence interval for N_e estimates as generated by IBDNe. The plots show results from one representative set of $n = 3$ replicated simulation sets. Note that log scales are used for all y axes.

S10 Fig. Performance of IBD callers with empirical data sets before parameter optimization (in comparison with post-optimization performance shown in [Fig 6](#)). a, IBD coverage and detected selection signals in the SEA data set using different IBD callers (rows 1 to 5). Annotations and corresponding vertical dotted lines at the top indicate the centers of known and putative drug resistance genes and genes related to sexual commitment; red shading indicates regions that are inferred to be under positive selection (see Methods for definitions). Note the difference in the peak height of IBD and the counts of hap-IBD and phased IBD compared to [Fig 6a](#). b, N_e estimates of the SEA data set based on IBD inferred from different callers. The line plots are the point estimates; the shaded areas around the line plots indicate confidence intervals based on bootstrapping (generated by IBDNe). c, Inference of the population structure of the structured data set by the InfoMap community detection algorithm using IBD network inferred from different IBD

callers. The rows of the heatmap are geographic regions of isolates, the columns of the heatmap are the largest, inferred communities, labeled as c1 to c6. The color of the heatmap represents the number of isolates in each block with the given row and column labels. The columns are rearranged such that diagonal blocks tend to have the largest values per row for better visualization. Note the difference in size of the main detected communities for hap-IBD and phased IBD for parasites from Africa when compared with [Fig 6c](#). Note that log scales are used for all y axes in (b).

S11 Fig. Comparison of maximum computer memory used for IBD calling process for different callers. a, Memory usage for different IBD callers to detect IBD from genomes of different sizes in single-thread mode. The comparison is based on *Pf* genomes simulated under the single-population model. The x-axis tick labels include the number of pairs of genomes simulated and analyzed (below the plot, reflecting the number of computation units) and the number of haploid genomes (above the plot, representing sample size) analyzed. The line styles and markers for different callers/tools are provided in the legend box on the far right of the figure, which is shared across the two subplots. Values on the y axis represent means (markers) and standard deviations (vertical error lines) from $n = 3$ sets of independent simulations. Note that at each sample size, the x values for different IBD callers are slightly staggered to prevent the error lines from overlapping. b, Memory usage under the multi-threading mode. (b) is organized similarly to (a), except that the IBD calling processes were run under the multithreading mode with 10 threads. Note that log scales are used on y axes and bottom x axes.

S1 Table. The default and optimized values for parameters used in inferring IBD segments via different callers. A link to the source code and a citation of the corresponding article are provided for each IBD caller. For hap-IBD, Refined IBD, and hmmIBD, the parameters are used on the command line except that the parameter `rec_rate` of hmmIBD needs to be specified in the source code file `hmmIBD.c`. For phased IBD and isoRelate, the parameters are specified within a Python or R script. The details of how the parameters are specified can be found in the scripts on GitHub (https://github.com/bgao068/bmibdcaller_simulations/tree/main/bin). Note that `mincm` and `minmaf` are values shared across IBD callers to allow fair comparisons.

S2 Table. Grid/Line search to optimize IBD caller-specific parameters. Column 2 lists the optimized parameters (other parameters are not explored). Column 3 shows the tested values for each parameter. Some parameters are optimized via two steps, such as coarse search (coarse) and fine-tuning (fine-tune). Column 4 provides comments on the impact of the values on IBD accuracy (measured as false positive rates (FP) and false negative rates (FN)).

S3 Table. Isolates in the “multi-population” data set used in empirical validation. Rows are counts of isolates from different locations (“Population” labels from MalariaGEN *Pf*7 meta information table).

S4 Table. Isolates in the “single-population” data set for “AS-SE-E” population used in empirical validation. Rows are counts of isolates from different locations (“Country” labels from MalariaGEN *Pf*7 meta information table); columns are counts of isolates collected in a given year.

S5 Table. Isolates in the “single-population” data set for “AF-W-Ghana” population used in empirical validation. Rows are counts of isolates from different locations (“admin level 1” labels from MalariaGEN *Pf*7 meta information table); columns are counts of isolates collected in a given year.

S1 Data. Detailed simulation and IBD calling parameters and values. The Excel file contains 18 tabs. Tab 1 provides a summary of the tables, including a schematic of the parameter combinations for 3 demographic models and 5 IBD callers, parameter descriptions, and a list of the names of the other tabs. Tab 2 details the simulation replications at the chromosome and

simulation set levels. Tab 3 presents all parameter values used in the simulations and IBD detection. Tabs 4-18 contain the parameter values for each demographic model-IBD caller combination.

S2 Data. Detailed IBD-level benchmarking results shown in heatmaps. Each panel (a-r) represents the FN/FP rates for a specific combination of IBD callers (hap-IBD, hmmIBD, isoRelate, phased IBD, and Refined IBD), demographic models (single-population model, multiple-population model, and UK human demographic model), and recombination rates (human *versus* *Pf*) as indicated in the text above the panel. Note that benchmarking for genomes with human recombination rates was not performed for hmmIBD and isoRelate as they did not scale well for large genome sizes (in base pairs). For each heatmap, the searched parameters and their values are indicated as the *x* and *y* labels and tick labels, respectively; the labels in grey background on the top indicate the IBD length bin that was used to calculate FN/FP rates (labels in grey on the left); the bold labels in grey at the bottom show either different groups (e.g., coarse search or fine-tune) or the third parameter searched (such as min-maf=0.01 and min-maf=0.00).

Funding

National Institutes of Health (1R01AI145852)

Additional files

Supplementary Information [↗](#)

S1 Data [↗](#)

