

Tests for segregation distortion in higher ploidy F1 populations

David Gerard ^{1,*} Guilherme Bovi Ambrosano ² Guilherme da Silva Pereira ³ Antonio Augusto Franco Garcia ²

¹Department of Mathematics and Statistics, American University, Washington, DC 20016-8002, United States

²Department of Genetics, Luiz de Queiroz College of Agriculture, University of São Paulo, Piracicaba 13418-900, Brazil

³Department of Agronomy, Federal University of Viçosa, Viçosa, Minas Gerais 36570-900, Brazil

*Corresponding author: Department of Mathematics and Statistics, American University, Washington, DC 20016-8002, United States. Email: dgerard@american.edu.

F1 populations are widely used in genetic mapping studies in agriculture, where known pedigrees enable rigorous quality control measures such as segregation distortion testing. However, conventional tests for segregation distortion are inadequate for polyploids, as they fail to account for double reduction, preferential pairing, and genotype uncertainty, leading to inflated type I error rates. Prior work developed a statistical framework to address these issues in tetraploids. Here, we extend these methods to higher even ploidy levels and introduce additional strategies to mitigate the influence of outliers. Through extensive simulations, we demonstrate that our tests maintain appropriate type I error control while retaining power to detect true segregation distortion. We further validate our approach using empirical data from a hexaploid mapping population. Our methods are implemented in the *segtest* R package, available on the Comprehensive R Archive Network (<https://doi.org/10.32614/CRAN.package.segtest>).

Keywords: double reduction; genotype likelihoods; genotype uncertainty; likelihood ratio tests; partial preferential pairing

Introduction

Polyploidy (whole genome duplication) is a common feature of many plant genomes. In addition to being a major driver of evolution in plants (Adams and Wendel 2005; Soltis et al. 2009), many important crops are polyploids (Udall and Wendel 2006). Notable examples include staple and specialty crops such as wheat (Matsuoka 2011), potatoes (Rosyara et al. 2016; Endelman et al. 2018), sweetpotatoes (Shirasawa et al. 2017; Mollinari et al. 2020), and blueberries (Ferrão et al. 2021). Because these crops are central to global food security and agricultural economies, methods for studying and analyzing polyploid genomic data are critical for understanding their biology, improving breeding programs, and supporting future crop development.

Such genomic methods must contend with the fact that meiosis in polyploids is rather complicated (Stift et al. 2008; Voorrips and Maliepaard 2012; Soares et al. 2021). One complication, called “preferential pairing,” arises from hom(o)ologous chromosomal pairing patterns. After replication, hom(o)ologous chromosomes can form either bivalents (as in diploids) or multivalents. Under strict bivalent formation, if homologous chromosomes always pair and homoeologous chromosomes never pair, resulting in disomic inheritance, this is considered allopolyploidy (using the genetic, not taxonomic, definition Doyle and Sherman-Broyles 2017). Conversely, if all chromosomes are homologous and bivalent pairing is nonpreferential, resulting in polysomic inheritance, this is considered autopolyploidy. A continuum of partial preferential pairing is possible, giving rise to intermediate forms called segmental allopolyploidy, among other terms (Bourke et al. 2017). Although segregation patterns

at a locus vary with the degree of preferential pairing, many current genomic methods assume strictly disomic or polysomic models, with limited capacity to handle partial preferential pairing (Bourke et al. 2018; Mollinari et al. 2020).

Another complication of polyploid meiosis, called “double reduction,” is possible in the presence of multivalent formation. If an odd number of crossovers occur between the centromere and a locus, followed by the sister chromatid segments migrating to the same pole in both the first and second meiotic divisions, then the two sister chromatid segments may end up in the same gamete (Stift et al. 2010). A visual illustration of this for a simplex locus in a tetraploid is provided in [Supplementary Fig. S1](#) in File S1. Consequently, for example, a parent with a simplex marker can produce duplex gametes. More generally, segregation patterns differ for loci where double reduction is possible. A clear treatment of this phenomenon is presented in Huang et al. (2019). However, most methods assume the absence of double reduction (Bourke et al. 2018; Mollinari et al. 2020).

Polyploid data are also more susceptible to genotype uncertainty than diploid data. This is because there are more possible genotypes to distinguish in polyploids. Whereas diploid genotyping at a locus with alleles A and a must distinguish among AA, Aa, and aa, tetraploid genotyping must distinguish among AAAA, AAAa, AAaa, Aaaa, and aaaa. Differentiating among more categories is inherently more difficult, and this challenge is compounded by features of modern sequencing technologies (Baird et al. 2008; Elshire et al. 2011), such as allele bias and overdispersion (Gerard et al. 2018). Treating genotypes as known without error, when in fact they are error-prone estimates, can lead to biased inference. Consequently, researchers have been

developing uncertainty-aware methods for many analyses (Li 2011; Korneliusen et al. 2014; Gerard 2021a, 2021b).

Even in the face of these complications, the ubiquity of polyploid crops has driven the use of experimental polyploid populations in agricultural research for tasks such as linkage mapping, quantitative trait locus mapping, and genomic selection (Shirasawa et al. 2017; Bourke et al. 2018; de C Lara et al. 2019; Mollinari and Garcia 2019; da Silva Pereira et al. 2020; Gemenet et al. 2020; Mollinari et al. 2020; Amadeu et al. 2021; Ferrão et al. 2021; Lau et al. 2022). These sophisticated statistical methods require high-quality genotype data to achieve optimal performance. Consequently, genotype data are routinely subjected to rigorous quality control procedures (e.g. Bourke et al. 2015; Cappai et al. 2020; Mollinari et al. 2020; Batista et al. 2021). Given that experimental populations possess well-defined family pedigrees, a frequent check involves using a χ^2 -test to compare observed offspring genotype frequencies with those predicted by Mendelian segregation (Mendel 1866), as is done by the `mappoly` software (Mollinari et al. 2020). However, this method has notable limitations: it fails to incorporate features specific to polyploid meiosis, such as double reduction and partial preferential pairing as described above, and does not account for the increased genotype uncertainty associated with modern sequencing technologies.

Recent approaches to testing for segregation distortion by Bourke et al. (2018) and Gerard et al. (2025) have been specifically tailored to polyploid data. However, both methods exhibit certain limitations. Bourke et al. (2018), through the `checkF1()` function of their `polymapR` software, implement a χ^2 -test for all polysomic and disomic segregation patterns in polyploids. Additionally, their method permits a specified proportion of individuals to have “invalid” genotypes (e.g. offspring genotypes greater than or equal to 2 are deemed invalid when the parent genotypes are 0 and 1). Despite these features, the approach of Bourke et al. (2018) has several shortcomings: (i) it does not account for partial preferential pairing or double reduction, (ii) it addresses genotype uncertainty in an ad hoc manner, which results in elevated type I error rates (Gerard et al. 2025), and (iii) it conducts separate tests for invalid genotypes and segregation distortion, an approach difficult to generalize when accounting for genotype uncertainty. See Gerard et al. (2025) for a detailed description of the `checkF1()` function.

Gerard et al. (2025) addressed some of these limitations by developing tests that incorporate both double reduction and partial preferential pairing, while systematically accounting for genotype uncertainty through the use of genotype likelihoods. However, the methods proposed by Gerard et al. (2025) are restricted to tetraploids and do not accommodate a proportion of invalid genotypes. Allowing for a certain degree of “wobble room” in invalid genotypes can be important, as a SNP may be otherwise valid except for a few individuals with anomalous results. Discarding an entire SNP based solely on the presence of a few outliers could result in unnecessary data loss.

In this article, we extend the work of Gerard et al. (2025) to develop tests for segregation distortion applicable to higher (even) ploidies, accounting for partial preferential pairing, invalid genotypes, and genotype uncertainty. Our approach integrates invalid genotypes into a single test for segregation distortion using a mixture model of genotype frequencies. Genotype uncertainty is addressed in a principled manner through the implementation of likelihood ratio tests (LRTs), utilizing

genotype likelihoods. Our methods further account for the effects of double reduction through a novel result for their gamete frequencies at simplex loci, and we demonstrate that our methods are robust to moderate levels of double reduction at non-simplex loci.

Materials and methods

Models for genotype frequencies

In this section, we develop a null model for genotype frequencies in an F1 population, defining segregation distortion as deviations from these expected frequencies. The model includes a parameter β to account for double reduction at simplex loci, a parameter γ to account for partial preferential pairing (which is only an important modeling consideration at nonsimplex loci), and a parameter π that specifies the tolerated proportion of outliers (i.e. “invalid genotypes”).

At a biallelic locus with alleles A and a, a K-ploid individual's genotype is their number of a alleles. The genotype frequencies $\mathbf{q} = (q_0, q_1, \dots, q_K)$ of an F1 population, where q_k is the probability an offspring has genotype k, are given by the discrete linear convolution of parental gamete frequencies $\mathbf{p}_j = (p_{j0}, p_{j1}, \dots, p_{j,K/2})$, where p_{jk} is the probability parent j's gamete has genotype k:

$$q_k = \sum_{i=\max(0, k-K/2)}^{\min(k, K/2)} p_{1i} p_{2, k-i}. \quad (1)$$

Different models for meiosis correspond to different models for the p_{jk} 's.

