

Linkage Analysis and Haplotype Phasing in Experimental Autopolyploid Populations with High Ploidy Level Using Hidden Markov Models

Marcelo Mollinari* and Antonio Augusto Franco Garcia^{†,1}

*Department of Horticultural Science, Bioinformatics Research Center, North Carolina State University, Raleigh, North Carolina, and [†]Department of Genetics, University of São Paulo/ESALQ, Piracicaba, São Paulo, Brazil

ORCID IDs: 0000-0002-7001-8498 (M.M.); 0000-0003-0634-3277 (A.A.F.G.)

ABSTRACT Modern SNP genotyping technologies allow measurement of the relative abundance of different alleles for a given locus and consequently estimation of their allele dosage, opening a new road for genetic studies in autopolyploids. Despite advances in genetic linkage analysis in autotetraploids, there is a lack of statistical models to perform linkage analysis in organisms with higher ploidy levels. In this paper, we present a statistical method to estimate recombination fractions and infer linkage phases in full-sib populations of autopolyploid species with even ploidy levels for a set of SNP markers using hidden Markov models. Our method uses efficient two-point procedures to reduce the search space for the best linkage phase configuration and reestimate the final parameters by maximizing the likelihood of the Markov chain. To evaluate the method, and demonstrate its properties, we rely on simulations of autotetraploid, autohexaploid and autooctaploid populations and on a real tetraploid potato data set. The results show the reliability of our approach, including situations with complex linkage phase scenarios in hexaploid and octaploid populations.

KEYWORDS

Polyploidy
Recombination
Fraction
Bivalent Pairing
Multilocus
Analysis

Polyploids are organisms with more than two sets of chromosomes. They are very important in agriculture and play a fundamental role in evolutionary processes, such as differentiation of species (Soltis *et al.* 2014a). The number of sets of chromosomes in an organism is called *ploidy level*. These multiple chromosome sets can originate from the combination of genomes from different, but related species, or from duplicated genomes from the same species (Birchler 2012; Comai 2005). In the first scenario, they are called *allopolyploids*; in the second, *autopolyploids*. Polyploid organisms are also characterized according to their pattern of inheritance. In general, allopolyploids exhibit diploid-like (or *disomic*) segregation, since homologous chromosomes, or homologs, tend to form bivalents within each

sub-genome. Autopolyploids, however, have more than two homologs per homology group, forming either random bivalents or multivalents during meiosis, resulting in *polysomic segregation* (Sybenga 1975; Soltis *et al.* 1993; Osborn *et al.* 2003). Since the molecular mechanics of polyploid organisms are quite complex, this dichotomy is often broken, and polyploids can display intermediate modes of inheritance (Otto and Whitton 2000; Osborn *et al.* 2003). Throughout this paper, the term autopolyploid (or autotetraploid, autohexaploid, etc.) will refer to polyploid organisms that exhibit polysomic segregation.

Despite advances in genetic studies in autotetraploids, (Mather 1936; Fisher 1943, 1947; Hackett *et al.* 2001; Luo *et al.* 2004, 2006; Wu *et al.* 2004; Leach *et al.* 2010; Li *et al.* 2010; Hackett *et al.* 2013; Xu *et al.* 2013; Rehmsmeier 2013; Zheng *et al.* 2016), there is still a shortage of statistical methods to address organisms with higher ploidy levels, such as sweet potato (Kriegner *et al.* 2003; Arizio *et al.* 2014; Shirasawa *et al.* 2017), sugarcane (Wang *et al.* 2010; Garcia *et al.* 2013), some ornamental flowers and forage crops (reviewed in (Soltis *et al.* 2014b)). In this work, we denote as *high-level autopolyploids* those autopolyploid organisms with ploidy level greater than four. A fundamental class of statistical methods that have lagged behind in high-level autopolyploid studies is the construction of genetic maps. A reliable genetic map is a crucial step in quantitative trait loci (QTL) analysis,

Copyright © 2019 Mollinari, Garcia

doi: <https://doi.org/10.1534/g3.119.400378>

Manuscript received June 3, 2019; accepted for publication August 8, 2019; published Early Online August 12, 2019.

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

Supplemental material available at FigShare: <https://doi.org/10.25387/g3.8218325>.

¹Corresponding author: Av. Pádua Dias, 11. CEP 13418-900 - Piracicaba, SP, Brazil.

E-mail: augusto.garcia@usp.br

as well as the assembly of reference genomes and the study of evolutionary processes (Lewin *et al.* 2009; Luo *et al.* 2013; Lemmon and Doebley 2014). Although understanding the concept of genetic mapping is rather easy, the construction of such maps in high-level autopolyploids is challenging. Even under bivalent pairing, there are many possible configurations during meiosis, and the number of possibilities gets exponentially larger as the ploidy level increases. Denoting m as the ploidy level, it is possible to find up to m different alleles for a locus in one individual. Furthermore, if some of those alleles are not distinguishable, it is necessary to consider the number of copies of each different allelic form, also known as *dosage*.

The construction of a genetic map in a full-sib population can be summarized in five basic steps: *i*) estimation of pairwise recombination fractions and associated statistical tests; *ii*) separation of markers into linkage groups; *iii*) ordering of markers within each linkage group using an optimization technique; *iv*) parental phasing, recombination fraction updating and likelihood computation (or other objective function) and *v*) if the order is optimal, the map is complete, otherwise, return to step *iii*. Historically, genetic maps in high-level autopolyploids have been constructed using only alleles present in one homolog, called *single-dose* or *simplex* markers (Wu *et al.* 1992; Sorrells 1992). In a full-sib population, these markers segregate in a 1:1 ratio (if they are present only in one parent), or in a 1:2:1 ratio (if present in both parents, also called *double simplex*). Given this level of simplification, it is possible to use the five-step procedure coupled with a standard software suitable for diploid populations. Nevertheless, it is well accepted that the use of single-dose markers imposes limitations on the construction of adequate genetic maps. These approaches sub-sample the genome (Hackett *et al.* 2013; Garcia *et al.* 2013), which precludes further consideration of multiallelic effects in models for QTL mapping and subsequent studies. Moreover, there is low statistical power to detect linkage when markers are in repulsion phase configurations (Wu *et al.* 1992; Ripol *et al.* 1999). Although some authors have addressed this problem by including multi-dose (or multiplex) markers when constructing genetic maps and performing QTL mapping (Ripol *et al.* 1999; Doerge and Craig 2000), the limitations on the genotyping technologies at the time required that the allelic dosage had to be inferred based on expected segregation rates. Because of the high amount of hidden information imposed by marker systems on those studies (Wu *et al.* 1992; Ripol *et al.* 1999), the estimation of recombination fraction between multi-dose markers was highly impaired.

Quantitative genotyping technologies for single nucleotide polymorphism (SNPs) evaluation have opened the door for further genetic mapping studies in high-level autopolyploids. It is now possible to measure the abundance of specific alleles within a locus in a polyploid genome (Voorrips *et al.* 2011; Serang *et al.* 2012; Hackett *et al.* 2013; Garcia *et al.* 2013; Bargary *et al.* 2014; Mollinari and Serang 2015). This technology, combined with the genotypic distribution in the population, makes it possible to infer the allelic dosage by using the ratio between the abundances of the two alternative alleles (Serang *et al.* 2012). Once the dosage of the markers is estimated, the construction of linkage maps can be significantly improved by taking this information into account, as previously done in autotetraploids by Hackett *et al.* (2013, 2014).

Genetic linkage maps can be constructed based on two-point or multipoint estimates of the recombination fraction. Two-point methods use information on pairs of markers, and even though they are less computationally demanding than multipoint methods, they require a higher amount of information in the markers to provide reliable results. Only recently, using a two-point-based method, van Geest *et al.* (2017) published an integrated hexaploid chrysanthemum genetic map using scripts implemented in the R package

polymapR (Bourke *et al.* 2018). Multipoint approaches, instead, use information of multiple markers present in a linkage group, increasing the statistical efficiency of the analysis (Lander and Green 1987; Jiang and Zeng 1997; Mollinari *et al.* 2009; Leach *et al.* 2010). This feature is particularly important in polyploid linkage analysis, where markers are mostly partially informative. One widely used procedure to obtain multipoint estimates is the hidden Markov model (HMM) (Lander and Green 1987). The construction of the genetic map using this method provides the estimates of the recombination fractions between all adjacent markers in a linkage group, as well as the multipoint likelihood, which has been shown to be an excellent criterion to evaluate and compare linkage phase configurations and orders of markers (Mollinari *et al.* 2009). Leach *et al.* (2010) presented a statistical framework in which HMMs were applied to reconstruct genetic linkage maps, but it was limited to autotetraploids. Recently, software packages such as polymapR (Bourke *et al.* 2018), pergola (Grandke *et al.* 2017) and netgwas (Behrouzi and Wit 2017), have been developed to build genetic maps in high-level autopolyploids. However, only polymapR is capable of estimating recombination fractions and inferring parental linkage phases in outcrossing populations, though it does not use multipoint procedures to perform those tasks.

The main challenges we address in this paper are the inference of the haplotypes of the multiple homologs and the multipoint estimation of recombination fractions in high-level autopolyploids. Although Zheng *et al.* (2016) proposed a probabilistic multilocus haplotype reconstruction model for autotetraploids considering double reduction, this remains as an open question for organisms with higher ploidy levels. Our method relies on an HMM and is developed for species with even ploidy levels under random chromosome segregation (complete polysomic inheritance). We also present a two-point method which is capable of dealing with hundreds of markers even in high ploidy level scenarios. Hence, we are proposing solutions for steps *i* and *iv* in high-level autopolyploids. Step *ii* is straightforward from step *i* using clustering algorithms, as proposed by (Van Ooijen and Jansen 2013). Even though step *iii* is a challenging task in genetic mapping, it can be addressed using pairwise recombination fractions or the resulting likelihood of the Markov model as it has been proposed by several studies (Lander *et al.* 1987; Buetow and Chakravarti 1987; Doerge 1996; Van Os *et al.* 2005; Wu *et al.* 2008; Preedy and Hackett 2016; Wang *et al.* 2016). To evaluate our method, and to show its properties, we rely on simulations of autotetraploid, autohexaploid, and autooctaploid data and on a real tetraploid potato data set. We also perform a set of hexaploid simulations to compare our method to the one implemented in polymapR (Bourke *et al.* 2018). The R computer codes to reproduce all simulations and analysis are publicly available.