References

1. World Health Organization (2024) **World Malaria Report** World Health Organization <https://www.who.int/teams/global-malaria-programme/reports/world-malaria-report-2024> | [Google Scholar](#)
2. Neafsey D. E., Taylor A. R., MacInnis B. L (2021) **Advances and Opportunities in Malaria Population Genomics** *Nature Reviews Genetics* :1–16 [Google Scholar](#)
3. Wesolowski A., et al. (2018) **Mapping Malaria by Combining Parasite Genomic and Epidemiologic Data** *BMC medicine* **16**:190–190 [Google Scholar](#)
4. Shetty A. C., et al. (2019) **Genomic Structure and Diversity of Plasmodium Falciparum in Southeast Asia Reveal Recent Parasite Migration Patterns** *Nature Communications* **10**:1–11 [Google Scholar](#)
5. Taylor A. R., Jacob P. E., Neafsey D. E., Buckee C. O (2019) **Estimating Relatedness between Malaria Parasites** *Genetics* **212**:1337–1351 [Google Scholar](#)
6. Henden L., Lee S., Mueller I., Barry A., Bahlo M (2018) **Identity-by-Descent Analyses for Measuring Population Dynamics and Selection in Recombining Pathogens** *PLoS genetics* **14**:e1007279–e1007279 [Google Scholar](#)
7. Gerlovina I., Gerlovin B., Rodríguez-Barraquer I., Greenhouse B (2024) **Dcifer: An IBD-based Method to Calculate Genetic Distance between Polyclonal Infections** *Genetics* **222**:iyac126 [Google Scholar](#)
8. Schaffner S. F., Taylor A. R., Wong W., Wirth D. F., Neafsey D. E (2018) **HmmIBD: Software to Infer Pairwise Identity by Descent between Haploid Genotypes** *Malaria Journal* **17**:10–13 [Google Scholar](#)
9. Schaffner S. F., et al. (2024) **Malaria Surveillance Reveals Parasite Relatedness, Signatures of Selection, and Correlates of Transmission across Senegal** *Nature Communications* **14**:7268 [Google Scholar](#)
10. Guo B., et al. (2024) **Strong Positive Selection Biases Identity-by-Descent-Based Inferences of Recent Demography and Population Structure in Plasmodium Falciparum** *Nature Communications* **15**:2499 [Google Scholar](#)
11. Browning S. R., Browning B. L (2020) **Probabilistic Estimation of Identity by Descent Segment Endpoints and Detection of Recent Selection** *The American Journal of Human Genetics* **107**:895–910 [Google Scholar](#)
12. Morgan A. P., et al. (2020) **Falciparum Malaria from Coastal Tanzania and Zanzibar Remains Highly Connected despite Effective Control Efforts on the Archipelago** *Malaria Journal* **19**:1–14 [Google Scholar](#)
13. Nait Saada J., et al. (2023) **Identity-by-Descent Detection across 487,409 British Samples Reveals Fine Scale Population Structure and Ultra-Rare Variant Associations** *Nature Communications* **14** [Google Scholar](#)

14. Al-Asadi H., Petkova D., Stephens M., Novembre J. (2019) **Estimating Recent Migration and Population-Size Surfaces** *PLOS Genetics* **15**:e1007908 [Google Scholar](#)
15. Camponovo F., Buckee C. O., Taylor A. R (2024) **Measurably Recombining Malaria Parasites** *Trends in Parasitology* **39**:17–25 [Google Scholar](#)
16. Guo B., Rowley E., O'Connor T. D., Takala-Harrison S (2025) **Potential and Pitfalls of Using Identity-by-Descent for Malaria Genomic Surveillance** *Trends in Parasitology* **41**:387–400 <https://www.sciencedirect.com/science/article/pii/S1471492225000923> | [Google Scholar](#)
17. Wong W., et al. (2025) **MalKinID: A Classification Model for Identifying Malaria Parasite Genealogical Relationships Using Identity-by-Descent** *Genetics* **229**:iyae197 <https://doi.org/10.1093/genetics/iyae197> | [Google Scholar](#)
18. Taylor A. R., et al. (5595) **Resolving the Cause of Recurrent Plasmodium Vivax Malaria Probabilistically** *Nature Communications* **10** [Google Scholar](#)
19. Amambua-Ngwa A., et al. (2019) **Major Subpopulations of Plasmodium Falciparum in Sub-Saharan Africa** *Science* **365**:813–816 [Google Scholar](#)
20. Zhou Y., Browning S. R., Browning B. L (2020) **A Fast and Simple Method for Detecting Identity-by-Descent Segments in Large-Scale Data** *American Journal of Human Genetics* **106**:426–437 [Google Scholar](#)
21. Freyman W. A., et al. (2021) **Fast and Robust Identity-by-Descent Inference with the Templated Positional Burrows-Wheeler Transform** *Molecular Biology and Evolution* **38**:2131–2151 [Google Scholar](#)
22. Browning S. R., Browning B. L (2015) **Accurate Non-parametric Estimation of Recent Effective Population Size from Segments of Identity by Descent** *American Journal of Human Genetics* **97**:404–418 [Google Scholar](#)
23. Gutenkunst R. N., Hernandez R. D., Williamson S. H., Bustamante C. D (2009) **Inferring the Joint Demographic History of Multiple Populations from Multidimensional SNP Frequency Data** *PLOS Genetics* **5**:e1000695 [Google Scholar](#)
24. Joy D. A., et al. (2023) **Early Origin and Recent Expansion of Plasmodium Falciparum** *Science* **300**:318–321 [Google Scholar](#)
25. Gardner M. J., et al. (2002) **Genome Sequence of the Human Malaria Parasite Plasmodium Falciparum** *Nature* [Google Scholar](#)
26. Su, X, et al. (1999) **A Genetic Map and Recombination Parameters of the Human Malaria Parasite Plasmodium Falciparum** *Science* **286**:1351–1353 [Google Scholar](#)
27. Jiang H., et al. (2011) **High Recombination Rates and Hotspots in a Plasmodium Falciparum Genetic Cross** *Genome Biology* [Google Scholar](#)
28. Miles A., et al. (2016) **and Recombination Drive Genomic Diversity in Plasmodium Falciparum** *Genome Research* **26**:1288–1299 [Google Scholar](#)