For nonsimplex loci, we account for (partial) preferential pairing using the pairing configuration model of Gerard et al. (2018), which is graphically represented in Fig. 1. During meiosis, a parent's K chromosome copies form K/2 pairs. Let $\mathbf{m} = (m_0, m_1, m_2)$ be the pairing configuration, where m_j is the number of pairs containing j copies of a. Given a pairing configuration, the gamete frequencies are:

$$p_k = \binom{m_1}{k - m_2} \frac{1}{2^{m_1}}. \quad (2)$$

All possible values for the p_k 's from (2) for ploidies 2 through 12 are presented in Supplementary Table S1 in File S1. Since the pairing configuration is unknown, let γ_i denote the probability of pairing configuration $i \in \{1, \dots, b_{K\ell}\}$, where $b_{K\ell}$ is the number of possible configurations for a parent with ploidy K and genotype ℓ (Gerard et al. 2018):

$$b_{K\ell} = \begin{cases} \frac{K}{2} - \lceil \frac{\ell}{2} \rceil + 1 & \text{if } \ell \geq K/2, \\ \lfloor \frac{\ell}{2} \rfloor + 1 & \text{if } \ell < K/2. \end{cases} \quad (3)$$

Marginalizing over pairing configurations, the gamete frequencies become:

$$p_k = \sum_{i=1}^{b_{K\ell}} \gamma_i \binom{m_{i1}}{k - m_{i2}} \frac{1}{2^{m_{i1}}}. \quad (4)$$

Model (4) generalizes the gamete frequencies of allopolyploids (disomic inheritance) and bivalent pairing autopolyploids (polysomic inheritance) (Doyle and Egan 2010; Parisod et al. 2010). For allopolyploids, γ is a one-hot vector (only one pairing configuration is possible). The value of γ that corresponds to bivalent pairing

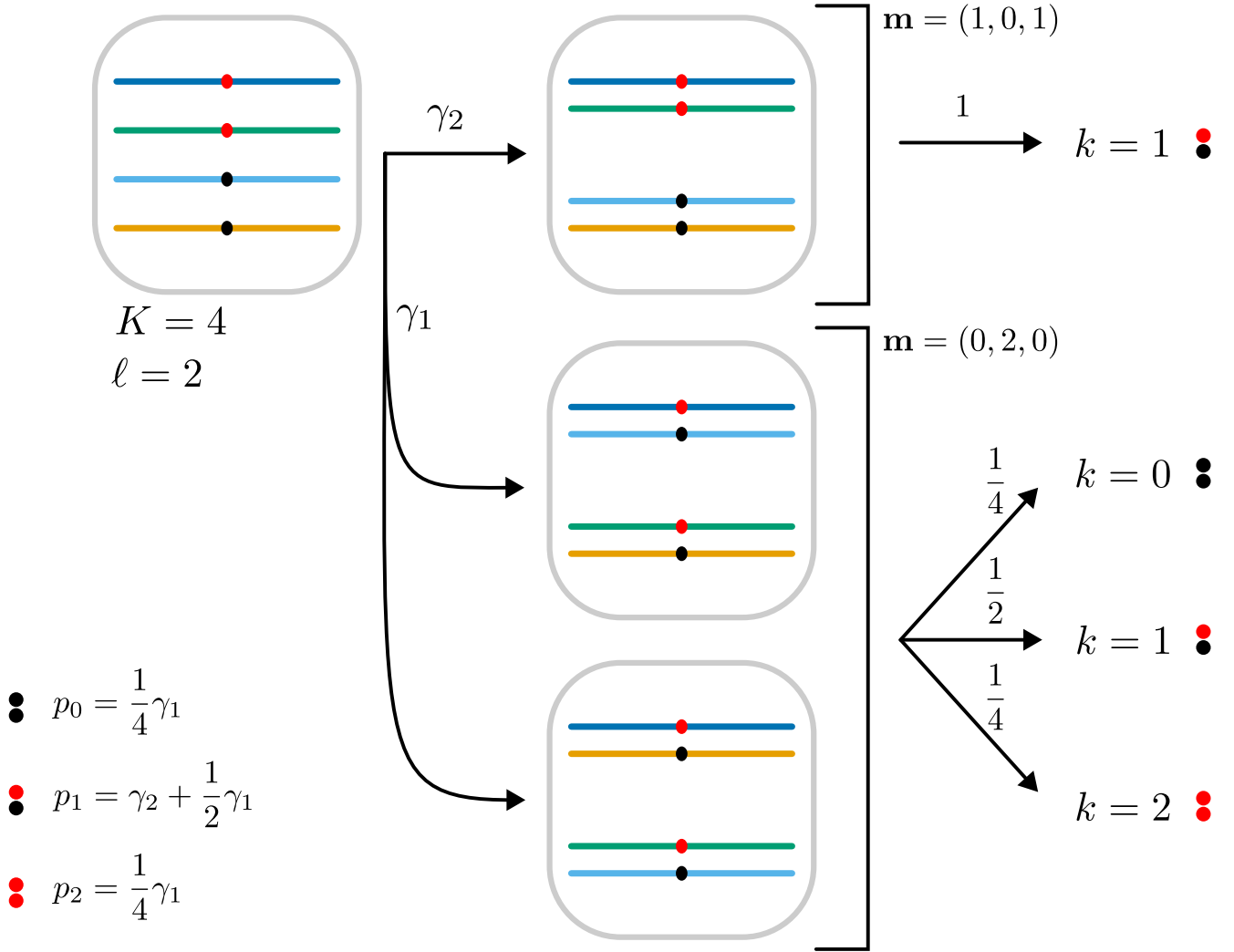


Fig. 1. A graphical illustration of the pairing configuration model described in *Models for genotype frequencies* section. A $K=4$ tetraploid parent with red and black alleles, and thus a dosage of $\ell=2$, is shown. Two pairing configurations are possible: $\mathbf{m} = (1, 0, 1)$, where the red alleles pair and the black alleles pair, and $\mathbf{m} = (0, 2, 0)$, where each pairing consists of one red and one black allele. The configuration $\mathbf{m} = (0, 2, 0)$ occurs with probability γ_1 , and $\mathbf{m} = (1, 0, 1)$ occurs with probability γ_2 . Configuration $\mathbf{m} = (1, 0, 1)$ produces only gametes with a dosage of $k=1$, while $\mathbf{m} = (0, 2, 0)$ produces gametes with dosages $k=0, 1, 2$ in a ratio of 1:2:1. The parent's gamete frequencies are a mixture of the frequencies from each configuration, weighted by their probabilities (lower left). The allopolyploid case occurs if only one pairing configuration is possible, whereas the autopolyploid case occurs if $\gamma_1 = 2/3$ and $\gamma_2 = 1/3$.

autopolyploids is given by Gerard et al. (2018). For example, for autotetraploids where $\ell=2$, the probability of configuration $\mathbf{m}_1 = (1, 0, 1)$ is $1/3$ and the probability of configuration $\mathbf{m}_2 = (0, 2, 0)$ is $2/3$, resulting in an autotetraploid segregation pattern of

$$\begin{aligned} p_1/3 + p_2 \times 2/3 \\ = (0, 1, 0)/3 + (1/4, 1/2, 1/4) \times 2/3 = (1, 4, 1)/6. \end{aligned} \quad (5)$$

More generally, this model accommodates segmental allopolyploids with arbitrary levels of partial preferential pairing.

At simplex loci, only one pairing configuration exists (and so $\gamma=1$). This simplifies modeling of double reduction at these loci. Theorem 1 states the gamete frequencies at simplex loci under arbitrary levels of double reduction (proof in Supplementary Section S1 in File S1). Upper bounds under two models for meiosis (Huang et al. 2019; Gerard 2022) appear in Supplementary Table S2 in File S1.

Theorem 1

Let α_i be the probability that there are i pairs of alleles in a gamete that are identical by double reduction. Then the gamete frequencies at a simplex locus are

$$\begin{aligned} p_0 &= \frac{1}{2} + \beta, \\ p_1 &= \frac{1}{2} - 2\beta, \\ p_2 &= \beta, \text{ and} \\ p_3 &= \dots = p_{K/2} = 0. \end{aligned} \quad (6)$$

where

$$\beta = \frac{1}{K} \sum_{i=1}^{\lfloor K/4 \rfloor} i\alpha_i. \quad (7)$$

To incorporate invalid genotypes, we mix the genotype frequencies with a discrete uniform distribution:

$$q_k = (1 - \pi) \sum_{i=\max(0, k-K/2)}^{\min(k, K/2)} p_{1i} p_{2, k-i} + \pi/(K+1), \quad (8)$$

where π represents the proportion of offspring with anomalous genotypes. By default, we set $\pi \leq 0.03$, but this is user-adjustable. This enables integrated modeling of genotype frequencies and outliers, unlike the two-step approach in `polymapR` (Bourke et al. 2018). This makes it possible to model outliers in the presence of genotype uncertainty in Likelihood ratio tests for segregation distortion section.

In summary, we model the genotype frequencies as a mixture of the discrete uniform distribution and the convolution of gamete frequencies (8). The gamete frequencies can be specified by the user to correspond to full autopolyploidy, full allopolyploidy, or segmental allopolyploidy through different constraints on γ in (4). This model offers several advantages over `polymapR` (Bourke et al. 2018), including support for segmental allopolyploidy (which `polymapR` does not accommodate) and the direct incorporation of invalid genotypes into the model, rather than addressing them through a two-step procedure. Additionally, this approach extends the methods of Gerard et al. (2025) to higher ploidy levels while incorporating the ability to handle invalid genotypes. We account for double reduction at simplex markers, where the effects of partial preferential pairing are absent. For clarity, in Supplementary Section S2 in File S1 we explicitly write out our model for the gamete frequencies for even ploidies 2 through 12.

Likelihood ratio tests for segregation distortion

In this section, we construct LRTs under the null hypothesis that the genotype frequencies correspond to one of the models of Models for genotype frequencies section. We will construct these tests either using known genotypes or by accounting for genotype uncertainty through genotype likelihoods.