MATERIALS AND METHODS

In this section, we define the notation used throughout this article and present the probabilistic model for the gamete formation in autopolyploids. The mathematical derivation of the HMM, including the estimation of the model parameters, is based on the work of Rabiner (1989), which presents the hidden Markovian process using three major elements, namely, the *transition probability* function (Equation 6), the *initial state* function (Equation 7), the *emission probability* function (Equations 8 and 9). In the genetic mapping context, the first connects adjacent marker loci in function of their recombination fraction, the second is the prior probability of the genotypes in the mapping population, and the third connects the observed marker dosage to the complete multi-allelic hidden genotypic states. While these ideas are widespread in the genetic mapping literature, for instance in Lander

and Green (1987); Jiang and Zeng (1997); Hackett and Broadfoot (2003); Hackett *et al.* (2013); Leach *et al.* (2010), here we present a generalization to any even ploidy level. We avoid using matrix notation throughout our model derivation since their high dimensionality would precludes the application of HMM techniques in polyploids with high ploidy level. We conclude this section by explaining the complexity of estimating linkage phases between markers, presenting an efficient two-point algorithm that simplifies the problem in a way that allows the phasing to be inferred using real data.

Notation

Consider a mapping population derived from a cross between two autopolyploid individuals P and Q with the same ploidy level (full-sib family). The ploidy level is denoted by m , and can be any even number greater than zero. Let the vectors $\mathcal{P}_k^m = \{P_k^i\}$ and $\mathcal{P}_{k+1}^m = \{P_{k+1}^i\}$, and $\mathcal{Q}_k^m = \{Q_k^i\}$ and $\mathcal{Q}_{k+1}^m = \{Q_{k+1}^i\}$, $i = 1, \dots, m$, denote the genotype of two adjacent multiallelic loci k and $k + 1$ in P and Q , respectively. The superscript i indicates one of the possible alleles for the loci, and each locus has m different alleles in each parent. For example, for a cross between two autohexaploid individuals, $\mathcal{P}_k^6 = \{P_k^1, P_k^2, \dots, P_k^6\}$; similarly, this can be done for $\mathcal{P}_{k+1}^6$, $\mathcal{Q}_k^6$ and $\mathcal{Q}_{k+1}^6$. All alleles denoted by the same superscript number are in the same homologs (e.g., P_k^1 and P_{k+1}^1 are in homolog 1, etc). The following assumptions are made to ensure random chromosome segregation (Muller 1914; Haldane 1930) and no double reduction (Burnham 1962): *i*) there is only formation of bivalents during the meiosis; *ii*) there is no preferential pairing during the formation of bivalents; *iii*) all bivalents have the same recombination fraction between loci k and $k + 1$; *iv*) bivalents are independent; *v*) there is separation of sister chromatids during meiosis II and *vi*) there is no chromatid interference. Consequences of violations of these assumptions will be addressed later using simulations. Although each bivalent is composed by a pair of chromosomes with two sister chromatids each, given assumption *vi*, we will consider only one sister chromatid per homolog during the derivation of our model.

Bivalent formation

Bivalent formation occurs during meiosis I (more specifically, at the pachytene stage of prophase). In diploid cells, there is only one possible pairing configuration: two duplicated homologs from a homology group pair to form one bivalent. However, in autopolyploid cells, given the previous assumptions, the expected number of possible pairing configurations, *i.e.*, the number of possible bivalent chromosomal pairings for a given homology group during meiosis can be obtained by sequentially choosing pairs out of m homologs without replacement, divided by all possible permutations of the chosen pairs

$$w_m = \frac{1}{\frac{m!}{2}} \prod_{i=1}^{\frac{m}{2}} \binom{2i}{2} \quad (1)$$

The orientation of the bivalents does not affect the expected frequencies of each gamete type, and therefore will not be considered. For example, as showed by Hackett (2001) in autotetraploids, there are two bivalents and three possible bivalent configurations: homolog pair as 1 with 2, and 3 with 4; or, 1 with 3 and 2 with 4; or 1 with 4 and 2 with 3. We denote $\Psi = \{\psi_j\}$, $j = 1, \dots, w_m$ a set of all bivalent configurations for a given ploidy level.

Expected gametic frequency for a given bivalent configuration

We will present the expected gametic frequencies considering parent P . Since parent Q undergoes a similar process, it is possible to combine

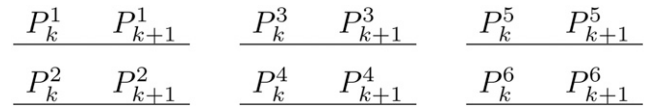


Figure 1 One possible pairing configuration in an autohexaploid, namely ψ_1 . P_k^i denotes one allele present in homolog i for locus k in parent P . Notice that some allelic configurations, such as $(\{P_k^1, P_k^2, P_k^5\}, \{P_{k+1}^1, P_{k+1}^2, P_{k+1}^5\})$, are impossible to be obtained in this bivalent pairing. In this case, the homologs containing alleles P_k^1 and P_k^2 will migrate to opposite poles of the cell during meiosis I. Therefore, P_k^1 and P_k^2 will not be present in the same gamete.

the expected gametic frequencies to obtain the expected genotypic frequency in the full-sib population. Each of the bivalents obtained for a given configuration ψ_j can result in two types of chromosomes for loci k and $k + 1$: *parental*, which results from bivalents with zero or any other even number of recombinations between k and $k + 1$; and *recombinants*, which results from bivalents with any odd number of recombinations. As presented by Doerge and Craig (2000), the probabilities of all chromosome types for any single bivalent can be represented always as

$$\mathbf{V} = \begin{bmatrix} \Pr(P_k^i, P_{k+1}^i) & \Pr(P_k^i, P_{k+1}^{i'}) \\ \Pr(P_k^{i'}, P_{k+1}^i) & \Pr(P_k^{i'}, P_{k+1}^{i'}) \end{bmatrix} = \begin{bmatrix} \frac{1-r_k}{2} & \frac{r_k}{2} \\ \frac{r_k}{2} & \frac{1-r_k}{2} \end{bmatrix}$$

where r_k is the recombination fraction between k and $k + 1$, $i \neq i'$. For a given configuration ψ_j , the expected frequencies for all possible gametes derived from that configuration is

$$\mathbf{V}_1 \otimes \dots \otimes \mathbf{V}_{\frac{m}{2}}$$

where $\otimes$ denotes the Kronecker product of matrices and subscripts in $\mathbf{V}$ indicate the corresponding bivalent. All elements of this product are of the form

$$\frac{(1-r_k)^{\frac{m}{2}-l}(r_k)^l}{2^{\frac{m}{2}}}$$

where l denotes the number of total recombinant bivalents between loci k and $k + 1$, $l \in \{0, \dots, m/2\}$. From this, we can define the probability of observing any gamete (for two loci) given a bivalent configuration ψ_j as

$$\Pr(\mathbf{p}_k, \mathbf{p}_{k+1} | \psi_j) = \begin{cases} \frac{(1-r_k)^{\frac{m}{2}-l}(r_k)^l}{2^{\frac{m}{2}}} & \text{if } \psi_j \text{ is consistent with } \{\mathbf{p}_k, \mathbf{p}_{k+1}\} \\ 0 & \text{otherwise} \end{cases} \quad (2)$$

where vectors $\mathbf{p}_k$ and $\mathbf{p}_{k+1}$ denote a subset of $\frac{m}{2}$ alleles present in $\mathcal{P}_k^m$ and $\mathcal{P}_{k+1}^m$, respectively; $\{\mathbf{p}_k, \mathbf{p}_{k+1}\}$ indicates a gamete for loci k and $k + 1$ from parent P . *Consistent* means that the gamete can be produced from bivalent configuration ψ_j . Notice that some gametes cannot be obtained from ψ_j once the bivalents are formed.

Since we assume that alleles with the same superscript are in the same homolog, l can be obtained by a simple examination of superscripts of elements contained in $\mathbf{p}_k$ and $\mathbf{p}_{k+1}$. Consider, for example, $\psi_1 = \{(1, 2), (3, 4), (5, 6)\}$ ($m = 6$, Figure 1). If one observes $\mathbf{p}_k = \{P_k^1, P_k^3, P_k^5\}$ and $\mathbf{p}_{k+1} = \{P_{k+1}^1, P_{k+1}^4, P_{k+1}^6\}$, the number of recombinant chromosomes is $l = 2$. Therefore, $\Pr(\{P_k^1, P_k^3, P_k^5\}, \{P_{k+1}^1, P_{k+1}^4, P_{k+1}^6\} | \psi_1) = \frac{(1-r_k)^2(r_k)^2}{2^3}$. On the other hand, $\Pr(\{P_k^1, P_k^2, P_k^5\}, \{P_{k+1}^1, P_{k+1}^2, P_{k+1}^5\} | \psi_1) = 0$, since it is impossible

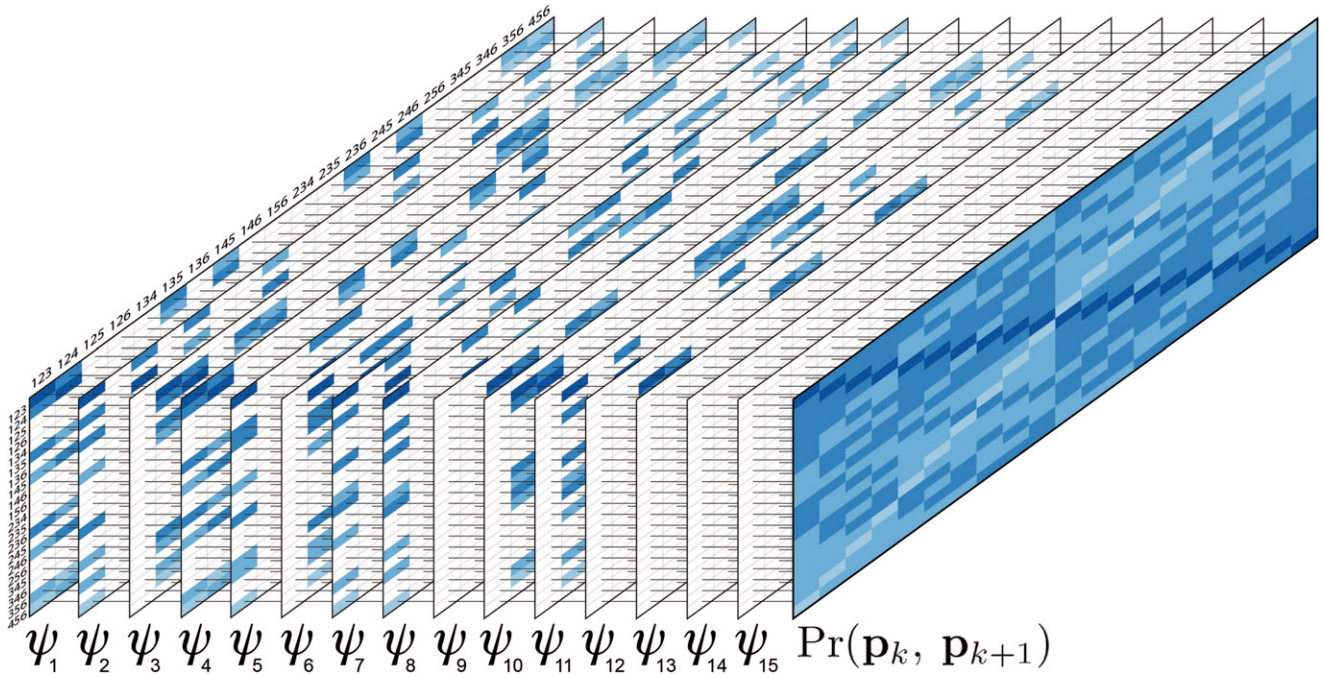


Figure 2 Graphical representation of Equations 2 and 3 for autohexaploid gametes. The first 15 tables represent the gametic probabilities given different bivalent configurations ψ . (Equation 2). The rows and the columns indicate gametic configurations for loci k and $k+1$, respectively. For simplification, only the superscripts of the gametic configurations were presented. For example, row 123, column 123, represent the gamete $(\{P_k^1, P_k^2, P_k^3\}; \{P_{k+1}^1, P_{k+1}^2, P_{k+1}^3\})$. Colored cells indicate the probability of gametic configurations consistent with the bivalent configuration ψ . The color scale indicates the number of recombinant bivalents associated to the gametic probability varying from 0 (dark blue) to 3 (light blue). Blank cells indicate non-consistent configurations. The far right full table represents the sum over all ψ configurations, weighted by their probability (Equation 3).

to obtain this gamete from configuration ψ_1 , i.e., it is not consistent with ψ_1 .