29. Kong A., et al. (2002) **A High-Resolution Recombination Map of the Human Genome** *Nature genetics* **31**:241–247 [Google Scholar](#)
30. Bopp S. E. R., et al. (2023) **Mitotic Evolution of Plasmodium Falciparum Shows a Stable Core Genome but Recombination in Antigen Families** *PLOS Genetics* **9**:e1003293 [Google Scholar](#)
31. Hamilton W. L., et al. (2017) **Extreme Mutation Bias and High AT Content in Plasmodium Falciparum** *Nucleic Acids Research* **45**:1889–1901 [Google Scholar](#)
32. McDew-White M., et al. (2023) **Mode and Tempo of Microsatellite Length Change in a Malaria Parasite Mutation Accumulation Experiment** *Genome Biology and Evolution* **11**:1971–1985 [Google Scholar](#)
33. Huber J. H., Johnston G. L., Greenhouse B., Smith D. L., Perkins T. A. (2016) **Quantitative, Model-Based Estimates of Variability in the Generation and Serial Intervals of Plasmodium Falciparum Malaria** *Malaria Journal* **15**:490 [Google Scholar](#)
34. Churcher T. S., et al. (2014) **Public Health. Measuring the Path toward Malaria Elimination** *Science (New York, N.Y.)* **344**:1230–1232 [Google Scholar](#)
35. Campbell C. D., et al. (2024) **Estimating the Human Mutation Rate Using Autozygosity in a Founder Population** *Nature Genetics* **44**:1277–1281 [Google Scholar](#)
36. Taliun D., et al. (2024) **Sequencing of 53,831 Diverse Genomes from the NHLBI TOPMed Program** *Nature* **590**:290–299 [Google Scholar](#)
37. MalariaGEN, et al. (2023) **Pf7: An Open Dataset of Plasmodium Falciparum Genome Variation in 20,000 Worldwide Samples** *Wellcome Open Research* **8**:22 [Google Scholar](#)
38. Tang K., Naseri A., Wei Y., Zhang S., Zhi D (2022) **Open-Source Benchmarking of IBD Segment Detection Methods for Biobank-Scale Cohorts** *GigaScience* **11**:giac111 [Google Scholar](#)
39. Shemirani R., et al. (2021) **Rapid Detection of Identity-by-Descent Tracts for Mega-Scale Datasets** *Nature Communications* **12**:3546 [Google Scholar](#)
40. Browning B. L., Browning S. R. (2011) **A Fast, Powerful Method for Detecting Identity by Descent** *The American Journal of Human Genetics* **88**:173–182 [Google Scholar](#)
41. Haller B. C., Messer P. W (2019) **SLiM 3: Forward Genetic Simulations Beyond the Wright-Fisher Model** *Molecular Biology and Evolution* **36**:632–637 [Google Scholar](#)
42. Haller B. C., Galloway J., Kelleher J., Messer P. W., Ralph P. L (2019) **Tree-Sequence Recording in SLiM Opens New Horizons for Forward-Time Simulation of Whole Genomes** *Molecular Ecology Resources* **19**:552–566 [Google Scholar](#)
43. Baumdicker F., et al. (2022) **Efficient Ancestry and Mutation Simulation with Msprime 1.0** *Genetics* **220**:iyab229 [Google Scholar](#)
44. Browning B. L., Browning S. R (2013) **Improving the Accuracy and Efficiency of Identity-by-Descent Detection in Population Data** *Genetics* **194**:459–471 [Google Scholar](#)

45. Rosvall M., Axelsson D., Bergstrom C. T. (2009) **The Map Equation** *European Physical Journal Special Topics* **178**:13–23 [Google Scholar](#)
46. Csardi G., Nepusz T (2006) **The Igraph Software Package for Complex Network Research** *Inter-Journal, complex systems* :1–9 [Google Scholar](#)
47. Thompson E. A (2013) **Identity by Descent: Variation in Meiosis, Across Genomes, and in Populations** *Genetics* **194**:301–326 [Google Scholar](#)
48. Browning S. R., Browning B. L (2012) **Identity by Descent Between Distant Relatives: Detection and Applications** *Annual Review of Genetics* **46**:617–633 [Google Scholar](#)
49. Kelleher J., et al. (2019) **Inferring Whole-Genome Histories in Large Population Datasets** *Nature Genetics* **51**:1330–1338 [Google Scholar](#)
50. Speidel L., Forest M., Shi S., Myers S. R (2019) **A Method for Genome-Wide Genealogy Estimation for Thousands of Samples** *Nature Genetics* **51**:1321–1329 [Google Scholar](#)
51. Harris D. N., et al. (2023) **Evolutionary History of Modern Samoans** *Proceedings of the National Academy of Sciences* **117**:9458–9465 [Google Scholar](#)
52. Miotto O., et al. (2013) **Multiple Populations of Artemisinin-Resistant Plasmodium Falciparum in Cambodia** *Nature Genetics* **45**:648–655 [Google Scholar](#)
53. Koenderink J. B., Kavishe R. A., Rijpma S. R., Russel F. G. M (2023) **The ABCs of Multidrug Resistance in Malaria** *Trends in Parasitology* **26**:440–446 [Google Scholar](#)
54. Martin R. E., Kirk K (2004) **The Malaria Parasite's Chloroquine Resistance Transporter Is a Member of the Drug/Metabolite Transporter Superfamily** *Molecular Biology and Evolution* **21**:1938–1949 [Google Scholar](#)
55. Brooks D. R., et al. (1994) **Sequence Variation of the Hydroxymethyldihydropterin Pyrophosphokinase: Dihydropteroate Synthase Gene in Lines of the Human Malaria Parasite, Plasmodium Falciparum, with Differing Resistance to Sulfadoxine** *European journal of biochemistry* **224**:397–405 [Google Scholar](#)
56. Early A. M., et al. (2022) **Declines in Prevalence Alter the Optimal Level of Sexual Investment for the Malaria Parasite Plasmodium Falciparum** *Proceedings of the National Academy of Sciences* **119**:e2122165119 [Google Scholar](#)
57. Miotto O., et al. (2015) **Genetic Architecture of Artemisinin-Resistant Plasmodium Falciparum** *Nature Genetics* **47**:226–234 [Google Scholar](#)
58. Ishigohoka J., Liedvogel M. (2024) **High-Recombining Genomic Regions Affect Demography Inference** *bioRxiv* [Google Scholar](#)
59. Mehra S., Neafsey D. E., White M., Taylor A. R. (2024) **Systematic Bias in Malaria Parasite Relatedness Estimation** *G3* [Google Scholar](#)
60. MalariaGEN, et al. (2021) **An Open Dataset of Plasmodium Falciparum Genome Variation in 7,000 Worldwide Samples** *Wellcome Open Research* **6**:42 [Google Scholar](#)
61. Amir A., Cheong F. W., de Silva J. R., Liew J. W. K., Lau Y. L (2024) **Plasmodium Knowlesi Malaria: Current Research Perspectives** *Infection and Drug Resistance* **11**:1145–1155 [Google](#)