The likelihood for the model depends on whether genotypes are known or unknown. In the known genotype case, we let x_k be the count of individuals with genotype $k \in \{0, 1, \dots, K\}$, which we organize into the vector $\mathbf{x} = (x_0, x_1, \dots, x_K)$, with total sample size $n = \sum_{k=0}^K x_k$. Given genotype frequencies $\mathbf{q}$, $\mathbf{x}$ follows a multinomial distribution, with the following likelihood

$$f(\mathbf{x} | \mathbf{q}) = \frac{n!}{\prod_{k=0}^K x_k!} \prod_{k=0}^K q_k^{x_k}. \quad (9)$$

In the unknown genotype case, we represent genotype uncertainty through genotype likelihoods (Li 2011). Let g_{ik} denote the genotype likelihood for individual $i = 1, 2, \dots, n$ and genotype $k = 0, 1, \dots, K$. Specifically, g_{ik} is the probability of the observed data (e.g. sequencing or microarray) for individual i , assuming their genotype is k . Given these genotype likelihoods, the likelihood function is

$$f(\mathbf{G} | \mathbf{q}) = \prod_{i=1}^n \sum_{k=0}^K g_{ik} q_k. \quad (10)$$

In the known genotype case (9), the maximum likelihood estimate of $\mathbf{q}$ under the alternative hypothesis is $\hat{\mathbf{q}}_A = \mathbf{x}/n$. In the unknown genotype case (10), $\hat{\mathbf{q}}_A$ can be calculated via the expectation–maximization (EM) algorithm of Li (2011). By maximizing either (9) or

(10) over the null parameter space, we obtain $\hat{\mathbf{q}}_0$. Using this, we calculate the likelihood ratio statistic (2 times the difference of the maximized log-likelihoods under the null and alternative), and we compare it to an appropriate χ^2 distribution to compute a P-value.

Our strategy to maximize over the null parameter space depends on the null hypothesis under consideration. If the null hypothesis corresponds to a true autopolyploid, the only unknown parameter is the invalid mixing proportion π in (8), which we maximize over using Brent's method (Brent 2013). If the null hypothesis corresponds to a true allopolyploid (2), the two unknowns are the latent pairing configuration and the invalid mixing proportion π in (8). We iterate over each of the $b_{K\ell}$ possible pairing configurations, optimizing over π via Brent's method for each configuration, to identify the pairing configuration and π that maximize the likelihood. If the null hypothesis corresponds to a segmental allopolyploid (4), the unknowns include the mixing proportions for pairing configurations in both parents, γ_1 and γ_2 , as well as the invalid genotype mixing proportion π . This optimization is performed using the method of Powell (2009), with initial values obtained via the global optimization method of Kucherenko and Sytsko (2005). The simplex constraints on the mixing proportions for pairing configurations are incorporated using the parameterization of Betancourt (2012).

The number of degrees of freedom for the null χ^2 distribution is the difference between the number of parameters under the alternative and under the null. The number of parameters under the alternative is just the ploidy of the offspring. Though, we can improve the power of the LRT if we subtract off the number of dosages whose frequencies are estimated to be 0 both under the null and under the alternative (Gerard et al. 2025).

Calculating the number of parameters under the null is tricky for two reasons: (i) some null parameters may lie on or near the boundary of the null parameter space, requiring extra consideration (Self and Liang 1987; Mitchell et al. 2019; Leung and Sturma 2024) and (ii) the null parameter space is weakly identified. To address the boundary problem (i), we adapt the data-dependent degrees of freedom strategy of Susko (2013). Specifically, the number of parameters under the null is estimated to be at most the number of parameters in θ that lie in the interior of the null parameter space.

For (ii), to define weak identification, let θ be a vector containing the null parameters (the γ 's, β 's, and π), then the Jacobian of the function $\theta \rightarrow \mathbf{q}$ is approximately low rank for some values of θ , which results in issues for the asymptotic performance of the LRT (Kleibergen 2005; Han and McCloskey 2019). To address this weak identifiability problem, we calculate the Jacobian matrix for the internal parameters at the MLE of the function $\theta \rightarrow \mathbf{q}$. The numeric rank of this Jacobian matrix is used as the number of free parameters under the null. By default, the numeric rank we used is the number of singular values that are at least one-thousandth as large as the largest singular value. This heuristic is based on the literature that shows that the number of parameters in an exactly nonidentified parameter space is the exact rank of the Jacobian (Catchpole and Morgan 1997; Viallefont et al. 1998; Schmittmann et al. 2010). Thus, we use the approximate rank of the Jacobian to approximate the number of parameters in a weakly identified parameter space. This heuristic works well in practice (Results section), though we have not seen its use in prior literature.

The above scenarios assume the parental genotypes are known. In the case that parental genotypes are not known, we first maximize the likelihood jointly over the parental genotypes and the unknown null parameters (the γ 's, β 's, and π). We then proceed

with the LRT using the MLEs of the parent genotypes as if they were the true parent genotypes.

Implementation features

All methods described in this manuscript are implemented in the user-friendly `segtest` R package on the Comprehensive R Archive Network (<https://doi.org/10.32614/CRAN.package.segtest>). The package includes additional features that may be useful for applied researchers:

- 1) Bayesian Information Criterion (BIC): `segtest` computes the BIC (Schwarz 1978) of the fitted null model using the estimated number of parameters under the null. Researchers can visualize BIC distributions across loci to assess model suitability, particularly when segregation patterns are unknown (e.g. to determine whether partial preferential pairing is present). The model with lower BIC values is generally preferred.
- 2) Outlier Detection: For each individual, `segtest` calculates the posterior probability of having an anomalous genotype at a given locus, allowing researchers to treat such genotypes as missing in downstream analyses. This follows a standard mixture model approach. For a K -ploid individual, let $q_0, \dots, q_K$ be genotype likelihoods under the assumption of no outliers, computed via (1). The likelihoods for an individual as a nonoutlier (f_1) and as an outlier (f_0) are defined as $f_1 = q_k$ and $f_0 = \frac{1}{K+1}$ when the genotype $k \in \{0, \dots, K\}$ is known. In the unknown genotype case, given individual genotype likelihoods $g_0, \dots, g_K$, these become $f_1 = \sum_{k=0}^K g_k q_k$ and $f_0 = \frac{1}{K+1} \sum_{k=0}^K g_k$. For a given outlier proportion π , the posterior probability that an individual is an outlier is then

$$\frac{\pi f_0}{\pi f_0 + (1 - \pi) f_1}. \quad (11)$$

- 3) Integration with `updog`: `segtest` includes helper functions to format output from `updog` (Gerard et al. 2018; Gerard and Ferrão 2019) for direct input.
- 4) Parallelization Support: The package supports parallel execution via the `future` package (Bengtsson 2021), with additional compatibility for high-performance computing through `future.batchtools`.

Copy editing

GPT-4o (<https://chatgpt.com/>) was used in this manuscript for light copy editing.

Results

Null simulations

We assessed the type I error control of our new LRT and the `polymapR` test through simulations under the null hypothesis of no segregation distortion. The following parameters were varied:

- Ploidy: $K \in \{4, 6, 8\}$
- Parent genotypes: $\ell_1 \in \{0, \dots, K\}$ and $\ell_2 \in \{\ell_1, \dots, K\}$
- Sample size: $n \in \{20, 200\}$
- Read depth: 10 (unknown genotype case) or infinity (known genotype case)
- Outlier proportion: $\pi \in \{0, 0.015, 0.03\}$
- Mixing proportions for the pairing configuration (for non-nulliplex and nosimplex loci):

- If the number of mixture components is 2: $(\gamma_1, \gamma_2) \in \{(1, 0), (0.5, 0.5), (0, 1)\}$
- For $K = 8$ and $\ell = 4$, where three mixture components exist:

$$(\gamma_1, \gamma_2, \gamma_3) \in \{(1, 0, 0), (0, 1, 0), (0, 0, 1), (1, 1, 0)/2, (1, 0, 1)/2, (0, 1, 1)/2, (1, 1, 1)/3\}$$

- Double reduction adjustment for simplex loci: $\beta \in \{0, a/2, a\}$, where a is the maximum value under the complete equation-
al segregation (CES) model in [Supplementary Table S2](#) in File S1.

This yielded a total of 7,920 distinct simulation scenarios to evaluate the null performance of our methods. For each scenario, genotype frequencies $\mathbf{q}_0$ were determined. In the known genotype case, genotype counts were drawn from a multinomial distribution (9). In the unknown genotype case, these simulated counts were further processed to generate individual sequencing read counts according to the model of Gerard et al. (2018), assuming no allele bias, a sequencing error rate of 0.01, and an overdispersion parameter of 0.01. Genotype likelihoods were then obtained using the method of Gerard et al. (2018). Each replication involved fitting either our LRT or the `polymapR` test (Bourke et al. 2018). A total of 200 replications were performed per simulation scenario.

Figure 2 presents histograms of the type I error rates across the 7,920 simulation scenarios at a nominal significance level of 0.05. Since the expected type I error rate should not exceed 0.05, any observed values above this threshold should only be attributable to the finite number of simulations. Specifically, the distribution of type I errors should be stochastically bounded by $X/200$, where X follows a binomial distribution with size 200 and success probability 0.05. For reference, Fig. 2 includes the 99th percentile of this distribution as a blue dotted line. The results indicate that `segtest` adequately controls type I error for large sample sizes ($n = 200$) and for nearly all scenarios with small sample sizes ($n = 20$). In some small-sample scenarios, slight anticonservatism is observed, which is expected given that the LRT only asymptotically controls type I error.

In contrast, `polymapR` fails to control type I error in many scenarios, particularly for large sample sizes when genotypes are unknown. To illustrate this issue, Fig. 3 presents results from select simulation scenarios. These simulations were conducted with a sample size of $n = 200$ octoploid ($K = 8$) offspring, where one parent had a genotype of 0 and the other had a genotype of 6. The figure contains quantile–quantile (Q–Q) plots of P -values against the uniform distribution, stratified by method (`segtest` or `polymapR`), read depth (10 or infinity), and pairing configuration mixing proportions for the second parent ((0.5, 0.5) or (1, 0)). Under the null hypothesis, the Q–Q plots should lie entirely at or above the $y = x$ line (black), indicating appropriate type I error control.