Expected gametic frequency unconditional to bivalent configurations

In reality ψ_j is unknown, thus the conditional probability given by Equation (2) must be considered for all possible ψ_j . The probability of observing a gamete $\{p_k, p_{k+1}\}$, unconditional to ψ_j , can be expressed as

$$\Pr(p_k, p_{k+1}) = \sum_{j=1}^{w_m} \Pr(p_k, p_{k+1} | \psi_j) \Pr(\psi_j) \quad (3)$$

It is important to notice that only a subset of Ψ is consistent with the observed gamete, and consequently $\Pr(p_k, p_{k+1} | \psi_j) > 0$ only for some ψ_j 's. Figure 2 shows a graphical representation of Equations 2 and 3 for autohexaploid gametes.

The probability of observing a specific gamete is always the same for each ψ_j in this consistent subset (Equation 2). Therefore, under random pairing (assumption ii), our task reduces to finding the number of elements in this subset that are consistent with the observed gamete and multiply $\Pr(p_k, p_{k+1} | \psi_j) \Pr(\psi_j)$ by this number. The result is the probability of observing a gamete unconditional to the bivalent configuration.

For every gamete, l can change from zero to $m/2$ recombinant homologs. The observed gamete is the result of homologs that migrate to one pole of the cell at anaphase I with a subsequent migration to opposite poles at anaphase II. Since we are assuming that there is separation of sister chromatids during anaphase II, if $l = 0$ (all chromosomes are of parental type), there is no information about the pairing configuration of the homologs that migrate to the opposite pole of the cell. In this situation, there are $\left(\frac{m}{2}\right)!$ possible pairing

configurations, and the number of possible ψ_j that can produce gametes with $l = 0$ is $\left(\frac{m}{2}\right)!$. Therefore, for $l > 0$, there are $\left(\frac{m}{2} - l\right)!$ possible pairing configurations of parental chromosomes. For the remaining l recombinant chromosomes, the number of possible pairing configurations is $l!$. Thus, the total number of possible pairing configurations that can produce a specific gamete is $l! \left(\frac{m}{2} - l\right)!$. This is precisely the number of elements in the subset of Ψ consistent with the observed gamete. Given the assumption of no preferential pairing during the formation of bivalents, $\Pr(\psi_j) = \frac{1}{w_m}$, the probability of a gamete $\{p_k, p_{k+1}\}$, unconditional to ψ_j , can be simplified to

$$\Pr(p_k, p_{k+1}) = \frac{l! \left(\frac{m}{2} - l\right)! (1 - r_k)^{\frac{m}{2} - l} (r_k)^l}{w_m 2^{\frac{m}{2}}} \quad (4)$$

Map reconstruction via hidden Markov model

The construction of a genetic map involves the estimation of the genetic distance and order between markers within linkage groups. If the origin of the haplotypes (i.e., linkage phase) for the parents of the mapping population is unknown, it also needs to be estimated. For several years, hidden Markov models have been proven to be an excellent avenue for obtaining these estimates (Lander and Green 1987; Jiang and Zeng 1997; Mollinari *et al.* 2009; Leach *et al.* 2010). The multipoint likelihood obtained using HMMs is employable as a criterion to compare marker orders and judge which one is best, and also to provide a reliable estimation of recombination fraction and linkage phases. (Rabiner 1989) defines an HMM as a generative process composed of three well-defined probability distributions: *transition*, *initial state* and *emission*. In genetic mapping context, the transition probability distribution is defined as the probability of having a particular genotype

at position $k + 1$, given the genotype at position k . Using Equation (4) the gametic transition probabilities $\Pr(\mathbf{p}_{k+1}|\mathbf{p}_k)$, or the conditional probability of a gamete genotype at locus $k + 1$ given the gamete genotype at locus k , is simply

$$\Pr(\mathbf{p}_{k+1}|\mathbf{p}_k) = \frac{\Pr(\mathbf{p}_k, \mathbf{p}_{k+1})}{\Pr(\mathbf{p}_k)}$$

Under random chromosome segregation, both $\mathbf{p}_k$ and $\mathbf{p}_{k+1}$ can have $\binom{m}{2}$ different genotypes. Let $\Theta_p^m = \{\theta_{p,i}^m\}, i = 1, \dots, \binom{m}{2}$ denote all possible genotypes that $\mathbf{p}_k$ can assume for locus k . Also, assume that genotypes in Θ_p^m are arranged according to the lexicographical order of their superscripts. For example, in an autotetraploid, $\Theta_p^4 = \{(P_k^1, P_k^2), (P_k^1, P_k^3), (P_k^1, P_k^4), (P_k^2, P_k^3), (P_k^2, P_k^4), (P_k^3, P_k^4)\}$ for locus k . After some simplifications (see Supplementary Information, File S1) the transition probability, *i.e.*, the conditional probability of a gametic genotype $\theta_{p,i'}^m$ in locus $k + 1$ given the gametic genotype $\theta_{p,i}^m$ in locus k , is

$$\Pr(\mathbf{p}_{k+1} = \theta_{p,i'}^m | \mathbf{p}_k = \theta_{p,i}^m) = \frac{(1-r_k)^{\frac{m}{2}-l}(r_k)^l}{\binom{m}{2}} \quad (5)$$

where $i, i' \in \left\{1, \dots, \binom{m}{2}\right\}$. The initial state and the emission probability distributions will be addressed in the next section (Equations 7 to 9).

Including information from both parents

Any given individual in a full-sib population is formed by the union of gametes from both parents, P and Q . Each parent can form $\binom{m}{2}$ different gametes for locus k . Since the formation of gametes in both parents is independent, the genotypic transition probability distribution can be written as

$$\begin{aligned} \Pr(\mathcal{G}_{k+1,j'}^m | \mathcal{G}_{k,j}^m) &= \Pr(\mathbf{p}_{k+1} = \theta_{p,i'}^m | \mathbf{p}_k = \theta_{p,i}^m) \Pr(\mathbf{q}_{k+1} = \theta_{q,h'}^m | \mathbf{q}_k = \theta_{q,h}^m) \\ &= \frac{(1-r_k)^{m-l_p-l_q}(r_k)^{l_p+l_q}}{\binom{m}{2} \binom{m}{2}} \end{aligned} \quad (6)$$

where $\mathcal{G}_{k,j}^m$ denotes the genotype of an individual derived from the union of gametes $\theta_{p,i}^m$ and $\theta_{q,h}^m$ at locus k . The same reasoning applies to $\mathcal{G}_{k+1,j'}^m$; $i, i', h, h' \in \left\{1, \dots, \binom{m}{2}\right\}$, $j = (i-1)\binom{m}{2} + h$ and $j' = (i'-1)\binom{m}{2} + h'$. l_p and l_q denote the number of recombinant bivalents between loci k and $k + 1$ in parents P and Q , respectively. Let $g_m = \binom{m}{2}$ denote the number of possible genotypes derived from the cross between individuals P and Q . For simplification and without loss of generality, let $t_k(j, j') = \Pr(\mathcal{G}_{k+1,j'}^m | \mathcal{G}_{k,j}^m)$. For a comprehensive autotetraploid example of the transition probabilities

(similar to that presented by Hackett (2001)) and the indexation used in Equation 6, see Table S3.8, File S3 in Supplementary Information.

Given a ploidy level m and a recombination fraction r_k , the only information required to obtain $t_k(j, j')$ in Equation (6) is l_p and l_q . Since the genotypes in Θ_p^m and Θ_q^m are arranged according to the lexicographical order of their superscripts, it is possible to obtain (l_p, l_q) for any given pair (j, j') using the algorithm presented in Supplementary Information, File S2. Although the number of possible transitions between positions k and $k + 1$ is $(g_m)^2$, which can be a very large number even for modest ploidy levels, it is possible to obtain the transition between any specific genotypes in j and j' without computing the entirety of the transition space.

The initial state distribution is the probability of observing a specific genotype. Given the assumption that there is no preferential pairing during the formation of bivalents, a uniform probability density function can be employed as the initial state probability function

$$\gamma_j = \Pr(\mathcal{G}_{1,j}^m) = \frac{1}{g_m}, j \in \{1, \dots, g_m\} \quad (7)$$

To this point, both transition and initial state distributions consider different allelic variants for all m homologs in both parents. This scenario can only be achieved when using fully informative markers. In reality, autopolyploid species may have the same allelic variant in some homologs. Besides, even if a particular locus have different allelic form in all homologs, modern genotyping platforms are usually capable of detecting polymorphisms at the nucleotide level (SNPs), which are essentially biallelic. Due to this lack of identity between the observed data and the full transition space, we make use of the emission function, which is defined as the probability of observing a molecular phenotype given a genotype $\mathcal{G}_{k,j}^m$.

The detection of the allelic variants in modern genotyping platforms is based on the abundance of different alternative nucleotides. In the autopolyploid setting, this can be translated as the *dosage* of a SNP at a specific locus. The dosage of a SNP can be estimated using the ratio between the abundance of its two allelic forms. Several methods were proposed to perform this task including (Voorrips *et al.* 2011), (Serang *et al.* 2012) and (Bargary *et al.* 2014). Here we introduce a biallelic derivation of the emission probability distribution. Although the function presented here use biallelic information, other distributions can be derived for partial informative multiallelic marker systems following the same reasoning.