Scholar

62. Bright A. T., et al. (2014) **A High Resolution Case Study of a Patient with Recurrent Plasmodium Vivax Infections Shows That Relapses Were Caused by Meiotic Siblings** *PLoS neglected tropical diseases* **8**:e2882 [Google Scholar](#)
63. Ibrahim A., et al. (2023) **Population-Based Genomic Study of Plasmodium Vivax Malaria in Seven Brazilian States and across South America** *The Lancet Regional Health - Americas* **18**:100420 [Google Scholar](#)
64. Escalante A. A., Cornejo O. E., Rojas A., Udhayakumar V., Lal A. A (2004) **Assessing the Effect of Natural Selection in Malaria Parasites** *Trends in Parasitology* **20**:388–395 [Google Scholar](#)
65. Loy D. E., et al. (2024) **Out of Africa: Origins and Evolution of the Human Malaria Parasites Plasmodium Falciparum and Plasmodium Vivax** *International journal for parasitology* **47**:87–97 [Google Scholar](#)
66. Lee K.-S., et al. (2024) **Plasmodium Knowlesi: Reservoir Hosts and Tracking the Emergence in Humans and Macaques** *PLoS Pathogens* **7**:e1002015 [Google Scholar](#)
67. Stapley J., Feulner P. G. D., Johnston S. E., Santure A. W., Smadja C. M (2017) **Variation in Recombination Frequency and Distribution across Eukaryotes: Patterns and Processes.** *Philosophical Transactions of the Royal Society of London. Series B Biological Sciences* **372**:20160455 [Google Scholar](#)
68. Sivakumar T., Hayashida K., Sugimoto C., Yokoyama N (2025) **Evolution and Genetic Diversity of Theileria** *Infection, Genetics and Evolution* **27**:250–263 <https://www.sciencedirect.com/science/article/pii/S1567134814002421> | [Google Scholar](#)
69. Kent C. F., Minaei S., Harpur B. A., Zayed A (2025) **Recombination Is Associated with the Evolution of Genome Structure and Worker Behavior in Honey Bees** *Proceedings of the National Academy of Sciences* **109**:18012–18017 <https://www.pnas.org/doi/10.1073/pnas> | [Google Scholar](#)
70. Leroy T., et al. (2025) **Inferring Long-Term and Short-Term Determinants of Genetic Diversity in Honey Bees: Beekeeping Impact and Conservation Strategies** *Molecular Biology and Evolution* **41**:msae249 <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC11653568/> | [PubMed](#) | [Google Scholar](#)
71. Peter J., et al. (2025) **Genome Evolution across 1,011 Saccharomyces Cerevisiae Isolates** *Nature* **556**:339–344 <https://www.nature.com/articles/s41586-018-0030-5> | [Google Scholar](#)
72. Barton A. B., Pekosz M. R., Kurvathi R. S., Kaback D. B (2025) **Meiotic Recombination at the Ends of Chromosomes in Saccharomyces Cerevisiae** *Genetics* **179**:1221–1235 <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC2475728/> | [PubMed](#) | [Google Scholar](#)
73. Anderson T. J. C., Paul R. E. L., Donnelly C. A., Day K. P (2023) **Do Malaria Parasites Mate Non-Randomly in the Mosquito Midgut?** *Genetics Research* **75**:285–296 [Google Scholar](#)
74. Smukowski C. S., Noor M. A. F (2023) **Recombination Rate Variation in Closely Related Species** *Heredity* **107**:496–508 [Google Scholar](#)
75. Conway D. J., et al. (1999) **High Recombination Rate in Natural Populations of Plasmodium Falciparum** *Proceedings of the National Academy of Sciences of the United States of America*

96:4506–4511 [Google Scholar](#)

76. Browning B. L., Zhou Y., Browning S. R. (2018) **A One-Penny Imputed Genome from Next-Generation Reference Panels** *The American Journal of Human Genetics* **103**:338–348 [Google Scholar](#)

Author information

Bing Guo

Center for Vaccine Development and Global Health, University of Maryland School of Medicine, Baltimore, United States, Institute for Genome Sciences, University of Maryland School of Medicine, Baltimore, United States
ORCID iD: [0000-0002-3998-5981](#)

Shannon Takala-Harrison*

Center for Vaccine Development and Global Health, University of Maryland School of Medicine, Baltimore, United States
ORCID iD: [0000-0003-4674-8500](#)

For correspondence: stakala@som.umaryland.edu

*These authors jointly supervised this work

Timothy D O'Connor*

Institute for Genome Sciences, University of Maryland School of Medicine, Baltimore, United States
ORCID iD: [0000-0002-0276-1896](#)

For correspondence: timothydoconnor@gmail.com

*These authors jointly supervised this work

Editors

Reviewing Editor

Bavesh Kana

University of the Witwatersrand, Johannesburg, South Africa

Senior Editor

Bavesh Kana

University of the Witwatersrand, Johannesburg, South Africa

Reviewer #1 (Public review):

Summary:

Authors benchmarked five IBD detection methods (hmmIBD, isoRelate, hap-IBD, phasedIBD, and Refined IBD) in *Plasmodium falciparum* using simulated and empirical data. *Plasmodium falciparum* has a mutation rate similar to that of humans but a much higher recombination rate and lower SNP density. Thus, the authors evaluated how recombination rate and marker density affect IBD segment detection. Next, they performed parameter optimization for

Plasmodium falciparum and benchmarked the robustness of downstream analyses (selection detection and N_e inference) using IBD segments detected by each method. They also tracked the computational efficiency of these methods. The authors' work is valuable for the tested species, and the analyses presented support their claim that users should be cautious when calling IBD in contexts of low SNP density and high recombination rate.

Strengths:

The study design is convincing and well-structured. The authors chose to use *P. falciparum*, which presents an interesting case due to its high recombination rate and a mutation rate similar to that of humans. The authors note that despite the widespread use of IBD for genomic surveillance, comprehensive evaluation of these methods in high-recombination, low-marker-density contexts has been lacking. Furthermore, they also examined the performance of IBD detection methods developed specifically for *P. falciparum*, and evaluated it with phased data which broadened the applicability of the work.

Weaknesses:

The authors thoughtfully addressed our prior concerns by 1) expanding the simulations; 2) updating figures and methods for clarity; and 3) more clearly framing the broader utility of their benchmarking effort. These updates strengthen the manuscript and make the relevance of their findings beyond *Plasmodium falciparum* more apparent.

More specifically:

The authors added three full replicates per simulation scenario and updated figures to reflect uncertainty at relevant levels, which addresses earlier concerns about reproducibility. The limited number of replicates is due to computational intensity. In the future, broader generalizability and deeper exploration of variance in benchmarking accuracy across parameter space would further strengthen the conclusions/generalizability. The author's also emphasized that, while the study is centered on *Plasmodium falciparum*, the benchmarking framework, not the parameters, are broadly applicable to other sexually recombining species. Lastly, they extensively updated multiple figures to include simulation models, results from simulation replicates, and improved the figures from the previous version of the manuscript.

<https://doi.org/10.7554/eLife.101924.2.sa2>

Reviewer #2 (Public review):

Summary:

Guo et al. benchmarked and optimized methods for detecting Identity-By-Descent (IBD) segments in *Plasmodium falciparum* (Pf) genomes, which are characterized by high recombination rates and low marker density. Their goal was to address the limitations of existing IBD detection tools, which were primarily developed for human genomes and do not perform well in the genomic context of highly recombinant genomes. They first analysed various existing IBD callers, such as hmmIBD, isoRelate, hap-IBD, phased-IBD, and refinedIBD. They focused on the impact of recombination on the accuracy, which was calculated based on two metrics, the false negative rate and the false positive rate. The results suggest that high recombination rates significantly reduce marker density, leading to higher false negative rates for short IBD segments. This effect compromises the reliability of IBD-based downstream analyses, such as effective population size (N_e) estimation. They showed that the best tool for IBD detection in Pf is hmmIBD, because it has relatively low FN/FP error rates and is less biased for relatedness estimates. However, this method is less computationally efficient.