`Segtest` satisfies this criterion in all scenarios, demonstrating proper type I error control at any significance level. In contrast, `polymapR` fails to control type I error in three cases. The method does not account for nonabsolute preferential pairing, meaning that $\gamma = (0.5, 0.5)$ represents an alternative scenario under its test, explaining its failure to control type I error at any read depth in this case. Additionally, `polymapR` fails to control type I error at low read depths even when absolute preferential pairing ($\gamma = (1, 0)$) is assumed, a scenario that should fall within its null model. This failure arises because `polymapR` first estimates genotype counts before applying its known-genotype test, and this estimation process is biased (Gerard et al. 2025).

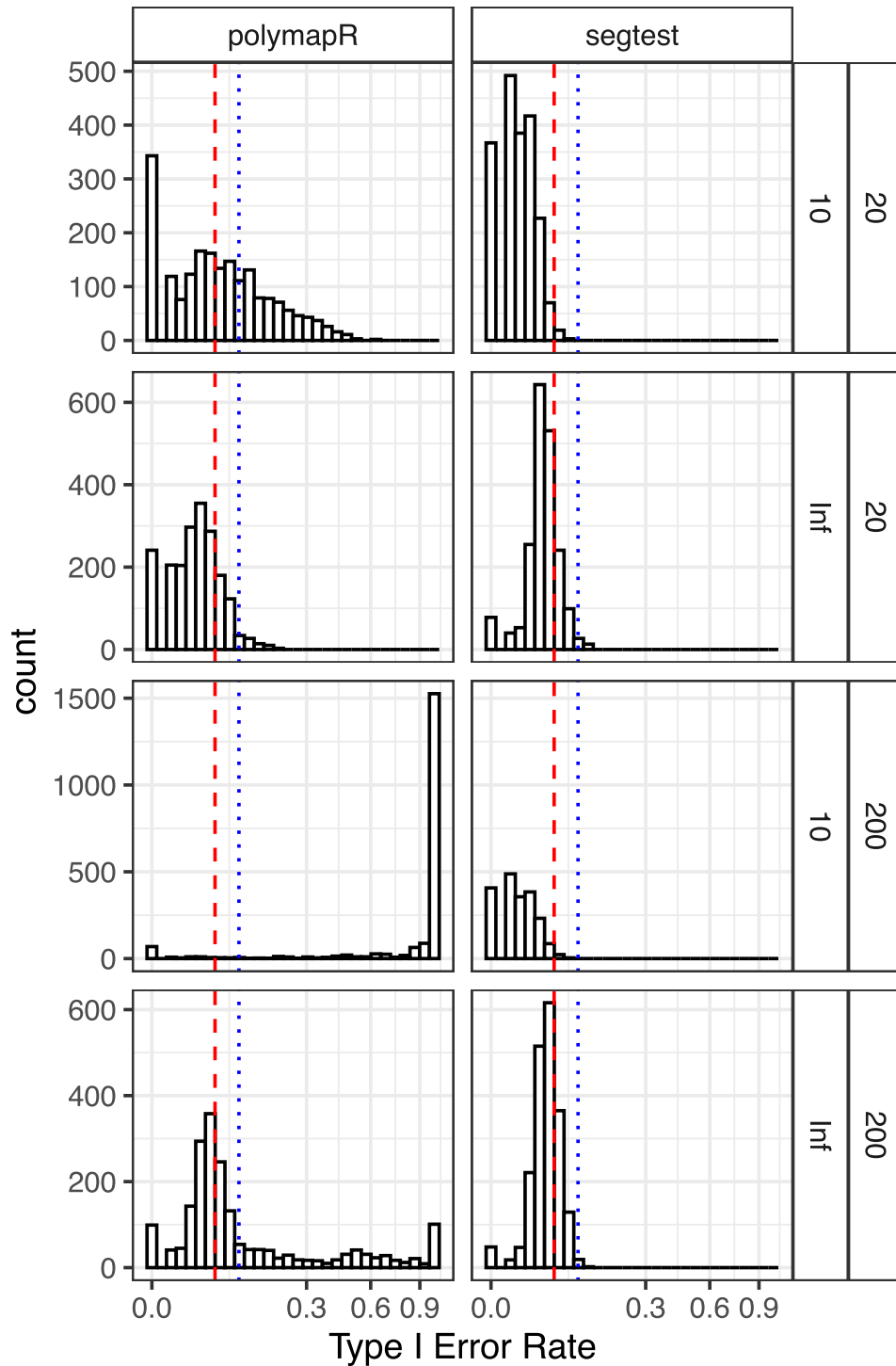


Fig. 2. Histograms of type I error rates (on the square root scale) from the simulations in *Null simulations* section. Tests were conducted at a nominal significance level of 0.05. Results are stratified by sample size ($n = 20$ or 200) and read depth (10 or infinity) in the rows, and by method (**segtest** or **polyploidR**) in the columns. Since the null hypothesis is true, the type I error rate should not exceed 0.05 (red dashed line), though small deviations are expected due to finite simulations. The blue dotted line marks the 99th percentile of expected variation under proper type I error control, based on binomial quantiles. **segtest** controls type I error in all scenarios for large $n = 200$ and most for small $n = 20$, as expected given its asymptotic guarantees. In contrast, **polyploidR** fails to control type I error at low read depths and for large sample sizes. It appears to control type I error at high read depths with small n , likely due to low power for small n .

Robustness to double reduction

Since our method does not explicitly account for double reduction at nonsimplex loci, we conducted simulations under the autopolyploid model of Huang et al. (2019), which allows for arbitrary levels of double reduction, though no preferential pairing. We varied the following parameters:

- Ploidy: $K \in \{4, 6, 8\}$
- Parent genotypes: $\ell_1 \in \{0, \dots, K\}$ and $\ell_2 \in \{2, \dots, K-2\}$
- Sample size: $n \in \{20, 200\}$
- Double reduction rate: $\alpha \in \{0, \alpha_m/2, \alpha_m\}$, where α_m is the maximum double reduction rate(s) under the CES model (Huang et al. 2019)

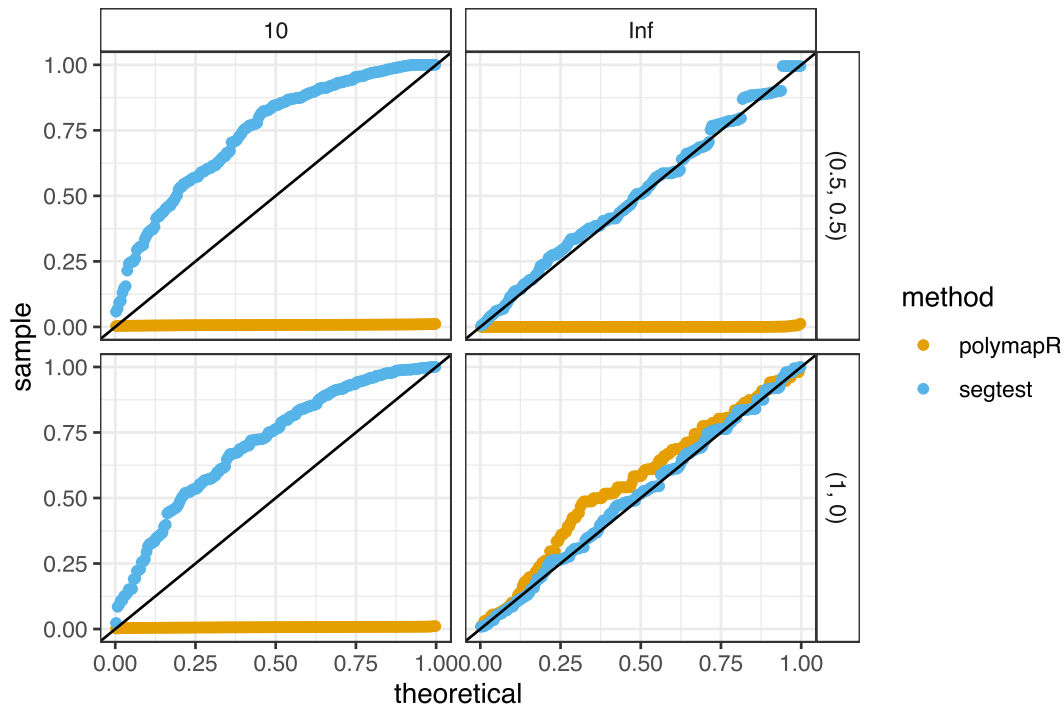


Fig. 3. Example results from the null simulations in *Null simulations* section. Q-Q plots compare P -values from *segtest* (blue) and *polymapR* (orange) against the uniform distribution. Plots are stratified by read depth (columns) and pairing configuration (γ_2 , rows). This scenario assumes a ploidy of $K = 8$, parent genotypes $\ell_1 = 0$ and $\ell_2 = 6$, and a sample size of $n = 200$. Under the null, tests that properly control type I error should lie at or above the $y = x$ line (black). *PolymapR* satisfies this condition only when $\gamma_2 = (1, 0)$ (bottom row) and genotypes are fully known (infinite read depth). However, it fails to control type I error when preferential pairing is not absolute (top row) or when genotypes are uncertain (left column), even under its assumed null scenario. This is due to its ad hoc genotype uncertainty adjustment. In contrast, *segtest* consistently controls type I error across all scenarios.

- Read depth: 10 (unknown genotype case) or infinity (known genotype case)

The outlier proportion was fixed at 0 (no outliers), resulting in 852 distinct simulation scenarios. Data were simulated as described in *Null simulations* section. For each replication, we applied the LRT assuming no outliers, the LRT allowing for outliers, and the *polymapR* test (Bourke et al. 2018), conducting 200 replications per scenario.