Let $d_p^k, d_q^k \in \{0, \dots, m\}$ denote the *observed dosage* of one allelic form in locus k for parents P and Q , respectively. The choice of the allelic form denoted by d_p^k is arbitrary, as long as the same allelic form is used in d_q^k . The dosage observed in parent P can be originated from alleles present in d_p^k of the m homologs. Let $\phi_p^k = \{\varphi_p^k : \varphi_p^k \subseteq \mathcal{P}_k^m, \#\{\varphi_p^k\} = d_p^k\}$ denote a set of size $\binom{m}{d_p^k}$ containing all possible subsets in $\mathcal{P}_k^m$ that originate the observed dosage d_p^k . The operator $\#\{\cdot\}$ is the cardinality of a set. The same reasoning applies for ϕ_q^k . For instance, in an autotetraploid, if $d_p^k = 3$, the three doses present in locus k can be derived from four distinct subsets $\phi_p^k = \{(P_k^1, P_k^2, P_k^3), (P_k^1, P_k^2, P_k^4), (P_k^1, P_k^3, P_k^4), (P_k^2, P_k^3, P_k^4)\}$. Given two particular subsets φ_p^k and φ_q^k in ϕ_p^k and ϕ_q^k , each one of the g_m genotypic states in the full transition space can be associated to a dosage. The observed dosage O associated to the j -th state is obtained by counting the number of alleles present in the intersection between the parental allelic set $(\varphi_p^k \cup \varphi_q^k)$ and $\mathcal{G}_{k,j}^m$. Thus, the emission function can be defined as

$$b_j(O) = \Pr(O | \mathcal{G}_{k,j}^m, \varphi_P^k, \varphi_Q^k) = \begin{cases} 1 - \epsilon & \text{if } O = \delta(k, j) \\ \frac{\epsilon}{m} & \text{otherwise} \end{cases} \quad (8)$$

where $\delta(k, j) = |(\varphi_P^k \cup \varphi_Q^k) \cap \mathcal{G}_{k,j}^m|$ and ϵ denotes the global genotyping error rate. In addition to the point estimate of the dosage, the genotyping calling methods cited above also provide the probability distribution of the dosages for a particular marker for all individuals of the biparental population. If this information is available, a more general emission function can be derived. Instead of modeling a global error rate ϵ , we use the prior information provided by the genotyping calling procedure. Let $\pi_k = \{\pi_i^k\}_{(1 \times m+1)}$ denote the probability distribution vector associated to the dosages $0, \dots, m$ at position k for a particular individual in the biparental population. For example, $\pi_k = \{0, \frac{1}{6}, \frac{2}{3}, \frac{1}{6}, 0\}$ denotes a tetraploid individual with probabilities $\frac{1}{6}, \frac{2}{3}$ and $\frac{1}{6}$ of having one, two and three doses, respectively, and zero for the remaining ones. Then, the emission probability function can be written as

$$b_j(O) = \Pr(O | \mathcal{G}_{k,j}^m, \varphi_P^k, \varphi_Q^k, \pi_k) = \pi_{\delta(k,j)+1}^k \quad (9)$$

In this case, the observation O can be any dosage from 0 to m and the information about the genotypes will be contained in the probability distribution of the dosages π_k . Thus, the probability of observing any dosage given a genotype $\mathcal{G}_{k,j}^m$ associated to a particular dosage $\delta(k, j)$ can be obtained by simply assessing the corresponding value in the probability distribution provided by the genotype calling procedure. Notice that Equation 8 can be reduced to Equation 9 using the appropriate π_k . For example, in autotetraploids, when the observed dosage for locus k is one, $O = 1$, $\pi_k = \{\frac{\epsilon}{m}, 1 - \epsilon, \frac{\epsilon}{m}, \frac{\epsilon}{m}, \frac{\epsilon}{m}\}$. Moreover, for missing values, it is possible to use the probability distribution of the genotypic classes under polysomic segregation, as presented by Serang *et al.* (2012).

Multipoint likelihood and the estimation of recombination fraction

Suppose there are z markers in a homology group in a known order represented by $M_1, \dots, M_k, \dots, M_z$. Let $\mathbf{r} = (r_1, \dots, r_k, \dots, r_{z-1})$ denote the recombination fraction vector between all marker intervals in this sequence. Also, assume linkage phase configurations in parents P and Q denoted respectively by $\Phi_P = (\varphi_P^1, \dots, \varphi_P^k, \dots, \varphi_P^z)$ and $\Phi_Q = (\varphi_Q^1, \dots, \varphi_Q^k, \dots, \varphi_Q^z)$. The sequence of observations for the z markers is denoted by $(O_1, \dots, O_k, \dots, O_z)$ and its underlying probability distributions are denoted by $\Pi = (\pi_1, \dots, \pi_k, \dots, \pi_z)$. The likelihood of $M_1, \dots, M_k, \dots, M_z$ can be obtained using Equations (6), (7) and (9) following the classical *forward procedure* (Rabiner 1989). Let $\alpha_k(j) = \Pr(O_1, \dots, O_k | \mathcal{G}_{k,j}^m, \mathbf{r}, \Phi_P, \Phi_Q, \Pi)$ denote the probability of the partial observation sequence $(O_1, \dots, O_k)$ and genotype $\mathcal{G}_{k,j}^m, j \in \{1, \dots, g_m\}$ given the sequence of recombination fractions $\mathbf{r}$, the linkage phase configurations Φ_P and Φ_Q and the probability distributions for the sequence of observations Π . The forward procedure follows the steps below:

1. Initialization:

$$\alpha_1(j) = \gamma_j b_j(O_1), \quad j = 1, \dots, g_m \quad (10)$$

2. Induction:

$$\alpha_{k+1}(j') = \left[\sum_j \alpha_k(j) t_k(j, j') \right] b_{j'}(O_{k+1}) \quad (11)$$

where $k = 1, \dots, z-1$ and $j' = 1, \dots, g_m$

3. Termination:

$$\Pr(O_1, \dots, O_z | \mathbf{r}, \Phi_P, \Phi_Q, \Pi) = \sum_{j=1}^{g_m} \alpha_z(j) \quad (12)$$

Then, the likelihood of the model is defined as

$$\prod_{i=1}^n \Pr(O_{1,i}, \dots, O_{z,i} | \mathbf{r}, \Phi_P, \Phi_Q, \Pi_i) \quad (13)$$

where n is the number of individuals in the full-sib population, $O_{1,i}, \dots, O_{z,i}$ is the sequence of marker observations for individual i and Π_i is a $(m+1) \times z$ matrix where the k -th column denotes the probability distributions associated to the marker M_k , individual i . The multipoint maximum likelihood estimate of $\mathbf{r}$ can be obtained using the *forward-backward* procedure coupled with the EM algorithm (Rabiner 1989). For the backward procedure, consider the variable $\beta_k(j) = \Pr(O_{k+1}, \dots, O_z | \mathcal{G}_{k,j}^m, \mathbf{r}, \Phi_P, \Phi_Q, \Pi)$ as the probability of the partial observation sequence from $k+1$ to z , given the genotype $\mathcal{G}_{k,j}^m$, the recombination fraction vector $\mathbf{r}$, the linkage phase configurations Φ_P and Φ_Q and the probability distributions for the sequence of observations Π . The solution to $\beta_k(j)$ was also described by (Rabiner 1989) as follows:

1. Initialization:

$$\beta_z(j) = 1, \quad j = 1, \dots, g_m \quad (14)$$

2. Induction:

$$\beta_k(j) = \sum_{j'}^{g_m} t_k(j, j') b_{j'}(O_{k+1}) \beta_{k+1}(j') \quad (15)$$

where $k = z-1, z-2, \dots, 1$ and $j = 1, \dots, g_m$

To estimate the recombination fraction for all intervals in the marker sequence we need to define $\xi_k(j, j')$ as the probability of state $\mathcal{G}_{k,j}^m$ at position k and state $\mathcal{G}_{k+1,j'}^m$ at position $k+1$ given the sequence of observations $O_1, \dots, O_z$ and their underlying probability distributions Π , the recombination fraction vector $\mathbf{r}$ and the linkage phase configurations Φ_P and Φ_Q

$$\begin{aligned} \xi_k(j, j' | \mathbf{r}) &= \Pr(\mathcal{G}_{k,j}^m, \mathcal{G}_{k+1,j'}^m | O_1, \dots, O_z, \Pi, \mathbf{r}, \Phi_P, \Phi_Q) \\ &= \frac{\alpha_k(j) t_k(j, j') b_{j'}(O_{k+1}) \beta_{k+1}(j')}{\sum_{j=1}^{g_m} \sum_{j'=1}^{g_m} \alpha_k(j) t_k(j, j') b_{j'}(O_{k+1}) \beta_{k+1}(j')} \end{aligned} \quad (16)$$

The recombination frequency r_k can be estimated through an iterative process using

$$r_k^{s+1} = \sum_{i=1}^n \sum_{j=1}^{g_m} \sum_{j'=1}^{g_m} \frac{\xi_k(j, j' | \mathbf{r}^s) \phi(j, j')}{n} \quad (17)$$

where $\xi_k(j, j' | \mathbf{r}^s)$ is calculated for individual i , $\phi(j, j') = \frac{(l_P + l_Q)}{m}$ is the proportion of recombinations between markers k and $k+1$ for individuals with genotypes $\mathcal{G}_{k,j}^m$ and $\mathcal{G}_{k+1,j'}^m$ and $\mathbf{r}^s$ is the vector of recombination fractions in the iteration (s) and $\mathbf{r}^{s+1}$ is the updated recombination fraction vector (Broman and Sen 2009).

Estimation of linkage phase

Let the Cartesian product $\phi_P^1 \times \dots \times \phi_P^k \times \dots \times \phi_P^z = \{(\varphi_P^1, \dots, \varphi_P^k, \dots, \varphi_P^z) | \varphi_P^i \in \phi_P^i, i = 1, \dots, z\}$ denotes a set containing all possible linkage phase configurations in parent P . Also, let

$\Phi = \{\Phi^u\} = (\phi_P^1 \times \cdots \times \phi_P^k \times \cdots \times \phi_P^z) \times (\phi_Q^1 \times \cdots \times \phi_Q^k \times \cdots \times \phi_Q^z)$, $u = 1, \cdots, \prod_{k=2}^z \binom{m}{d_P^k} \binom{m}{d_Q^k}$, denote a set containing all possible linkage phase configurations in both parents. The probability of the linkage phase configurations can be obtained using Bayes' rule

$$\Pr(\Phi^u | \mathbf{O}, \mathbf{\Pi}, \mathbf{r}) = \frac{\prod_{i=1}^n \Pr(O_{1,i}, \cdots, O_{z,i} | \mathbf{r}, \mathbf{\Pi}_i, \Phi^u) \Pr(\Phi^u)}{\sum_{\Phi^u \in \Phi} \prod_{i=1}^n \Pr(O_{1,i}, \cdots, O_{z,i} | \mathbf{r}, \mathbf{\Pi}_i, \Phi^u) \Pr(\Phi^u)} \quad (18)$$

where $\mathbf{O}$ is an array containing the observation for z markers in n individuals, and $\mathbf{\Pi}$ is the underlying probability distribution for all marker observations. Since the prior probability $\Pr(\Phi^u)$ can be assumed to be uniform, the posterior probability is proportional to the likelihood of the model, which can be used to select the best linkage phase configuration. Depending on the dosage and number of markers, some of these configurations are equivalent and will result in the same likelihood. The search space for the best linkage phase configuration can be unwieldy depending on the ploidy level, dosage and number of markers. Also, the transition space on the HMM gets larger as the ploidy level increases. To circumvent these problems, we propose a very efficient two-point procedure to reduce the search space for linkage phases. With the estimates in hands, it is possible to compare two-point likelihoods of alternative linkage phase configurations, eliminating those that do not meet a given threshold. The remaining configurations will be evaluated using the multipoint approach. The adequate threshold level will be discussed in the Simulations section.