Their suggestion is to optimize human-oriented IBD methods and use hmmIBD only for the estimation of N_e .

Strengths:

Although I am not an expert on *Plasmodium falciparum* genetics, I believe the authors have developed a valuable benchmarking framework tailored to the unique genomic characteristics of this species. Their framework enables a thorough evaluation of various IBD detection tools for non-human data, such as high recombination rates and low marker density, addressing a key gap in the field.

This study provides a comparison of multiple IBD detection methods, including probabilistic approaches (hmmIBD, isoRelate) and IBS-based methods (hap-IBD, Refined IBD, phased IBD). This comprehensive analysis offers researchers valuable guidance on the strengths and limitations of each tool, allowing them to make informed choices based on specific use cases. I think this is important beyond the study of Pf.

The authors highlight how optimized IBD detection can help identify signals of positive selection, infer effective population size (N_e), and uncover population structure.

They demonstrate the critical importance of tailoring analytical tools to suit the unique characteristics of a species. Moreover, the authors provide practical recommendations, such as employing hmmIBD for quality-sensitive analyses and fine-tuning parameters for tools originally designed for non-P. falciparum datasets before applying them to malaria research.

Overall, this study represents a meaningful contribution to both computational biology and malaria genomics, with its findings and recommendations likely to have an impact on the field.

Weaknesses:

One weakness of the study is the lack of emphasis on the broader importance of studying *Plasmodium falciparum* as a critical malaria-causing organism. Malaria remains a significant global health challenge, causing hundreds of thousands of deaths annually.

While the study provides a thorough technical evaluation of IBD detection methods and their application to Pf, it does not adequately connect these findings to the broader implications for malaria research and control efforts. Additionally, the discussion on malaria and its global impact could have framed the study in a more accessible and compelling way, making the importance of these technical advances clearer to a broader audience, including researchers and policymakers in the fight against malaria. In the revised version, the authors have placed greater emphasis on this aspect, while still maintaining the methodological focus of the paper.

<https://doi.org/10.7554/eLife.101924.2.sa1>

Author response:

The following is the authors' response to the original reviews

Public Reviews:

Reviewer #1 (Public review):

Summary:

Authors benchmarked 5 IBD detection methods (hmmIBD, isoRelate, hap-IBD, phasedIBD, and Refined IBD) in Plasmodium falciparum using simulated and empirical data. Plasmodium falciparum has a mutation rate similar to humans but a much higher recombination rate and lower SNP density. Thus, the authors evaluated how recombination rate and marker density affect IBD segment detection. Next, they performed parameter optimization for Plasmodium falciparum and benchmarked the robustness of downstream analyses (selection detection and Ne inference) using IBD detected by each of the methods. They also tracked the computational efficiency of these methods. The authors work is valuable for the tested species and the analyses presented appear to support their claim that users should be cautious calling IBD when SNP density is low and recombination rate is high.

Strengths:

The study design was solid. The authors set up their reasoning for using P. falciparum very well. The high recombination rate and similar mutation rate to humans is indeed an interesting case. Further, they chose methods that were developed explicitly for each species. This was a strength of the work, as well as incorporating both simulated and empirical data to support their goal that IBD detection should be benchmarked in P. falciparum.

Weaknesses:

The scope of the optimization and application of results from the work are narrow, in that everything is finetuned for Plasmodium. Some of the results were not entirely unexpected for users of any of the tested software that was developed for humans. For example, it is known that Refined IBD is not going to do well with the combination of short IBD segments and low SNP density. Lastly, it appears the authors only did one largescale simulation (there are no reported SDs).

We thank the reviewer for highlighting the strengths and weaknesses of the study.

First, we would like to highlight that: (1) while we use *Plasmodium* as a model to investigate the impact of high recombination and low marker density on IBD detection and downstream analyses, our IBD benchmarking framework and strategies are widely applicable to IBD methods development for many sexually recombining species including both *Plasmodium* and non-*Plasmodium* species. (2) Although some results are not completely unexpected, such as the impact of low marker density on IBD detection, IBD-based methods have been increasingly used in malaria genomic surveillance research without comprehensive benchmarking for malaria parasites despite the high recombination rate. Due to the lack of benchmarking, researchers use a variety of different IBD callers for malaria research including those that are only benchmarked in human genomes, such as refined-ibd. Our work not only confirmed that low marker density (related to high recombination rate) can affect the accuracy of IBD detection, but also demonstrated the importance of proper parameter optimization and tool prioritization for specific downstream analyses in malaria research. We believe our work significantly contributes to the robustness of IBD segment detection and the enhancement of IBD-based malaria genomic surveillance.

Second, we agree that there is a lack of clarity regarding simulation replicates and the uncertainty of reported estimates. We have made the following improvements, including (1) running $n = 3$ full sets of simulations for each analysis purpose, which is in addition to the large sample sizes and chromosomal-level replications already presented in our initial submission, and (2) updating data and figures to reflect the uncertainty at relevant levels (segment level, genome-pair level or simulation set level).

Reviewer #2 (Public review):

Summary:

Guo et al. benchmarked and optimized methods for detecting Identity-By-Descent (IBD) segments in *Plasmodium falciparum* (Pf) genomes, which are characterized by high recombination rates and low marker density. Their goal was to address the limitations of existing IBD detection tools, which were primarily developed for human genomes and do not perform well in the genomic context of highly recombinant genomes. They first analysed various existing IBD callers, such as *hmmIBD*, *isoRelate*, *hap-IBD*, *phased-IBD*, *refinedIBD*. They focused on the impact of recombination on the accuracy, which was calculated based on two metrics, the false negative rate and the false positive rate. The results suggest that high recombination rates significantly reduce marker density, leading to higher false negative rates for short IBD segments. This effect compromises the reliability of IBD-based downstream analyses, such as effective population size (N_e) estimation. They showed that the best tool for IBD detection in Pf is *hmmIBD*, because it has relatively low FN/FP error rates and is less biased for relatedness estimates. However, this method is less computationally efficient. Their suggestion is to optimize human-oriented IBD methods and use *hmmIBD* only for the estimation of N_e .

Strengths:

Although I am not an expert on *Plasmodium falciparum* genetics, I believe the authors have developed a valuable benchmarking framework tailored to the unique genomic characteristics of this species. Their framework enables a thorough evaluation of various IBD detection tools for non-human data, such as high recombination rates and low marker density, addressing a key gap in the field. This study provides a

comparison of multiple IBD detection methods, including probabilistic approaches (*hmmIBD*, *isoRelate*) and IBS-based methods (*hap-IBD*, *Refined IBD*, *phased IBD*). This comprehensive analysis offers researchers valuable guidance on the strengths and limitations of each tool, allowing them to make informed choices based on specific use cases. I think this is important beyond the study of Pf. The authors highlight how optimized IBD detection can help identify signals of positive selection, infer effective population size (N_e), and uncover population structure. They demonstrate the critical importance of tailoring analytical tools to suit the unique characteristics of a species. Moreover, the authors provide practical recommendations, such as employing *hmmIBD* for quality-sensitive analyses and fine-tuning parameters for tools originally designed for non-P. *falciparum* datasets before applying them to malaria research.