Figure 4 presents histograms of type I error rates at a nominal level of 0.05 across the simulated scenarios. When double reduction was absent, *segtest*, regardless of whether it accounted for outliers, maintained appropriate type I error control. In contrast, *polymapR* failed to control type I error at low read depths, consistent with the findings in *Null simulations* section.

At moderate levels of double reduction, only *segtest* with outlier accommodation remained relatively robust. There were eight simulation scenarios where *segtest* exhibited poor type I error control (above the 99th percentile of the Binomial(200, 0.05) distribution). These cases exclusively involved nulliplex-by-duplex parental genotypes at ploidies 6 and 8, with large sample sizes ($n = 200$) and known genotypes (Supplementary Table S3 in File S1). Notably, genotype uncertainty appeared to render the test conservative enough (see *Null simulations* section) to maintain adequate type I error control in the presence of moderate double reduction. Genotype uncertainty is the most common scenario in applied work (Gerard et al. 2018; Gerard and Ferrão 2019).

At extreme levels of double reduction, only *segtest* with outlier handling performed reasonably well, though 30 of the 852 scenarios still exhibited poor type I error control. In such cases, type I error inflation was more pronounced. Though, extremely

high double reduction rates are believed to be rare (Stift et al. 2008; Bomblies et al. 2016). And in practice, such deviations would likely be identified by researchers in applied settings.

Power analysis

In this section, we evaluate the power of our methods under segregation distortion. We simulated genotype frequencies for ploidies $K \in \{4, 6, 8\}$ under two alternative scenarios:

- Easy scenario: The genotype frequency vector $\mathbf{q}$ was drawn uniformly from the K -dimensional unit simplex.
- Hard scenario: Gamete frequency vectors $\mathbf{p}_1$ and $\mathbf{p}_2$ were first drawn independently and uniformly from the $K/2$ -dimensional unit simplex and then convolved (1) to obtain $\mathbf{q}$.

The hard scenario generates alternative cases that more closely resemble the null, as the genotype frequencies are obtained by convolving randomly generated gamete frequencies, mirroring the null model structure. We further varied sample size ($n \in \{20, 200\}$) and read depth (10 or infinite). For each of 1,000 replications, we simulated $\mathbf{q}$, generated data as described in *Null simulations* section, and applied both the LRT from *segtest* and the *polymapR* test.

Figure 5 presents the power of *segtest* and *polymapR* at a nominal significance level of 0.05. Because *segtest* properly controls type I error while *polymapR* does not, some reduction in power for *segtest* is expected. However, we find that *segtest* typically exhibits only a modest decrease in power, particularly for large samples ($n = 200$), where both methods perform well. In some hard scenarios with small n , *segtest* shows a more pronounced power reduction. Nonetheless, hypothesis testing requires valid type I error control, which *segtest* ensures,

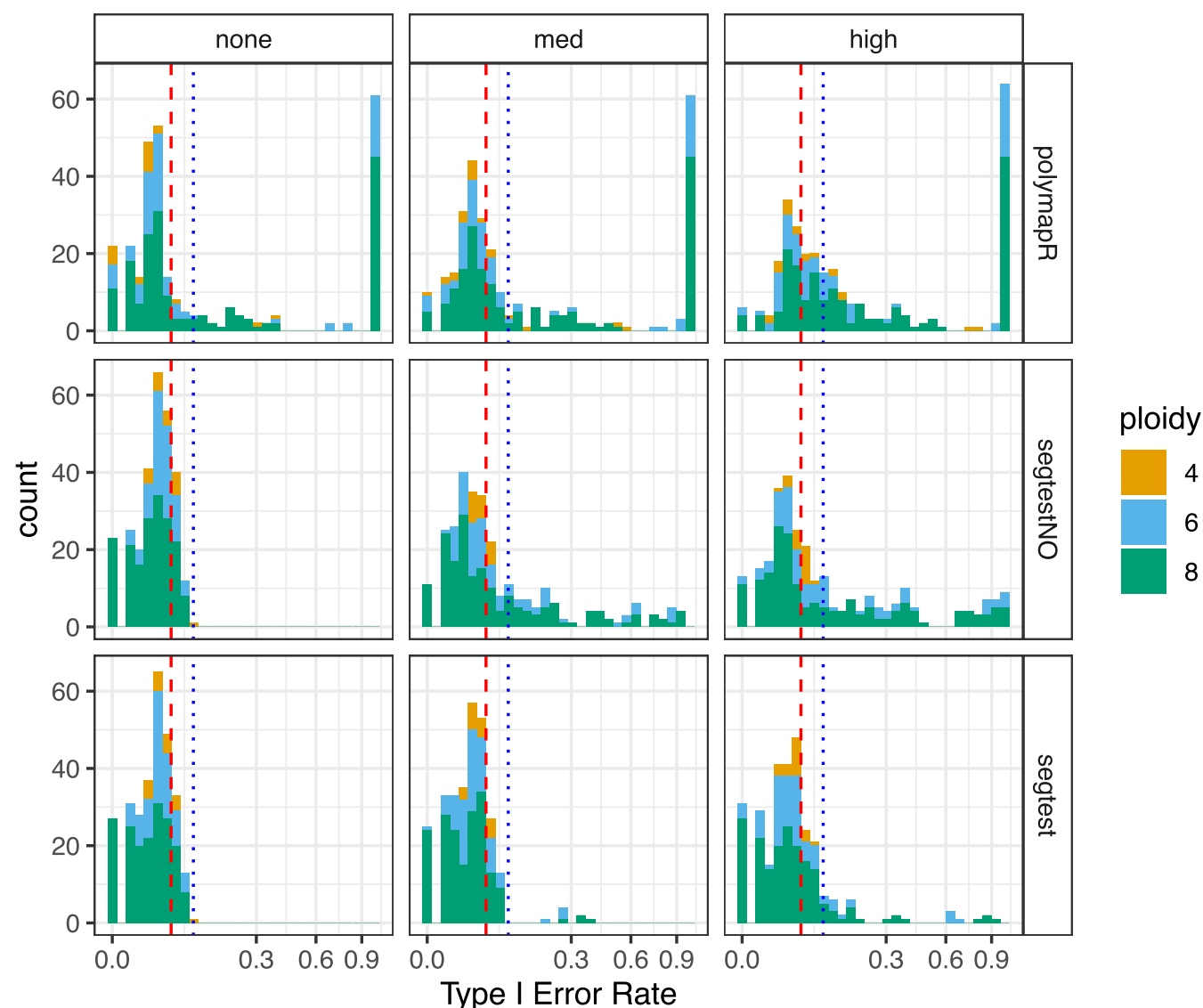


Fig. 4. Robustness to double reduction results from *Robustness to double reduction* section. Histograms of type I error rates (square root scale) at a nominal significance level of 0.05 (red dashed line). The blue dotted line indicates the 99th percentile of expected variation under proper type I error control, based on binomial quantiles. Results are stratified by the level of double reduction (none, medium, or high) along the columns and by method (`polymapR`, `segtest` assuming no outliers, and `segtest` allowing for outliers) along the rows. Ploidy is color-coded. Among the tested methods, only `segtest` with outlier accommodation remains relatively robust to the effects of double reduction, particularly at moderate levels.

whereas `polymapR` does not. For sufficiently large sample sizes, `segtest` remains consistent and maintains power against all tested scenarios.

A hexaploid F1 population

To validate our methods, we analyzed the hexaploid sweetpotato (*Ipomoea batatas*, $2n = 6x = 90$) mapping population from Mollinari et al. (2020), which consists of next-generation sequencing data from 315 full-sib individuals derived from a cross between the Beauregard and Tanzania varieties. These data include all SNPs from chromosome 8 (52,125 SNPs). After filtering out multiallelic loci (51,988 SNPs remaining), we performed genotyping using the method of Gerard et al. (2018) and Gerard and Ferrão (2019) with the “f1” model.

We tested for segregation distortion in these data using three different procedures. The first method emulated the default behavior of `mappoly` (Mollinari et al. 2020), following this procedure:

- Filtered out SNPs estimated to be nulliplex in both parents (46,434 SNPs remaining).
- For each SNP, filtered out individuals with a maximum genotype posterior probability below 0.95 (the default behavior of `mappoly`).
- Filtered out SNPs with an average read depth less than 20 (16,976 SNPs remaining), as done in Mollinari et al. (2020).
- Filtered out SNPs with at least 25% missing individuals (3,448 SNPs remaining), as done in Mollinari et al. (2020).

For the remaining SNPs, we used the posterior mode genotype as the “known” genotype and ran chi-squared tests for segregation distortion, assuming polysomic inheritance with bivalent pairing.