Two-point algorithm for high-level autopolyploids

When the linkage analysis is conducted only between two markers (two-point analysis), the information contained in these markers does not propagate into the rest of the chain. Thus, based on the dosage and linkage phase configuration of the markers involved in the analysis, the g_m genotypic states present in the full transition space can be collapsed into a small number of states, and a straightforward likelihood function can be derived. It is worthwhile to mention that the estimates obtained using the two-point procedure are the same as those obtained using the multipoint algorithm for two markers. However, the two-point computation is extremely faster.

Consider a biallelic marker in an autopolyploid biparental cross with ploidy m . The number of possible genotypic states in the progeny for a given locus at position k is $u(d_P^k) + u(d_Q^k) + 1$, where the operator $u(x) = \lfloor |x - \frac{m}{2}| \rfloor$ and $\lfloor \cdot \rfloor$ denotes module. For example, in an auto-hexaploid biparental cross, if the dosage of the marker at position k in parent P is two ($d_P^k = 2$) and in parent Q is three ($d_Q^k = 3$), the number of possible genotypic classes expected in the progeny is six. Depending on the linkage phase configuration, each of the g_m genotypic states in the full transition space corresponds to one of these expected genotypic classes, as presented in the emission function (Equations 8 and 9). Thus, in the previous example, all the g_m states could be collapsed into six different classes. To perform this reduction of dimensionality, let $\mathcal{D}_k^m \in \{0, \cdots, m\}$ denote one of the possible genotypes based on the observed dosage of one individual in the progeny of an autopolyploid biparental cross for position k with ploidy m . The joint probability of $\mathcal{D}_k^m$ and $\mathcal{D}_{k'}^m$, for a given genotypic configuration at positions k and k' can be written as

$$\Pr(\mathcal{D}_k^m, \mathcal{D}_{k'}^m | \phi_P^k, \phi_Q^k, \phi_P^{k'}, \phi_Q^{k'}) = \sum_{j \in T_k} \sum_{j' \in T_{k'}} \Pr(\mathcal{G}_{k,j}^m | \mathcal{G}_{k,j'}^m) \Pr(\mathcal{G}_{k,j}^m) \quad (19)$$

where $T_k = \{j | \delta(k, j) = \mathcal{D}_k^m, j = 1, \cdots, g_m\}$ and $\delta(k, j)$ was defined in Equation 8; the same applies to $T_{k'}$. Since in a two-point analysis the probability distribution of the genotypic states in locus k can be assumed to be uniform, i.e., $\Pr(\mathcal{G}_{k,j}^m) = \frac{1}{g_m}$, Equation (19) can be rewritten as a sum of weighted terms from Equation (6)

$$\Pr(\mathcal{D}_k^m, \mathcal{D}_{k'}^m | r_k, \phi_P^k, \phi_Q^k, \phi_P^{k'}, \phi_Q^{k'}) = \sum_{l_P=0}^{\frac{m}{2}} \sum_{l_Q=0}^{\frac{m}{2}} \xi_{T_k, T_{k'}}(l_P, l_Q) \times \frac{(1-r_k)^{m-l_P-l_Q} (r_k)^{l_P+l_Q}}{\binom{m}{2} \binom{m}{2}} \quad (20)$$

where

$$\xi_{T_k, T_{k'}}(l_P, l_Q) = \frac{1}{g_m} \sum_{j \in T_k} \sum_{j' \in T_{k'}} h(j, j'; l_P, l_Q)$$

$h(j, j'; l_P, l_Q)$ is 1 if (j, j') corresponds to (l_P, l_Q) according to the procedure described in Supplementary Information, File S2, and zero otherwise. Equation 20 can be expressed in matrix form as

$$\mathbf{A}_{\phi_P^k, \phi_Q^k, \phi_P^{k'}, \phi_Q^{k'}}(r_k) = \left\{ \Pr(\mathcal{D}_k^m = i - 1, \mathcal{D}_{k'}^m = j - 1 | r_k, \phi_P^k, \phi_Q^k, \phi_P^{k'}, \phi_Q^{k'}) \right\}_{ij} \quad (21)$$

where $\mathbf{A}_{\phi_P^k, \phi_Q^k, \phi_P^{k'}, \phi_Q^{k'}}(r_k)$ is a $(m+1) \times (m+1)$ matrix. Yet, in a two-point analysis with biallelic markers, the linkage phase configuration can be summarized in an ordered pair $(w_P^{k,k'}, w_Q^{k,k'})$ indicating the number of homologs that share allelic variants for loci k and k' in parents P and Q , respectively. For a given pair $(\phi_P^k, \phi_P^{k'})$, $w_P^{k,k'} = \#\{x_P^k \cap x_P^{k'}\}$, where x_P^k and $x_P^{k'}$ denote the set of homologs inherited by parent P in positions k and k' , which can be assessed using the superscripts in ϕ_P^k and $\phi_P^{k'}$. $\#\{\cdot\}$ indicates the cardinality of the set. Notice that ϕ_P^k and $\phi_P^{k'}$ can assume several linkage phase configurations resulting in the same $w_P^{k,k'}$. Let $\Phi_P^{k,k'} = \phi_P^k \times \phi_P^{k'}$ denote a set containing all possible pairs $(\phi_P^k, \phi_P^{k'})$ for a given pair $(d_P^k, d_P^{k'})$. In this set, there are $\min\{u(d_P^k), u(d_P^{k'})\} + 1$ partitions, each one corresponding to a different $w_P^{k,k'}$. Figure 3 shows an example of $\Phi_P^{k,k'}$ for $(d_P^k = 2, d_P^{k'} = 2)$ in an autotetraploid homology group. The size of the set is 36, and it can be subdivided into three partitions where $w_P^{k,k'} = 2$, $w_P^{k,k'} = 1$ and $w_P^{k,k'} = 0$.

In a two-point context, the likelihood function derived from any of the configurations belonging to the same partition (same $w_P^{k,k'}$) will be the same. Thus, any of them can be used to obtain the likelihood function for a given $w_P^{k,k'}$. Let $(\phi_P^k, \phi_P^{k'})^*$ denote one of the possible pairs $(\phi_P^k, \phi_P^{k'})$ that correspond to $w_P^{k,k'}$. The same reasoning applies to parent Q . Without loss of generality, the two-point likelihood function of biallelic observed molecular phenotypes for markers k and k' given $w_P^{k,k'}$ and $w_Q^{k,k'}$ is

$$L(r_k | w_P^{k,k'}, w_Q^{k,k'}) = \prod_{i=1}^n \pi^k \mathbf{A}_{(\phi_P^k, \phi_P^{k'})^*, (\phi_Q^k, \phi_Q^{k'})^*}(r_k) (\pi^{k'})^T \quad (22)$$

where n is the number of individuals and T denotes transposition of a vector. In Equation (22), r_k can be estimated using iterative procedures such as EM or Newton-Raphson. As in Equation (18), it is possible to list all linkage phase configurations and evaluate them

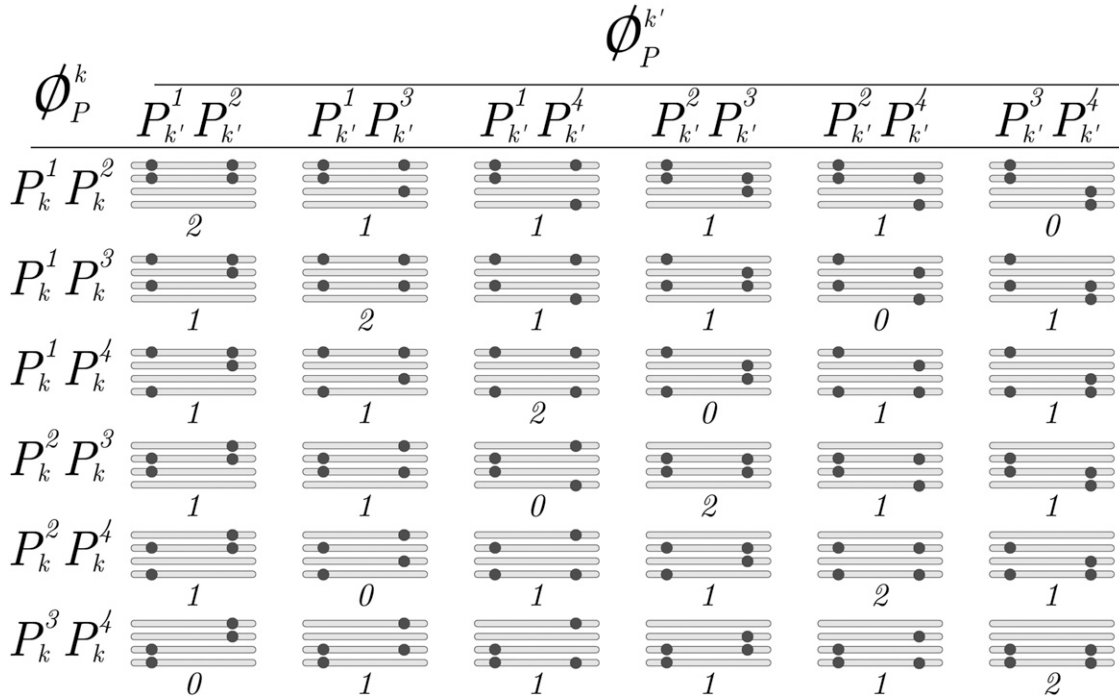


Figure 3 Example of $\Phi_P^{k,k'} = \phi_P^k \times \phi_P^{k'}$ for an autotetraploid homology group with observed dosages $d_P^k = 2$ and $d_P^{k'} = 2$ homologs sharing alleles. In this case, ϕ_P^k denotes a set of size six, containing all possible subsets of size two in $\mathcal{P}_k^4 = \{P_k^1, P_k^2, P_k^3, P_k^4\}$. The same reasoning applies to $\phi_P^{k'}$. The horizontal bars represent homologs forming a homology group and the dots represent allelic variations of a biallelic marker. The number below each homology group represents the number of homologs that share allelic variants ($w_P^{k,k'}$). This defines three partitions: $w_P^{k,k'} = 2$, $w_P^{k,k'} = 1$ and $w_P^{k,k'} = 0$. Notice that, from a homology group within a specific partition, it is possible to obtain the same linkage phase configuration observed in another homology group within that partition by permuting the its homologs.

based on their likelihood. Here we use the LOD Score (base-10 logarithm of likelihood ratios) in relation to the highest likelihood. Thus, models with high likelihoods will yield LOD Scores close to zero. We also use the LOD Score to assess the evidence for linkage between the two markers using the ratio between the model under $H_a : r = \hat{r}$ and under the null hypothesis of no linkage $H_0 : r = 0.5$, given a linkage phase configuration.