Overall, this study represents a meaningful contribution to both computational biology and malaria genomics, with its findings and recommendations likely to have an impact on the field.

Weaknesses:

One weakness of the study is the lack of emphasis on the broader importance of studying *Plasmodium falciparum* as a critical malaria-causing organism. Malaria remains a significant global health challenge, causing hundreds of thousands of deaths annually. The authors could have introduced better the topic, even though I understand this is a methodological paper. While the study provides a thorough technical evaluation of IBD detection methods and their application to Pf, it does not adequately connect these findings to the broader implications for malaria research and control efforts. Additionally, the discussion on malaria and its global impact could have framed the study in a more accessible and compelling way, making the importance of these

technical advances clearer to a broader audience, including researchers and policymakers in the fight against malaria.

We thank the reviewer for highlighting the need to better contextualize the work and emphasize its relevance to malaria control and elimination efforts. We have edited the introduction and discussion sections to highlight the importance of studying *Plasmodium* as malaria-causing organisms and why IBD-based analysis is important to malaria researchers and policymakers. We believe the changes will better emphasize the public health relevance of the work and improve clarity for a general audience.

We would like to clarify that we are not recommending that researchers “optimize human-oriented IBD methods and use hmmIBD only for the estimation of N_e .” We recommended hmmIBD for N_e analysis; however, hmmIBD can be utilized for other applications, including population structure and selection detection. Thus, we generally recommend using hmmIBD for *Plasmodium* when phased genotypes are available. To avoid potential misunderstandings, we have revised relevant sentences in the abstract, introduction, and discussion. One reason to consider human-oriented IBD detection methods in *Plasmodium* research is that hmmIBD currently has limitations in handling large genomic datasets. Our ongoing research focuses on improving hmmIBD to reduce its computational runtime, making it scalable for large *Plasmodium* wholegenome sequence datasets.

Recommendations for the authors:

Reviewer #1:

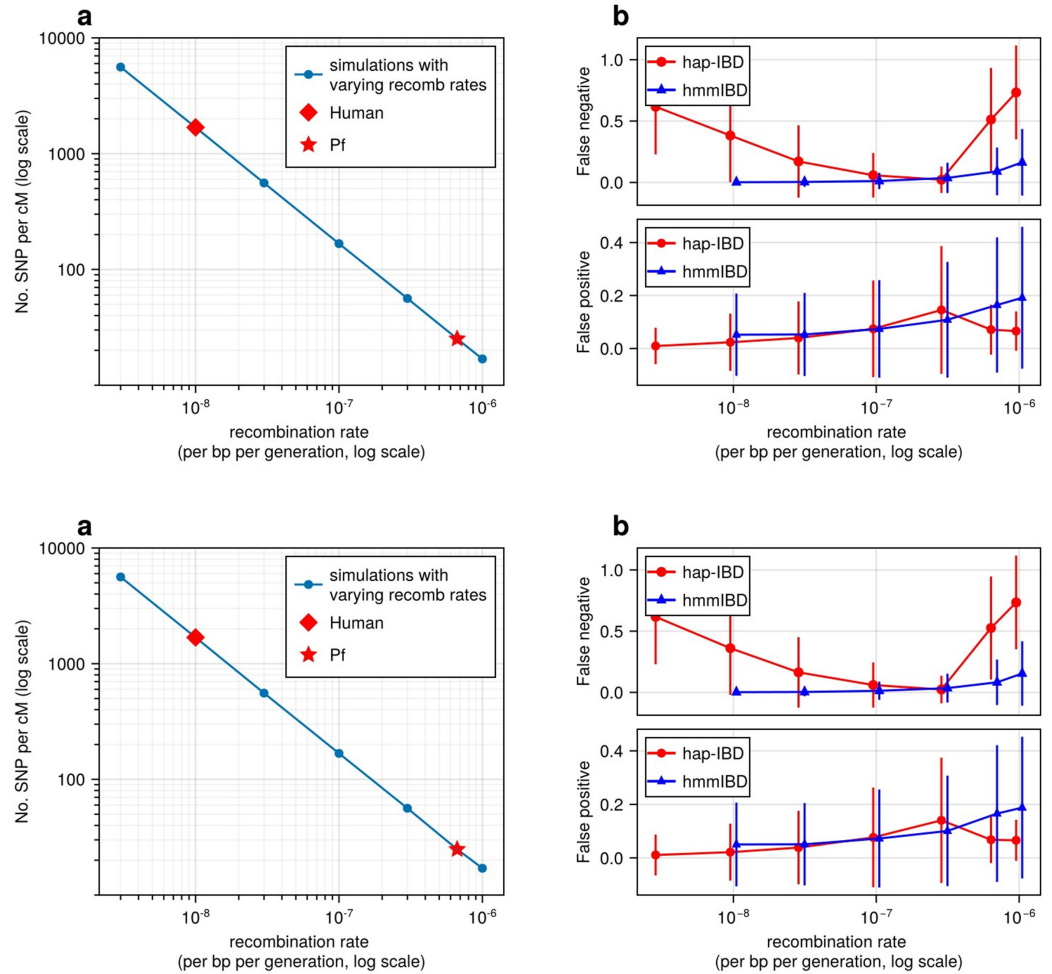
(1) Additional experiments

(i) More simulation replicates would be valuable here. The way that results are presented, it appears as though there are no replicates. Apologies if I am incorrect, but when looking through the authors code the --num_reps defaults to one simulation and there are no SDs reported for any figure. Perhaps the authors are bypass replicates by taking a random sample of lineages? Some clarification here would be great.

We agree with the reviewer’s constructive suggestions. We have increased the number of simulation sets to ($n = 3$) in addition to the existing replicates at the chromosomal level. We did not use a larger n for full sets of simulation replicates for two reasons: (1) full replication is quite computationally intensive ($n=3$ simulation sets already require a week to run on our computer cluster with hundreds of CPU cores). (2), the results from different simulation sets are highly consistent with each other, likely due to our large sample size ($n= 1000$ haploid genomes for each parameter combination). The consistency across simulation sets can be exemplified by the following figures (Author response image 1 and 2) based on simulation sets different from Figures and Supplementary Figures included in the manuscript.