The second method used the `checkF1()` function from `polymapR` (Bourke et al. 2018). Because this function accounts for genotype uncertainty, we only applied two filters: filtering out SNPs estimated to be nulliplex in both parents (46,434 SNPs remaining) and SNPs with an average read depth below 20 (16,105

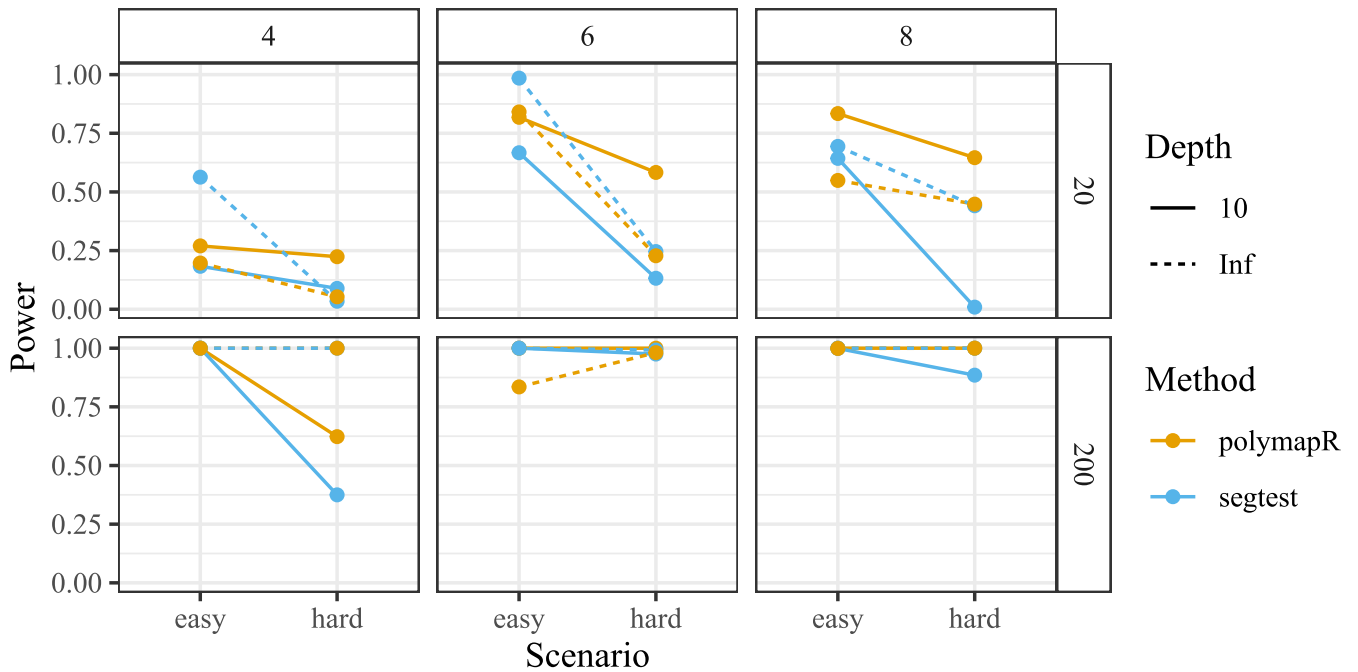


Fig. 5. Results of the alternative simulations from *Power analysis* section. The figure plots the power (y-axis) of *segtest* and *polymapR* (color) under scenarios where segregation distortion is present. Tests were performed at a nominal significance level of 0.05 across varying sample sizes ($n \in \{20, 200\}$) (row facets), ploidies ($K \in \{4, 6, 8\}$) (column facets), read depths (10 or infinite) (line type), and alternative scenarios (easy or hard) (x-axis). *segtest* exhibits strong power for large n and maintains reasonable power in most small n scenarios. While *segtest* has slightly lower power than *polymapR* in most settings, this reduction is expected due to its ability to control type I error, which *polymapR* does not.

SNPs remaining). We then applied the *polymapR* test using genotype posterior probabilities.

The third method applied our new approach based on genotype likelihoods, implemented in the *segtest* software. We applied the same prefilters as in *polymapR*, analyzing 16,105 SNPs. We ran our new likelihood ratio test (*Likelihood ratio tests for segregation distortion* section) using three null models: our general model (*Models for genotype frequencies* section), a polysomic inheritance model that allows for double reduction (Huang et al. 2019), and a polysomic inheritance model without double reduction (Serang et al. 2012). For each locus, we computed the BIC of the null model (*Implementation features* section). We compared BICs pairwise at loci where the more complex null model had a P-value greater than 0.1. The polysomic model had a lower BIC than the general model 59.74% of the time (95% CI: 58.70% to 60.77%), the double reduction model had a lower BIC than the general model 55.85% of the time (95% CI: 54.80% to 56.90%), and the double reduction model had a lower BIC than the simpler polysomic model 51.66% of the time (95% CI: 50.57% to 52.75%). These results suggest that the polysomic inheritance model allowing for double reduction best fits these data, and thus tests from *segtest* based on this model are used in the comparisons that follow.

Histograms of the P-values from the three methods are shown in Fig. 6. The P-value density for *polymapR* exhibits a mode at 1, which is generally not expected for unbiased P-values (Storey and Tibshirani 2003). In contrast, both *mappoly* and *segtest* display the more typical monotonically decreasing P-value densities. When controlling the false discovery rate at level 0.05, *mappoly* identifies far fewer SNPs (2,859 SNPs) not in segregation distortion than *polymapR* (13,796 SNPs) or *segtest* (10,058 SNPs). This suggests that *mappoly*'s filtering procedure, which does not account for genotype uncertainty, may be overly

conservative, potentially discarding informative SNPs and leading to unnecessary data loss. Among SNPs in common after prefiltering, the P-values from *segtest* and *mappoly* are correlated at 0.6996, those from *segtest* and *polymapR* at 0.4311, and those from *polymapR* and *mappoly* at 0.8169.

It is informative to examine cases where the methods disagree. [Supplementary Fig. S2](#) in File S1 shows genotype plots (Gerard et al. 2018) for four SNPs where *mappoly* produces P-values less than 0.05 while *segtest* produces P-values greater than 0.1, and four SNPs with the opposite pattern. The P-values for these SNPs are listed in [Supplementary Table S4](#) in File S1. Because *mappoly* filters out individuals with low genotype posterior probabilities before testing for segregation distortion, this alters the observed genotype frequencies and may lead to anticonservative results. For example, if we include those individuals in the *mappoly* test, the P-value becomes nonsignificant ([Supplementary Table S4](#) in File S1). Conversely, *segtest* often yields small P-values for SNPs with many outlier genotypes, which *mappoly* has prefiltered as “invalid.” These outlier individuals typically have low (or no) reference reads and may reflect null alleles (Gautier et al. 2013). When we remove individuals with fewer than 30 reference reads, *segtest* produces larger P-values ([Supplementary Table S4](#) in File S1). Whether such SNPs should be flagged is debatable, but in cases like SNP S8_2212186, where 21 of 315 individuals ($\approx 6.7\%$) have “invalid” genotypes, flagging seems appropriate. Alternatively, if a researcher does not wish to flag such SNPs, instead of mixing the genotype frequencies with a discrete uniform in (8), we can mix them with an outlier distribution that one would expect in the presence of null alleles (a pointmass at a genotype of 0). When doing so, we again obtain much larger P-values ([Supplementary Table S4](#) in File S1).

We conducted a similar comparison between *polymapR* and *segtest*. [Supplementary Fig. S3](#) in File S1 displays four SNPs where *polymapR* indicates segregation distortion (P-value < 0.05)

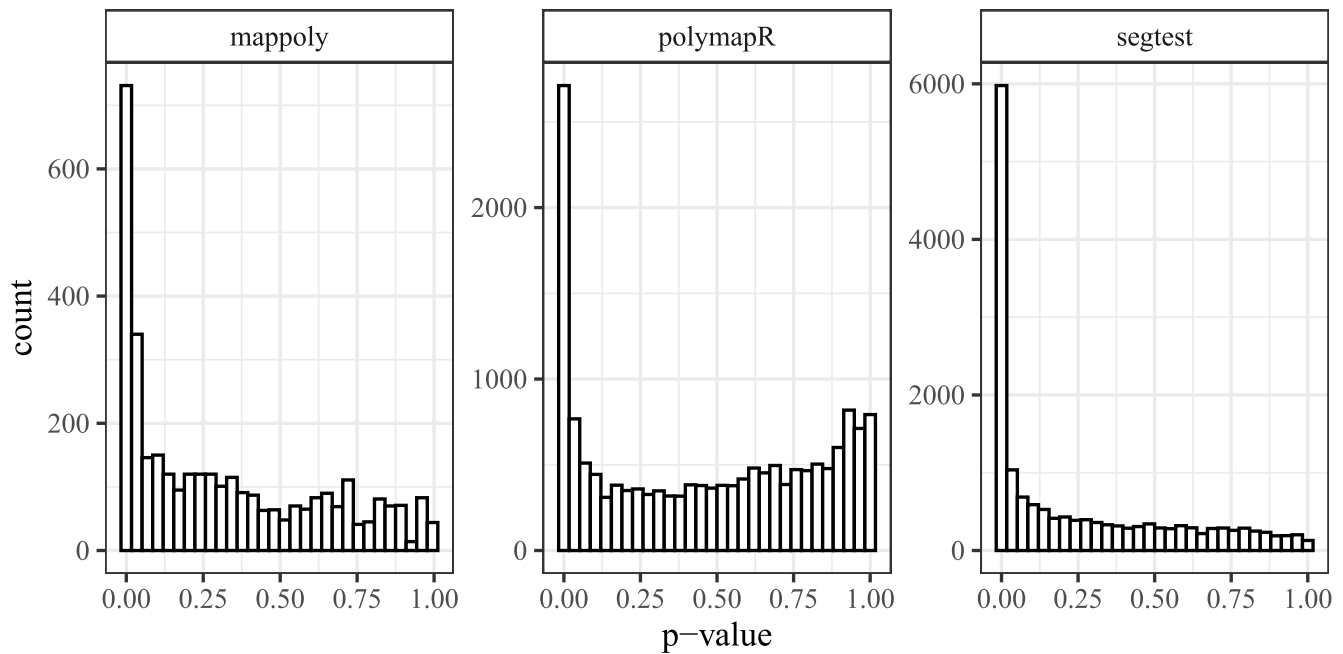


Fig. 6. Sweetpotato P-values. Histograms of the P-values from the segregation distortion tests performed by mappoly, polymapR, and segtest on the sweetpotato data described in A hexaploid F1 population section. Note that the y-axes differ between facets.

but segtest does not (P -value > 0.1), and four SNPs where the reverse is true. The P -values from both methods are reported in [Supplementary Table S5](#) in File S1. Like mappoly, polymapR removes “invalid” genotypes prior to testing. Thus, we again find that outliers explain the SNPs flagged by segtest but not by polymapR. When individuals with low reference reads are excluded, segtest returns much larger P -values ([Supplementary Table S5](#) in File S1). Segtest also produces much larger P -values when the outlier distribution is a pointmass at a genotype of 0, thereby accounting for null alleles ([Supplementary Table S5](#) in File S1). To understand cases where polymapR produces low P -values but segtest does not, we refer to [Supplementary Table S6](#) in File S1, which shows the null and alternative genotype frequencies compared by each method. Both use the same alternative model, so the alternative frequencies are similar, differing only slightly due to polymapR’s ad hoc estimation versus segtest’s maximum likelihood approach (Li 2011). The key difference lies in the null frequencies. Segtest can model double reduction, allowing its null frequencies to better match the alternative frequencies than those used by polymapR. In these cases, the ability to account for double reduction appears to explain why segtest does not detect segregation distortion, while polymapR does.