As previously shown, it is possible to enumerate all linkage phase configurations for parent P using the Cartesian product $\phi_P^1 \times \phi_P^2 \times \dots \times \phi_P^z$. To reduce this Cartesian space based on two-point analysis, we add a restriction where all pairs $(\phi_P^k, \phi_P^{k'})$ in a sequence of configurations $(\phi_P^1, \dots, \phi_P^z)$ must be contained in $\Phi_P^{k,k'}(\eta)$, where $\Phi_P^{k,k'}(\eta)$ is a subset of all partitions in $\Phi_P^{k,k'}$ in which the associated LOD Score is smaller than η . Thus, a reduced subset of linkage phases in parent P based on two-point analysis can be obtained using

$$\Phi_P(\eta) = \left\{ (\phi_P^1, \dots, \phi_P^z) \mid \phi_P^i \in \phi_P^i \wedge (\phi_P^k, \phi_P^{k'}) \in \Phi_P^{k,k'}(\eta), \right. \\ \left. \forall k, k' \in (1, \dots, z), k > k' \right\} \quad (23)$$

It is important to note that it is not necessary to represent the whole Cartesian space $\{\Phi_P\}$ to restrict the linkage phase configurations to the condition $(\phi_P^k, \phi_P^{k'}) \in \Phi_P^{k,k'}(\eta)$. This procedure can be done through the sequential addition of markers from M_1 to M_z . For each marker $M_{k'}$ added to the end of the chain, the ordered pair (k, k') , $k' = 2, \dots, z$ and $k = k' - 1, \dots, 1$, is evaluated and only linkage phase configurations that meet the condition $(\phi_P^k, \phi_P^{k'}) \in \Phi_P^{k,k'}(\eta)$ $\forall k \in \{k' - 1, \dots, 1\}$ are considered.

Some of the configurations selected using the previous procedure can be equivalent once they are products of a permutation of the same set of

homologs. In order to remove this redundancy, let each one of the selected configurations be represented as a binary matrix of dimensions $(m \times k')$ such as

$$\mathbf{H}_k^u = \{h_{ij}\}_{(m \times k')} = \begin{cases} 1 & \text{if } P_j^i \in \phi_P^j \\ 0 & \text{otherwise} \end{cases} \quad (24)$$

where $u \in \{1, \dots, U\}$, U is the number of selected linkage phase configurations, and k' indicates that $M_{k'}$ was the last marker inserted in the chain. The rows of matrix $\mathbf{H}_k^u$ represent the homologs for the u -th linkage phase configuration with the insertion of the k' -th marker at the end of chain; 1 denotes the presence of an allelic variation, and 0 denotes its absence. If a matrix $\mathbf{H}_{k'}$ could be obtained from a matrix $\mathbf{H}_k^u$ just by permuting the rows (permuting the order of the homologs), these two linkage configurations yield the same likelihood. Thus, one of the configurations should be excluded from consideration. The same reasoning applies to parent Q . This procedure can be done recursively until all redundancy is eliminated. The reduced linkage phase configurations search space considering both parents is obtained using $\Phi(\eta) = \Phi_P(\eta) \times \Phi_Q(\eta)$, such as $\#\{\Phi(\eta)\} \ll \#\{\Phi\}$, combined with the redundancy elimination for homology groups. This sequential procedure results in a set of linkage phase configurations containing markers up to $M_{k'}$, which are evaluated using the HMM likelihood. A LOD Score threshold in relation to the most likely configuration is assumed to determine which configurations should be taken into consideration in the next round of marker inclusion (Figure 4). Additionally, it is possible to limit the two-point search space reduction to a window of SNPs in the terminal part of the chain to speed up the phasing process. Finally, with all markers inserted, the multipoint likelihood of the whole map is used

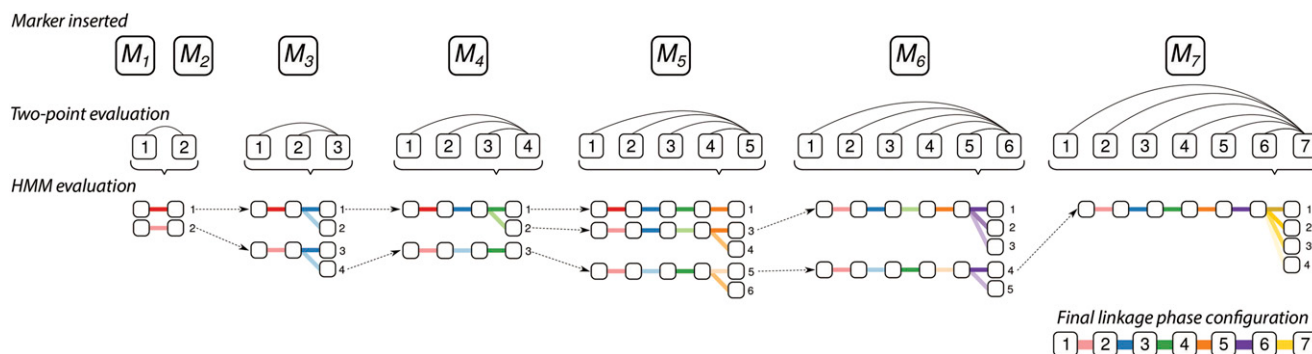


Figure 4 Example of linkage phase configuration estimation using sequential search space reduction and HMM evaluation. Only one parent is presented. The two-point search reduction is composed of two parts: the first one evaluates the LOD Scores obtained through pairwise recombination fraction likelihoods. The second detects equivalent configurations by performing all possible permutations of the homologs. The remaining configurations are evaluated using the HMM-based likelihood. In the first step, linkage phase configurations of M_1 and M_2 are evaluated using the two-point analysis. Color shades indicate different linkage phase configurations provided by the two-point analysis. In this example, there are two possible linkage phases represented by two shades of red. In the second step, we evaluate the linkage phases between markers M_3 and M_2 , and M_3 and M_1 . Configurations with LOD Scores smaller than η are maintained to be evaluated by HMM. There are two possible linkage phases given a certain η , represented by two shades of blue. These two configurations are combined with the configurations from the previous step, resulting in four configurations evaluated using HMM likelihood. Given a likelihood threshold, only configurations 1 and 4 are eligible for the next step. The same reasoning applies for the remaining markers. A final linkage phase configuration is obtained after inserting the last marker and choosing the one that yields the highest HMM-based likelihood.

to find the best configuration among the remaining ones, and the recombination fractions are reestimated. To demonstrate the mechanics of the two-point analysis coupled with the multipoint procedure, a simple example is presented in Supplementary Information, File S3.

Data availability

All the methods and procedures described here are available in the R package *MAPpoly*, which is freely available from <https://github.com/mmollina/mappoly>. R scripts to perform the simulations and the potato map construction presented in this article can be accessed at https://github.com/mmollina/Autopolyploid_Linkage. The tetraploid potato data set is available through the Solanaceae Coordinated Agricultural Project at http://solcap.msu.edu/potato_infinium.shtml. Supplemental material available at FigShare: <https://doi.org/10.25387/g3.8218325>.

RESULTS

Simulations

Simulation 1 - local performance under random bivalent pairing: the aim of this simulation study was to evaluate the local performance of the algorithm considering three ploidy levels ($m = 4$, $m = 6$ and $m = 8$) under the mapping model assumptions (*i.e.*, random pairing and bivalent formation). To be in accordance with molecular data that have been made available through sequence technologies, we simulated bi-allelic markers that can be observed in terms of dosage in parents and progeny. Three different linkage phase scenarios were simulated. In scenario A, for each marker, one of the allelic variants was assigned to the first homolog in the homology group and the remaining variants of the same type were assigned to the subsequent homologs. In scenario B, the allelic variant was randomly assigned to one of the first $\frac{m}{2}$ homolog and the remaining were assigned to the subsequent homologs. In scenario C, the allelic variants were randomly assigned to the m homologs. Thus, it is expected an increasing difficulty to detect recombination events from scenario A, where the allelic variants were concentrated in the same homologs, to scenario C, where they were randomly distributed.

For each combination of ploidy level and linkage phase scenario, we simulated five different parental haplotypes. In total, 45 parental configurations were considered ($3 \times 3 \times 5$, Supplementary Information, Figure S4). For autotetraploid and autohexaploid configurations, we simulated 1000 full-sib populations. For autooctaploids, this number was reduced to 200 due to the high demand of computer processing required to reconstruct such maps. Each population was comprised of 200 individuals with one linkage group containing 10 markers positioned at a fixed distance of 1 centimorgan (cM) between them. For each combination, the percentage of correctly estimated linkage phase configuration in each parent was recorded. Also, for the cases where the linkage phases were correctly estimated, we calculated the average Euclidean distance between the distances of the estimated and simulated maps using

$$\left\{ \frac{(\hat{\mathbf{d}} - \mathbf{d})^T (\hat{\mathbf{d}} - \mathbf{d})}{z-1} \right\}^{-\frac{1}{2}}$$

where $\hat{\mathbf{d}}$ is the vector of distances for an estimated map, $\mathbf{d}$ is the vector of distances for the simulated map, z is the number of markers and T indicates vector transposition. For example, a value of 1 cM indicates that the maps differ 1 cM in average from each other (Mollinari *et al.* 2009). We used the sequential two-point procedure to reduce the search space assuming that linkage phase configurations with associated $LOD < 3.0$ should be investigated using HMM multipoint strategies ($\eta = 3$). For the remaining configurations evaluated using HMM, we kept those with $LOD < 10.0$ to be evaluated in the next round of marker insertion.