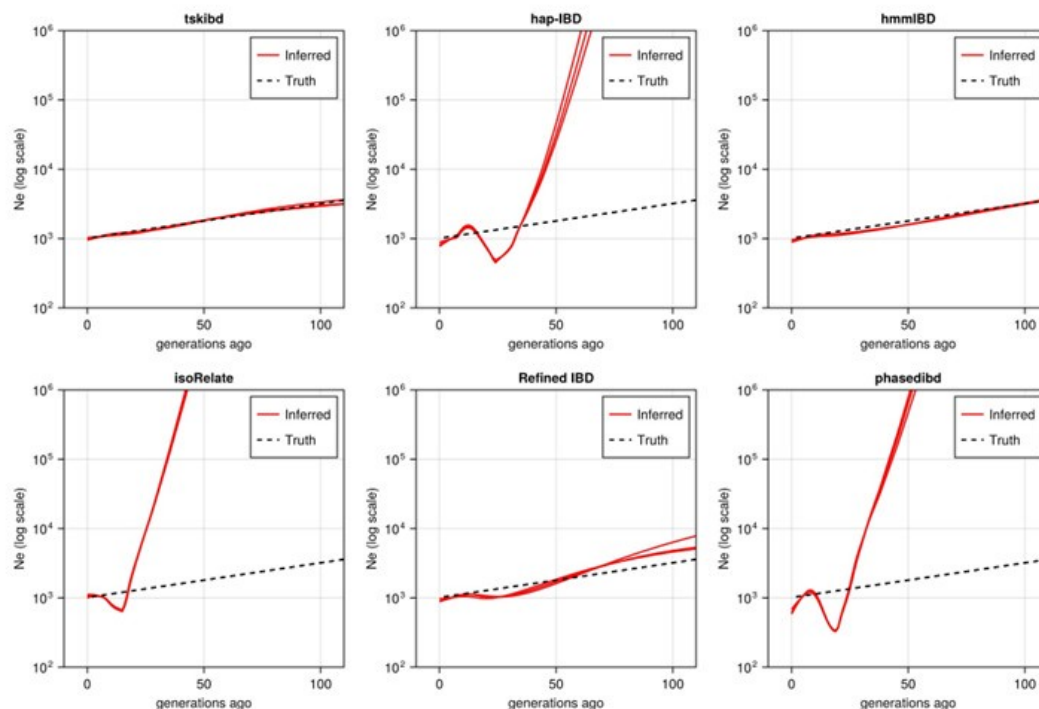
Author response image 1.

Additional simulation sets repeating experiments shown in Fig 2.



Author response image 2.

Post-optimization N_e estimates based on three independent simulation sets (Fig 5 shows data simulation set 1).



In our updated figures, we address the uncertainty of measurements as follows:

(1) For IBD accuracy based on overlapping IBD segments, we present the mean $\pm$ standard deviation (SD) at the segment level (IBD segment false positives and false negatives for each length bin) or genome-pair level (IBD error rates at the genome-wide level). Figures in the revised manuscript show results from one of the three simulation set replicates. The SD of IBD segment accuracy is included in all relevant figures. In the S2 Data file, we chose not to show SDs to avoid text overcrowding in the heatmaps; however, a detailed version, including SD plotting on the heatmap and across three simulation set replicates, is available on our GitHub repository at https://github.com/bgao068/bmibdcaller_simulations/tree/main/simulations/ext_data.

(2) For IBD-based genetic relatedness, the uncertainty is depicted in scatterplots.

(3) For IBD-based selection signal scans, we provide the mean $\pm$ SD of the number of true selection signals and false selection signals. The SD is calculated at the simulation set level (n=3).

(4) For IBD network community detection, the mean $\pm$ SD of the adjusted Rand index is reported at the simulation set level (n=3). A representative simulation set is randomly chosen for visualization purposes.

(5) For IBD-based Ne estimates, each simulation set provides confidence intervals via bootstrapping. We found Ne estimates across n=3 simulation sets to be highly consistent and decided to display Ne from one of the simulation sets.

(6) For the measurement of computational efficiency and memory usage, the mean $\pm$ SD was calculated across chromosomes from the same simulation sets.

We have included a paragraph titled "Replications and Uncertainty of Measures" in the methods section to clarify simulation replications. Additionally, a table of simulation

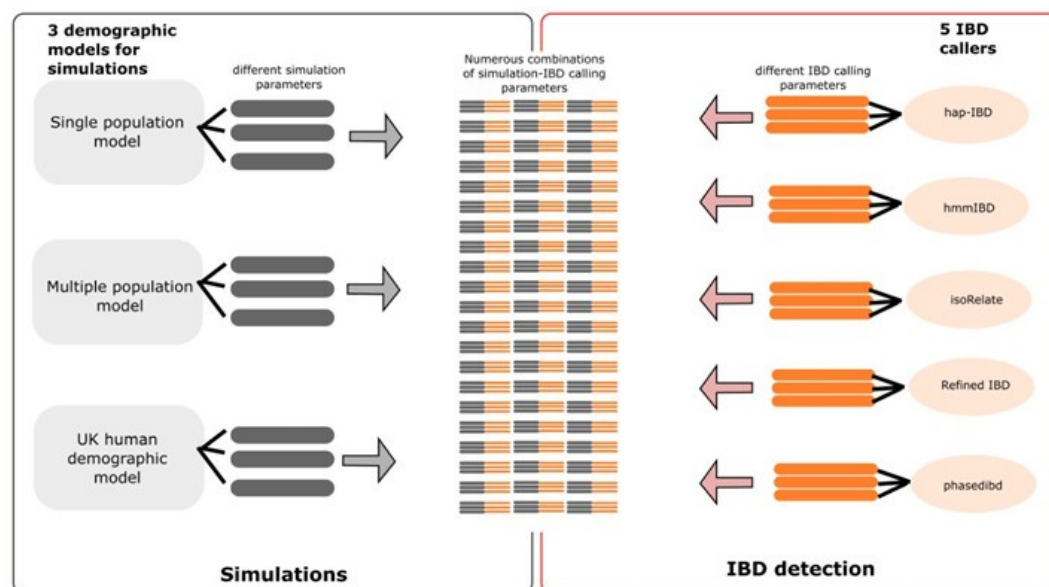
replicates is provided in the new S1 Data file under the sheet named “02_simulation_replicates.”

(ii) I might also recommend a table or illustrative figure with all the simulation parameters for the readers rather than them having to go to and through a previous paper to get a sense of the tested parameters.