Discussion

In this study, we developed statistical tests for segregation distortion in F1 populations of arbitrary (even) ploidy. Our methods account for varying levels of partial preferential pairing, double reduction at simplex loci, and a proportion of outliers. Effective modeling of outliers enhances robustness to moderate levels of double reduction at nonsimplex loci. Additionally, our approach accommodates genotype uncertainty through genotype likelihoods. While our methods have only asymptotic guarantees for controlling type I error, we improve finite-sample performance by adaptively estimating the degrees of freedom, by counting boundary parameters and approximating the rank of a particular

Jacobian. As a result, our LRT controls type I error reasonably well for sample sizes as small as 20 and performs robustly for samples of size 200.

For tetraploids, our methods fully and jointly account for double reduction and (partial) preferential pairing, as they generalize the approach of Gerard et al. (2025). That is, due to the unidentifiability of the double reduction rate and the preferential pairing parameters for tetraploids, as described in Gerard et al. (2025), one can without loss of generality set the double reduction rate to 0 (at duplex loci), which is essentially what our method does for tetraploids in this manuscript. However, for ploidies greater than four, our model does not explicitly account for double reduction at nonsimplex loci. For example, in hexaploids ($K = 6$), double reduction could theoretically produce gamete genotypes of three at a duplex locus, but our model does not allow for this possibility ([Supplementary Section S2](#) in File S1). Nonetheless, by incorporating a small proportion of outliers, our approach demonstrates robustness to moderate levels of double reduction (*Robustness to double reduction* section). Increasing the maximum outlier proportion may further enhance type I error control in cases of extreme double reduction, though this would likely come at the cost of reduced power.

Our results here focus on scenarios that permit arbitrary levels of preferential pairing while explicitly modeling double reduction at simplex loci and allowing for a small fraction of outliers. However, our segtest package offers greater flexibility, enabling segregation distortion testing under additional assumptions:

- 1) True autopolyploidy, assuming either complete bivalent pairing (no double reduction) or arbitrary levels of double reduction, modeled as in Huang et al. (2019).
- 2) True allopolyploidy with complete preferential pairing.
- 3) An unknown ploidy structure, where the data could correspond to either a true allopolyploid or a bivalent-pairing autopolyploid. This is the assumption used by polymapR.

Additionally, segtest allows for testing without permitting outliers, modeling segregation distortion when any individual exhibits an

“impossible” genotype. Each of these scenarios represents a special case of the general model studied here. If researchers have prior knowledge about the expected segregation patterns in their organism, they should specify the most appropriate null model.

In this manuscript, we focused on applying these tests as a quality control procedure for methods that use experimental populations, such as linkage or quantitative trait locus mapping. However, our methods may have broader applications. Because we implemented multiple models of meiosis and can calculate the BIC (Schwarz 1978) for each (Implementation features section), our approach could also be used to estimate the meiotic behavior in polyploids. We demonstrated this in *A hexaploid F1 population* section, where we found that a polysomic inheritance model with double reduction best fit the sweetpotato data. This result aligns with the findings of Mollinari et al. (2020), who reported no evidence of preferential pairing on Chromosome 8 (the chromosome we analyzed here) yet found some evidence of multivalent pairing, a prerequisite for double reduction (Stift et al. 2010). Further study is needed to assess whether our approach is broadly applicable for distinguishing meiotic models (allo-, auto-, or segmental polyploid).

We sought to improve the finite sample performance of the likelihood ratio test by adaptively estimating its degrees of freedom. An alternative approach would have been to use the bootstrap (Efron 1979), which is known to provide highly accurate finite-sample results across many model classes (Efron 1987). However, we did not pursue this approach primarily due to computational constraints. Our tests are designed for genomic applications, where they may be applied thousands or even millions of times within a single study. Any method requiring more than approximately one second per test imposes an undue burden on applied researchers. Nonetheless, as computational power continues to increase, the bootstrap remains a promising avenue for future exploration.

Mappoly assumes genotypes are known when performing genetic mapping, and so it is natural that it filters out individuals with high genotype uncertainty. However, our results in *A hexaploid F1 population* section suggest that filtering individuals prior to testing for segregation distortion can lead to biased tests and unnecessary data loss. We recommend reversing the order of operations: rather than filtering out uncertain genotypes before testing, researchers should first apply a test that accounts for genotype uncertainty (e.g. our approach), and then filter out individuals with high uncertainty. Understanding how different data quality control steps influence downstream applications such as genetic mapping is an important question for future study.

Data availability

All methods described in this article are implemented in the *segtest* R package on the Comprehensive R Archive Network under a GNU General Public License v3.0 or later: <https://doi.org/10.32614/CRAN.package.segtest>. All analysis scripts to reproduce the results of this article are available on Zenodo under a GNU General Public License v3.0 or later: <https://doi.org/10.5281/zenodo.15784734> and <https://doi.org/10.5281/zenodo.15784738>. The data used in this manuscript is from Mollinari et al. (2020) and is available on Figshare under a Creative Commons BY 4.0 license (Mollinari et al. 2019): <https://doi.org/10.25387/g3.10255844>.

Supplemental material available at G3 online.

Acknowledgments

Most analyses were performed using the R statistical language (R Core Team 2025).

Funding

None declared.

Conflicts of interest. None declared.

Literature cited

- Adams KL, Wendel JF. 2005. Polyploidy and genome evolution in plants. *Curr Opin Plant Biol.* 8:135–141. <https://doi.org/10.1016/j.pbi.2005.01.001>.
- Amadeu RR, Muñoz PR, Zheng C, Endelman JB. 2021. QTL mapping in outbred tetraploid (and diploid) diallele populations. *Genetics.* 219:iyab124. <https://doi.org/10.1093/genetics/iyab124>.
- Baird NA et al. 2008. Rapid SNP discovery and genetic mapping using sequenced RAD markers. *PLoS One.* 3:1–7. <https://doi.org/10.1371/journal.pone.0003376>.
- Batista LG, Mello VH, Souza AP, Margarido GR. 2021. Genomic prediction with allele dosage information in highly polyploid species. *Theor Appl Genet.* 135:1–17. <https://doi.org/10.1007/s00122-021-03994-w>.
- Bengtsson H. 2021. A unifying framework for parallel and distributed processing in R using futures. *R J.* 13:208–227. <https://doi.org/10.32614/RJ-2021-048>.
- Betancourt M. 2012. Cruising the simplex: Hamiltonian Monte Carlo and the Dirichlet distribution. *AIP Conf Proc.* 1443:157–164. <https://doi.org/10.1063/1.3703631>.
- Bombliès K, Jones G, Franklin C, Zickler D, Kleckner N. 2016. The challenge of evolving stable polyploidy: could an increase in “crossover interference distance” play a central role? *Chromosoma.* 125:287–300. <https://doi.org/10.1007/s00412-015-0571-4>.
- Bourke PM, et al.. 2017. Partial preferential chromosome pairing is genotype dependent in tetraploid rose. *Plant J.* 90:330–343. <https://doi.org/10.1111/tjp.13496>.
- Bourke PM et al. 2018. polypmapR-linkage analysis and genetic map construction from F1 populations of outcrossing polyploids. *Bioinformatics.* 34:3496–3502. <https://doi.org/10.1093/bioinformatics/bty371>.
- Bourke PM, Voorrips RE, Visser RGF, Maliepaard C. 2015. The double-reduction landscape in tetraploid potato as revealed by a high-density linkage map. *Genetics.* 201:853–863. <https://doi.org/10.1534/genetics.115.181008>.
- Brent R. 2013. Algorithms for minimization without derivatives. Dover Books on Mathematics.
- Cappai F et al. 2020. High-resolution linkage map and QTL analyses of fruit firmness in autotetraploid blueberry. *Front Plant Sci.* 11: 562171. <https://doi.org/10.3389/fpls.2020.562171>.
- Catchpole EA, Morgan BJ. 1997. Detecting parameter redundancy. *Biometrika.* 84:187–196. <https://doi.org/10.1093/biomet/84.1.187>.
- da Silva Pereira G et al. 2020. Multiple QTL mapping in autopolyploids: a random-effect model approach with application in a hexaploid sweetpotato full-sib population. *Genetics.* 215: 579–595. <https://doi.org/10.1534/genetics.120.303080>.
- de C Lara LA et al. 2019. Genomic selection with allele dosage in *Panicum maximum* jacq. G3 (Bethesda). 9:2463–2475. <https://doi.org/10.1534/g3.118.200986>.
- Doyle JJ, Egan AN. 2010. Dating the origins of polyploidy events. *New Phytol.* 186:73–85. <https://doi.org/10.1111/j.1469-8137.2009.03118.x>.
- Doyle JJ, Sherman-Broyles S. 2017. Double trouble: taxonomy and definitions of polyploidy. *New Phytol.* 213:487–493. <https://doi.org/10.1111/nph.14276>.
- Efron B. 1979. Bootstrap methods: another look at the jackknife. *Ann Stat.* 7:1–26. <https://doi.org/10.1214/aos/1176344552>.