Simulation 2 - chromosome-wide performance under preferential pairing and multivalent formation: In this simulation study, we evaluated the performance of the algorithm in dense maps, allowing for multivalent formation and preferential pairing. We used Scenario C from the previous study as a template to simulate five tetraploid and five hexaploid parental haplotypic configurations, each one comprising 200 equally spaced markers with a final length of 100.0 cM (Supplementary Information, Figure S5). For each parental configuration, we simulated 200 full-sib populations of 200 offspring considering a combination of three levels of preferential pairing (0.00, 0.25 and 0.50) and three levels of cross-like quadrivalent formation proportion

■ **Table 1** Percentage of data sets where linkage phase configuration was correctly estimated for five different parental P and Q haplotypes in simulation 1

Ploidy level	A		B		C	
	P	Q	P	Q	P	Q
Autotetraploid ($m = 4$)	100	100	99.7	99.8	100	100
	100	100	99.7	99.7	99.9	99.7
	100	100	100	100	99.7	99.8
	100	100	99.9	99.7	99.9	99.9
	100	100	99.9	99.8	100	100
Autohexaploid ($m = 6$)	100	100	96.6	97.2	96.1	94.6
	100	100	97.3	97.5	95.8	95.6
	100	100	96.5	96.7	94.7	94.6
	100	100	97.3	97.4	96.1	94.7
	100	100	97.2	97.4	95.2	94.5
Autooctaploid ($m = 8$)	100	100	93.6	94.4	93.2	95.7
	100	100	97.6	96.8	92.1	93.9
	100	100	96.8	97.6	90.4	89.2
	100	100	97.7	98.4	90.6	90.0
	100	100	96.9	94.6	88.8	90.6

(0.00, 0.25 and 0.50) with the position of the pairing partner switch varying across simulations. No hexavalents were simulated in this study. For autohexaploids, the multivalent configurations were always composed by a cross-like quadrivalent plus a bivalent. The centromere was positioned at 20.0 cM from the beginning of the chromosome (subtelocentric centromere with arms ratio 1:4) to study the effect of the double reduction which is more pronounced at the end of both chromosome arms. All simulations were conducted using the software PedigreeSim (Voorrips and Maliepaard 2012). In addition to the statistics recorded in Simulation 1, we computed the rate of double reduction observed in each marker for all constructed maps using the “founderalleles” file provided by PedigreeSim. We also evaluate two values for the LOD Score threshold associated to the two-point analysis ($\eta = 3$ and $\eta = 5$). We used a multipoint LOD Score threshold of 10.0 and also limited the two-point search to a 50 SNP window in the terminal part of the map. Markers presenting higher doses than the sum of the doses in both parents, originated from double reduced gametes were filtered-out and assumed as missing data.

Simulation results

Simulation 1: Table 1 shows the percentage of data sets where the linkage phase configuration was correctly estimated in both parents P and Q. In scenario (A) the method was capable of recovering the correct linkage phase configuration in all situations for all ploidy levels. In scenarios (B) and (C) there was a slight decrease in the ability to correctly estimate the linkage phase configuration, especially for $m = 6$ and $m = 8$. Although in these cases the percentages of correctly estimated linkage phases was lower, they were nevertheless high, varying from 100 to 88.8%. This indicates a very good performance to estimate the linkage phase configurations, even using the two-point procedure to narrow the search space.

Figure 5 shows the distributions of the average Euclidean distances between the estimated and simulated distance vectors for the correctly estimated linkage phase configuration. In all cases, the majority of the recombination fractions were consistently estimated once the medians of all distributions are very close 0.5 cM, with no practical problems in terms of mapping construction. These results show that, apart from a relatively small percentage of entangled linkage phase configurations, the method successfully performed the phasing and

managed to estimate the recombination fraction of 10 markers in all situations evaluated.

Simulation 2: The proportion of correctly estimated linkage phase configurations for the dense chromosome-wise map is shown in Table 2. In general, results for tetraploid maps were superior when compared to results for hexaploid maps. It is also possible to observe a better performance for the threshold level $\eta = 5$ in comparison to $\eta = 3$. Similarly to Simulation 1, maps resulting from configurations with no preferential pairing or quadrivalent formation showed a high proportion of correctly estimated linkage phase configurations. Results ranged from 100 to 99% for tetraploid maps and from 100 to 84% for hexaploid maps. Different levels of quadrivalent formation rate had no substantial influence in estimating the correct linkage phase configurations in tetraploids. Within the preferential pairing level 0.0, the percentage of maps with correct linkage phases varied from 100 to 90%. For hexaploids, there was a decrease in this percentage as the quadrivalent formation increases from 0.0 to 0.50, with proportions varying from 100 to 70.5%. Especially for autohexaploids, there was considerable variation between the five simulated configurations. This occurred because the effect of the quadrivalent formation can be more pronounced depending on the level of information contained in a particular configuration. Also, the use of a more stringent two-point threshold $\eta = 5$, improved the performance of the phasing algorithm.

Within the preferential pairing level 0.25, results showed decay of correctly estimated linkage phases, which was more pronounced for hexaploid cases with threshold level $\eta = 3$, reaching a minimum value of 52.5% for parent Q in configuration 1. Again, the use of a higher two-point threshold level, $\eta = 5$, helped to improve this number to 68.5%. For preferential pairing level 0.50, there was a clear distinction between the results in tetraploid and hexaploid cases. In the former, the effect was not as pronounced as it was in the latter, where in several cases, the proportion of correctly estimated linkage phases was close to zero. As expected, the usage of a higher threshold level of $\eta = 5$ helped to improve the number of corrected estimated linkage phase configurations. Interestingly, for both cases with preferential pairing (0.25 and 0.50), the formation of quadrivalents had an overall tendency to improve the algorithm’s performance. This improvement was expected because when a quadrivalent is formed, each chromosome involved can exchange segments with two others, providing more information regarding their phase configuration.

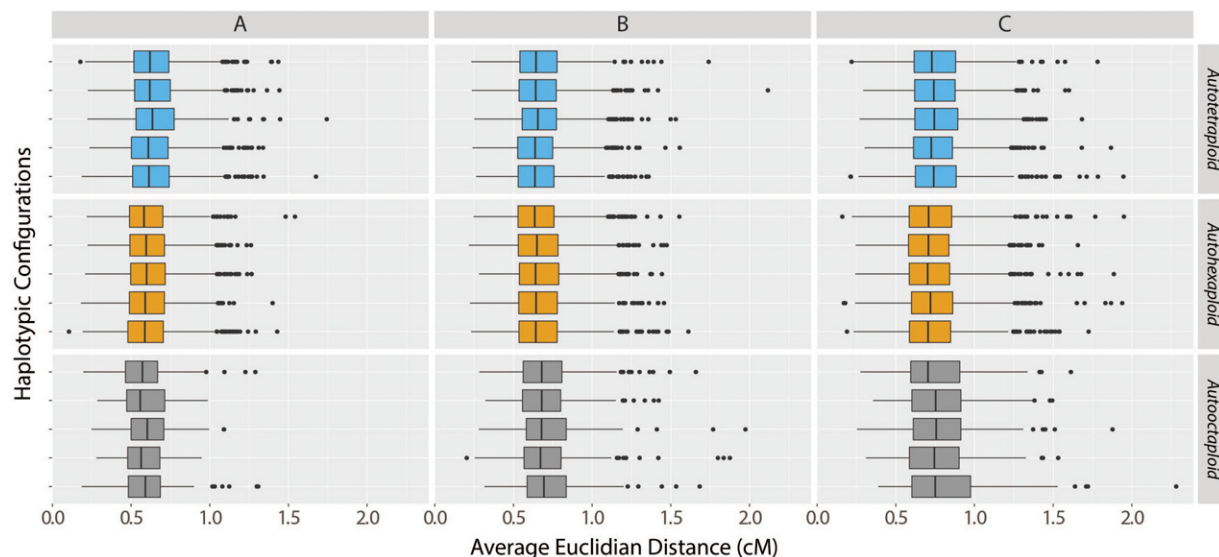


Figure 5 Distributions of the average Euclidean distances between the estimated and simulated distance vectors considering correctly estimated linkage phase configurations. The order of boxplots is the same as the order of haplotypes in Figure S4. Each column indicates the results for different linkage phase configuration scenarios, namely, A, B and C, and each row indicates a different haplotypic configuration within three ploidy levels.

Given a correctly estimated linkage phase, the recombination fractions were consistently estimated for all levels of preferential pairing with no quadrivalent formation (Figure 6). However, they were overestimated in the presence of quadrivalent formation. This effect was mainly observed at the terminal regions of the chromosome, especially in the long arm, where double reduction is more pronounced. In this case, tetraploid maps were the most affected. This is in agreement with our expectations since in autohexaploid simulations, there was always the formation of a bivalent which was not involved in the double reduction process (although the rates of double reduction were very similar in both ploidy levels, Figure 6). In addition to the quadrivalent, the bivalent serves as an extra source of information to access the recombination events. The average Euclidean distances reflect the overestimation of recombination fractions in cases with quadrivalent formation, showing distributions with higher medians and interquartile ranges in tetraploid cases when compared to hexaploids (Supplementary Information, Figure S6). Nevertheless, all the Euclidean distances distributions were located relatively close to zero, with a maximum value of 1.4 cM, indicating that although we observed overestimated recombination fractions toward the terminal ends of the chromosome, they were equally distributed, causing no severe disturbances in the final map. Figure S7, in Supplementary Information, shows an example of the effect of increasing quadrivalent formation rate in autotetraploid and autohexaploid maps. As the markers get further away from the centromere, the recombination fractions become overestimated.

Analysis of real tetraploid potato SNP data set

We applied our method to construct a genetic map of the B2721 population which is a cross between tetraploid potato varieties *Atlantic* and *B1829-5*. The population comprises 160 offsprings genotyped with the SolCAP Infinium 8303 potato array. The genotype calling was performed using fitTetra R package (Voorrips *et al.* 2011). We obtained 4017 SNPs and computed all the pairwise recombination fraction between them for all possible linkage phase configurations. For each pair, we selected the configuration that yields the higher

likelihood and applied the Unweighted Pair Group Method with Arithmetic Mean (UPGMA) clustering algorithm to assign markers into 12 linkage groups. Within each linkage group, we ordered the markers using the MDSMap R package (Preedy and Hackett 2016). We applied the unconstrained multidimensional scaling algorithm (MDS) to the pairwise marker distance obtained using Haldane's mapping function and LOD^2 weighting. We performed two rounds of marker removal by inspecting the nearest neighbor measure scatter plot and also the monotony of the resulting recombination fraction matrix. Given the marker order obtained for each group, we applied our algorithm using $\eta = 10$. For each round of marker inclusion, we limited our phase search to the last 100 markers inserted at the end of the map and eliminated markers that caused map inflation greater than 10 cM. The resulting map consisted of 3348 SNPs (58% simplex and double simplex markers and 42% multiplex) distributed in 12 linkage groups with lengths varying from 165.2 cM to 332.5 cM with no visible gaps between markers (Supplementary Information, Table S8.1). These values seem inflated compared to tetraploid potato maps available in the literature (Hackett *et al.* 2013; Sharma *et al.* 2013; Massa *et al.* 2015; Bourke *et al.* 2015; Rak *et al.* 2017). The map expansion observed here was mainly caused by two sources of error, namely, local marker misplacement and genotyping errors. Both factors cause the detection of spurious recombination events which are propagated through the HMM causing a global overestimation of recombination fractions.