We have now generated tables containing full lists of simulation/IBD calling parameters. We have organized the tables into two sections: simulation parameters and IBD calling parameters. For the simulations, we are using three demographic models: the single-population (SP) model, the multiple-population (MP) model, and the human population demography in the UK (UK) model, each with different sets of parameters. Parameters and their values are listed separately for each demographic model (SP, MP and UK). For the IBD calling, we have five different IBD callers, each with different parameters. We have provided lists of the parameters and their values separately for each caller. In total, there are 15 different combinations of 3 demographic models in simulation and five callers in IBD detection (Author response image 3). We provide a table for each of the 15 combinations. We also provide a single large table by concatenating all 15 tables. In the combined table, demographic model-specific or IBD caller-specific parameters are displayed in their own columns, with NA values (empty cells) appearing in rows where these parameters are not applied (see S2 Data file).

Author response image 3.

Schematic of combined parameters from simulations and IBD detection (also included in the S2 Data file)



(2) Recommendations for improving the writing and presentation

Overall, the writing was great, especially the introduction.

Three thoughts:

(i) It would be great if the authors included a few sentences with guidance on the approach one would take if their organism was not human or *P. falciparum*.

We have updated our discussion with the following statement: “Beyond Plasmodium parasites, there are many other high-recombining organisms such as Apicomplexan species like Theileria, insects like Apis mellifera (honeybee), and fungi like Saccharomyces cerevisiae (Baker's yeast). For these species, our optimized parameters may not be directly applicable, but the benchmarking framework established in this study can be utilized to prioritize and optimize IBD detection methods in a context-specific manner.”

(ii) I think there was a lot of confusion about the simulations as they were presented between the co-reviewer and I. Clarification on whether there were replicates and how sampling of lineages occurred would be helpful for a reader.

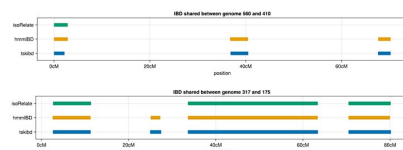
We have added a paragraph with heading “Replications and uncertainty of measures” under the method section to clarify simulation replicates. Please also refer to our response above for more details (Reviewer #1 (1) Additional experiments).

(iii) Maybe we missed it, but could the authors add a sentence or two about why isoRelate performed so poorly (e.g. lines 206-207) considering it was developed for Plasmodium? This result seems important.

IsoRelate assumes non-phased genotypes as input; therefore, even if phased genotypes are provided, the HMM model used in isoRelate (distinct from the hmmIBD model) may not utilize them. Below, we present examples of IBD segments between true sets and inferred sets from both isoRelate and hmmIBD, where many small IBD segments identified by tskibd (ground truth) and hmmIBD (inferred) are not detected by isoRelate (inferred), although isoRelate still captures very long IBD segments. These patterns are also illustrated in Fig. 3 and S3 Fig. We acknowledge that isoRelate may outperform other methods in the context of unphased genotypes. However, we chose not to benchmark IBD calling methods using unphased genotypes in simulations, as the results may be significantly influenced by the quality of genotype phasing for all other IBD detection methods. The characterization of deconvolution methods is beyond the scope of this paper. We have added a paragraph in the discussion to reflect the above explanation.

Author response image 4.

Example IBD segments inferred by *isoRelate* and *hmmIBD* compared to true IBD segments calculated by *tskibd*.



(3) Minor corrections to the text and figures

Lines 105-110 feel like introduction because the authors are defining IBD and goals of work

We have shortened these sentences and retained only relevant information for transition purposes.

Line 121-122 The definition of false positive is incorrect, it appears to be the exact text from false negative

We apologize for the typo and have corrected the definition, so that it is consistent with that in the methods section.

Lines 177-180 feels more like discussion than results

We have removed this sentence for brevity.

Figure 1:

Remove plot titles from the figure

Write out number in a

The legend in b overlaps the data so moving that inset to the right would be helpful

We have removed the titles from Figure 1. In Figure 1a, we have changed the format of the y-axis tick labels from scientific notation to integers. In Figure 1b, we have adjusted the size and location of the legend so that it does not overlap with the data points.

Figure 2-3 & S4-5:

It was hard to tell the difference between [3-4) and [10-18) because the colors and shapes are similar. It might be worth using a different color or shape for one of them?

We have changed the color for the [10-18) group so that the two groups are easier to distinguish.

Figure 3 & S3-5:

Biggest suggestion is that when an axis is logged it should not only be mentioned in the caption but also should be shown in the figure as well.

We have updated all relevant figures so that the log scale is noted in the figure captions (legends) as well as in the figures (in the x and/or y axis labels).

Supplementary Figure S2

(i) It would be nice to either combine it with the main text Figure 1 (I don't believe it would be overwhelming) or add in the other two methods for comparison

We have now plotted data for all five IBD callers in S1 Fig for better comparison.

(ii) the legend overlaps the data so relocating it to the top or bottom would be helpful

We have moved the legend to the bottom of the figure to avoid overlap with the data.

Reviewer #2:

I don't have any major comments on the paper. It is well-written, although perhaps a bit long and repetitive in some sections. Make sure not to repeat the same concepts too many times.

We have consolidated and removed several paragraphs to reduce repetition of the same concepts.

*I am not a methodological developer, but it seems you have addressed several challenges regarding IBD detection in *P. falciparum*. You have also acknowledged the study's caveats, which I agree with.*

Thank you for the positive comments.

Minor comments:

-In my opinion, the paper would benefit from including the workflow figure in the main text rather than keeping it in the supplementary materials. This would make it more accessible and useful for readers.

We have moved the original S1 Fig to be Fig 1 in the main text.

-Some of the figures (e.g. Fig. 2, 4) should be larger for better clarity and interpretation.

We have updated Fig 2 and Fig 4 (now labeled as Figure 3 and 5) to make them larger for improved clarity and interpretation.

*-While the focus on *P. falciparum* is understandable, it would have been valuable to include examples of other species and discuss the broader implications of the findings for a broader field.*

We have updated the third-to-last paragraph to discuss implications for other species, such as Apicomplexan species like *Theileria*, insects like *Apis mellifera* (honeybee), and fungi like *Saccharomyces cerevisiae* (Baker's Yeast). We acknowledge that optimal parameters and tool choices may vary among species due to differences in demographic history and evolutionary parameters. However, we emphasize that the methods outlined are adaptable for prioritizing and optimizing IBD detection methods in a context-specific manner across different species.

-Figure 6 is somewhat confusing and could use clearer labeling or additional explanation to improve comprehension.

We have updated the labels and titles in the figure to improve clarity. We also edited the figure caption for better clarity.

*-Although *hmmIBD* outperformed other tools in accuracy, its computational inefficiency due to single-threaded execution poses a significant challenge for scaling to large datasets. The trade-off between accuracy and computational cost could be discussed in more detail.*

We have added a paragraph in the discussion section to highlight the trade-off between accuracy and computation cost. We noted that we are developing an adapted tool to enhance the *hmmIBD* model and significantly reduce the runtime via parallelizing the IBD inference process.

<https://doi.org/10.7554/eLife.101924.2.sa0>