- Efron B. 1987. Better bootstrap confidence intervals. *J Am Stat Assoc.* 82:171–185. <https://doi.org/10.1080/01621459.1987.10478410>.
- Elshire RJ et al. 2011. A robust, simple genotyping-by-sequencing (GBS) approach for high diversity species. *PLoS One.* 6:1–10. <https://doi.org/10.1371/journal.pone.0019379>.
- Endelman JB et al. 2018. Genetic variance partitioning and genome-wide prediction with allele dosage information in autotetraploid potato. *Genetics.* 209:77–87. <https://doi.org/10.1534/genetics.118.300685>.
- Ferrão LFV, Amadeu RR, Benevenuto J, de Bem Oliveira I, Munoz PR. 2021. Genomic selection in an outcrossing autotetraploid fruit crop: lessons from blueberry breeding. *Front Plant Sci.* 12: 676326. <https://doi.org/10.3389/fpls.2021.676326>.
- Gautier M, et al. 2013. The effect of RAD allele dropout on the estimation of genetic variation within and between populations. *Mol Ecol.* 22:3165–3178. <https://doi.org/10.1111/mec.12089>.
- Gemenet DC et al. 2020. Quantitative trait loci and differential gene expression analyses reveal the genetic basis for negatively associated β -carotene and starch content in hexaploid sweetpotato [*Ipomoea batatas* (L.) lam.]. *Theor Appl Genet.* 133:23–36. <https://doi.org/10.1007/s00122-019-03437-7>.
- Gerard D. 2021a. Pairwise linkage disequilibrium estimation for polyploids. *Mol Ecol Resour.* 21:1230–1242. <https://doi.org/10.1111/1755-0998.13349>.
- Gerard D. 2021b. Scalable bias-corrected linkage disequilibrium estimation under genotype uncertainty. *Heredity (Edinb).* 127: 357–362. <https://doi.org/10.1038/s41437-021-00462-5>.
- Gerard D. 2022. Double reduction estimation and equilibrium tests in natural autopolyploid populations. *Biometrics.* 79:2143–2156. <https://doi.org/10.1111/biom.13722>.
- Gerard D, Ferrão LFV. 2019. Priors for genotyping polyploids. *Bioinformatics.* 36:1795–1800. <https://doi.org/10.1093/bioinformatics/btz852>.
- Gerard D, Ferrão LFV, Garcia AAF, Stephens M. 2018. Genotyping polyploids from messy sequencing data. *Genetics.* 210:789–807. <https://doi.org/10.1534/genetics.118.301468>.
- Gerard D, Thakkar M, Ferrão LFV. 2025. Tests for segregation distortion in tetraploid F1 populations. *Theor Appl Genet.* 138:30. <https://doi.org/10.1007/s00122-025-04816-z>.
- Han S, McCloskey A. 2019. Estimation and inference with a (nearly) singular Jacobian. *Quant Econom.* 10:1019–1068. <https://doi.org/10.3982/QE989>.
- Huang K et al. 2019. Genotypic frequencies at equilibrium for polysomic inheritance under double-reduction. *G3 (Bethesda).* 9: 1693–1706. <https://doi.org/10.1534/g3.119.400132>.
- Kleibergen F. 2005. Testing parameters in GMM without assuming that they are identified. *Econometrica.* 73:1103–1123. <https://doi.org/10.1111/j.1468-0262.2005.00610.x>.
- Korneliussen TS, Albrechtsen A, Nielsen R. 2014. ANGSD: analysis of next generation sequencing data. *BMC Bioinformatics.* 15:356. <https://doi.org/10.1186/s12859-014-0356-4>.
- Kucherenko S, Sytsko Y. 2005. Application of deterministic low-discrepancy sequences in global optimization. *Comput Optim Appl.* 30:297–318. <https://doi.org/10.1007/s10589-005-4615-1>.
- Lau J et al. 2022. Rose rosette disease resistance loci detected in two interconnected tetraploid garden rose populations. *Front Plant Sci.* 13:916231. <https://doi.org/10.3389/fpls.2022.916231>.
- Leung D, Sturma N. 2024. Singularity-agnostic incomplete U-statistics for testing polynomial constraints in Gaussian covariance matrices [preprint], arXiv, arXiv:2401.02112. <https://doi.org/10.48550/arXiv.2401.02112>.
- Li H. 2011. A statistical framework for SNP calling, mutation discovery, association mapping and population genetical parameter estimation from sequencing data. *Bioinformatics.* 27:2987–2993. <https://doi.org/10.1093/bioinformatics/btr509>.
- Matsuoka Y. 2011. Evolution of polyploid *Triticum* wheats under cultivation: the role of domestication, natural hybridization and allopolyploid speciation in their diversification. *Plant Cell Physiol.* 52:750–764. <https://doi.org/10.1093/pcp/pcr018>.
- Mendel G. 1866. Versuche über pflanzenhybriden. In: *Verhandlungen des Naturforschenden Vereins Brünn.* Band 4. p. 3–47.
- Mitchell JD, Allman ES, Rhodes JA. 2019. Hypothesis testing near singularities and boundaries. *Electron J Stat.* 13:2150–2193. <https://doi.org/10.1214/19-EJS1576>.
- Mollinari M et al. 2019. Supplemental material for Mollinari et al., 2020. *GSA Journals. Dataset.* <https://doi.org/10.25387/g3.10255844.v1>.
- Mollinari M et al. 2020. Unraveling the hexaploid sweetpotato inheritance using ultra-dense multilocus mapping. *G3 (Bethesda).* 10: 281–292. <https://doi.org/10.1534/g3.119.400620>.
- Mollinari M, Garcia AAF. 2019. Linkage analysis and haplotype phasing in experimental autopolyploid populations with high ploidy level using hidden Markov models. *G3 (Bethesda).* 9:3297–3314. <https://doi.org/10.1534/g3.119.400378>.
- Parisod C, Holderegger R, Brochmann C. 2010. Evolutionary consequences of autopolyploidy. *New Phytol.* 186:5–17. <https://doi.org/10.1111/j.1469-8137.2009.03142.x>.
- Powell MJ. 2009. The BOBYQA algorithm for bound constrained optimization without derivatives. Technical Report DAMTP 2009/NA06. Centre for Mathematical Sciences, University of Cambridge, UK.
- R Core Team. 2025. R: A Language and Environment for Statistical Computing. R Foundation for Statistical Computing, Vienna, Austria.
- Rosyara UR, De Jong WS, Douches DS, Endelman JB. 2016. Software for genome-wide association studies in autopolyploids and its application to potato. *Plant Genome.* 9:1–10. <https://doi.org/10.3835/plantgenome2015.08.0073>.
- Schmittmann VD, Dolan CV, Rajmakers ME, Batchelder WH. 2010. Parameter identification in multinomial processing tree models. *Behav Res Methods.* 42:836–846. <https://doi.org/10.3758/BRM.42.3.836>.
- Schwarz G. 1978. Estimating the dimension of a model. *Ann Stat.* 6: 461–464. <https://doi.org/10.1214/aos/1176344136>.
- Self SG, Liang KY. 1987. Asymptotic properties of maximum likelihood estimators and likelihood ratio tests under nonstandard conditions. *J Am Stat Assoc.* 82:605–610. <https://doi.org/10.1080/01621459.1987.10478472>.
- Serang O, Mollinari M, Garcia AAF. 2012. Efficient exact maximum a posteriori computation for Bayesian SNP genotyping in polyploids. *PLoS One.* 7:1–13. <https://doi.org/10.1371/journal.pone.0030906>.
- Shirasawa K et al. 2017. A high-density SNP genetic map consisting of a complete set of homologous groups in autohexaploid sweetpotato (*Ipomoea batatas*). *Sci Rep.* 7:44207. <https://doi.org/10.1038/srep44207>.
- Soares NR, Mollinari M, Oliveira GK, Pereira GS, Vieira MLC. 2021. Meiosis in polyploids and implications for genetic mapping: a review. *Genes (Basel).* 12:1517. <https://doi.org/10.3390/genes12101517>.
- Soltis DE et al. 2009. Polyploidy and angiosperm diversification. *Am J Bot.* 96:336–348. <https://doi.org/10.3732/ajb.0800079>.
- Stift M, Berenos C, Kuperus P, van Tienderen PH. 2008. Segregation models for disomic, tetrasomic and intermediate inheritance in tetraploids: a general procedure applied to *Rorippa* (yellow cress) microsatellite data. *Genetics.* 179:2113–2123. <https://doi.org/10.1534/genetics.107.085027>.

- Stift M, Reeve R, Van Tienderen PH. 2010. Inheritance in tetraploid yeast revisited: segregation patterns and statistical power under different inheritance models. *J Evol Biol.* 23:1570–1578. <https://doi.org/10.1111/j.1420-9101.2010.02012.x>.
- Storey JD, Tibshirani R. 2003. Statistical significance for genomewide studies. *Proc Natl Acad Sci U S A.* 100:9440–9445. <https://doi.org/10.1073/pnas.1530509100>.
- Susko E. 2013. Likelihood ratio tests with boundary constraints using data-dependent degrees of freedom. *Biometrika.* 100:1019–1023. <https://doi.org/10.1093/biomet/ast032>.
- Udall JA, Wendel JF. 2006. Polyploidy and crop improvement. *Crop Sci.* 46:S–3. <https://doi.org/10.2135/cropsci2006.07.0489tpg>.
- Viallefont A, Lebreton J, Reboulet AM, Gory G. 1998. Parameter identifiability and model selection in capture-recapture models: a numerical approach. *Biom J: J Math Methods Biosci.* 40:313–325. [https://doi.org/10.1002/\(ISSN\)1521-4036](https://doi.org/10.1002/(ISSN)1521-4036).
- Voorrips RE, Maliepaard CA. 2012. The simulation of meiosis in diploid and tetraploid organisms using various genetic models. *BMC Bioinformatics.* 13:248. <https://doi.org/10.1186/1471-2105-13-248>.

Editor: J. Birchler