While obtaining a *de novo* order based on multipoint likelihood is not feasible in our current implementation, we used the multipoint likelihood as an objective function to compare the map obtained using the genomic order from the *Solanum tuberosum* genome version 4.03 (Sharma *et al.* 2013) and the MDS based order. When using the genomic order, the length of all linkage groups are smaller, and the likelihoods were substantially superior when compared to the *de novo* MDS-based order (Supplementary Information, Table S8.1 and Figure S8.1). Furthermore, our algorithm estimated the same linkage phase in both cases, indicating the robustness of the phasing method to local marker misplacement. To address the genotyping errors, we used the approaches

■ **Table 2** Percentage of data sets where linkage phase configuration was correctly estimated for parents *P* and *Q* in simulation 2

Preferential pairing		0.00			0.25			0.50		
Quadrivalent formation		0.00	0.25	0.50	0.00	0.25	0.50	0.00	0.25	0.50
<i>Autotetraploid</i>										
$\eta = 3$	<i>P</i>	100.0	99.0	91.5	98.5	98.5	90.0	80.5	93.0	87.5
		100.0	99.5	99.5	98.5	99.5	97.5	57.5	88.5	97.0
		99.5	97.5	98.5	100.0	98.5	94.0	55.0	85.5	94.5
		100.0	100.0	99.5	99.0	98.0	98.0	60.5	86.5	93.0
		99.5	99.5	97.0	98.5	97.0	95.5	67.5	84.5	97.5
		100.0	98.5	90.0	100.0	97.0	90.0	60.0	91.5	86.0
	<i>Q</i>	100.0	100.0	98.0	99.5	100.0	99.0	65.0	89.0	93.5
		100.0	98.5	98.0	97.0	98.5	94.5	41.0	82.0	93.5
		100.0	100.0	99.0	99.5	98.0	98.0	56.5	84.5	90.0
		99.5	99.5	98.0	99.0	98.5	94.5	58.0	82.0	94.0
		100.0	99.5	93.0	100.0	99.5	95.0	98.0	99.0	95.0
		100.0	100.0	100.0	100.0	100.0	100.0	90.0	99.5	99.0
	$\eta = 5$	100.0	99.5	100.0	100.0	100.0	99.5	86.0	98.5	100.0
		100.0	100.0	100.0	99.5	100.0	99.5	86.5	98.5	100.0
		100.0	100.0	100.0	100.0	100.0	100.0	90.5	96.0	100.0
		100.0	99.5	93.0	100.0	99.0	94.0	88.0	98.5	95.5
		100.0	100.0	100.0	100.0	100.0	100.0	91.5	99.5	99.5
		100.0	99.5	100.0	99.5	100.0	99.0	85.0	98.0	100.0
<i>Autohexaploid</i>	$\eta = 3$	100.0	100.0	100.0	99.5	100.0	99.5	86.0	97.5	98.5
		100.0	100.0	99.5	100.0	100.0	100.0	92.0	96.0	99.0
	<i>P</i>	84.0	78.5	70.5	69.0	63.5	61.0	2.5	10.5	19.0
		99.0	94.0	91.0	93.0	84.5	80.0	6.5	16.0	22.0
		89.0	94.0	88.0	80.0	84.0	80.5	10.5	16.0	32.5
		93.0	90.5	86.0	88.5	84.0	80.0	9.0	16.5	28.5
	<i>Q</i>	96.0	92.5	91.5	89.5	94.0	87.5	19.0	30.5	44.5
		85.0	81.0	71.0	68.0	52.5	57.5	1.5	3.5	8.5
		99.0	95.0	91.0	86.5	90.0	88.5	9.0	28.0	37.5
		90.0	90.0	86.0	79.0	82.0	77.0	9.5	18.0	28.0
		96.5	92.5	89.5	90.0	89.0	89.0	25.5	35.5	41.0
		95.0	92.0	92.5	89.5	91.0	88.0	16.0	23.0	39.0
	$\eta = 5$	86.0	84.5	75.5	77.5	69.5	72.5	27.0	36.5	52.5
		100.0	97.5	96.5	98.5	98.0	91.0	55.5	70.5	74.5
		91.5	95.5	93.0	90.5	94.5	89.5	68.0	68.5	77.5
		96.5	94.0	91.0	99.5	99.0	96.5	65.0	78.5	85.0
		98.0	98.5	100.0	97.5	99.0	99.0	73.0	87.5	91.0
		86.5	83.5	75.0	69.5	68.5	72.0	17.5	20.0	39.5
$\eta = 5$	<i>Q</i>	100.0	99.5	99.0	100.0	99.5	100.0	74.0	81.0	92.5
		91.5	95.5	93.0	91.0	95.0	89.5	67.5	71.5	77.0
		99.0	97.5	93.5	100.0	100.0	99.5	80.0	89.0	92.0
		98.0	98.5	100.0	97.5	99.0	99.0	83.0	83.0	90.5

presented in Equations 8 and 9, i.e., (i) use the probability distribution provided by the mixture proportions of the doses from fitTetra software and (ii) assuming a global genotyping error; in this case we assumed an *ad hoc* error rate of 5%. We applied this prior information in both *de novo* MDS-based and the genome-based orders. The result also can be observed in Table S8.1 and Figure S8.1. Both approaches produced smaller maps when compared to their relative original maps. However, since (i) relies on the dosage proportions based in single SNPs, the adjustment of the map was not as flexible as observed in (ii), which assumes an equal global error for all markers. Thus the usage of global error allowed genotypes to conform to a global chromosomal structure, rather than restrain the markers in certain genotypic classes proposed by the classification method. Furthermore, the usage of a global error mitigates the effect of the local marker misplacement caused by the MDS algorithm by clustering markers into linkage disequilibrium blocks. This effect can be observed in linkage

groups 1, 5, 8, 10, 11, and 12, where the difference between *de novo* MDS-based and the genome-based order when modeling a global error was less than 10 cM.

Comparison with polymapR software

Among the available methods to construct maps in high-dose autopolyploids, namely, pergola (Grandke *et al.* 2017), netgwas (Behrouzi and Wit 2017), and polymapR (Bourke *et al.* 2018), only the latter is capable of inferring parental haplotypes and estimating recombination fraction in outcrossing populations. Thus, we limited our comparison to polymapR software. To assess the performance of the methods, we simulated 50 full-sib hexaploid populations with 200 each individuals in five marker density scenarios: 200, 400, 600, 800, and 1000 equally spaced markers in a 100 cM linkage group. Similarly to Simulation 2 described before, we randomly assigned allelic variants, from 0 to 3 doses, to the six homologs. Additionally, in this study, we considered two

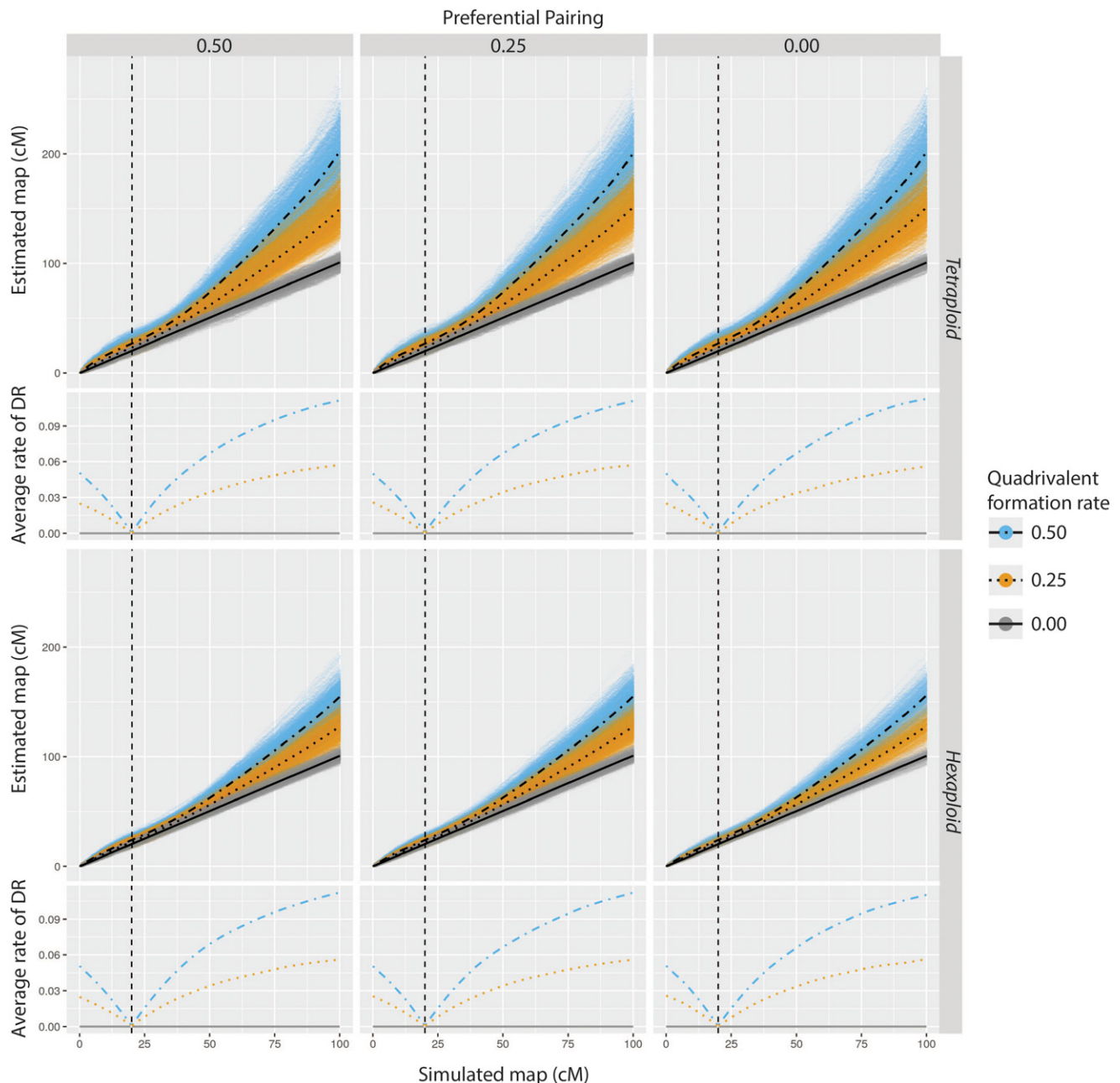


Figure 6 Comparison of estimated vs. simulated maps given a correct estimation of linkage phases in simulation 2. Smoothed conditional means of the observed average rate of double reduction is presented along with the simulated chromosome. The centromere was positioned at 20 cM from its beginning (vertical dashed line). Upper panels show the results for tetraploid simulations while lower panels show the results for hexaploid simulations. Three levels of preferential pairing (0.00, 0.25, 0.50) and three levels of quadrivalent formation rate (0.00, 0.25, 0.50) were simulated. The lines superimposed to the scatter plots are smoothed conditional means of the distances using a generalized additive model. Both two-point thresholds were considered since they only affect the phasing procedure.

different dosage proportions scenarios: the first one considers a higher proportion of simplex and double simplex markers, with 40% of the simulated markers being nulliplex, 40% simplex, 10% duplex and 10% triplex in both parents; the second, considers equal proportions for all doses, with 25% for all dosage types, from nulliplex to triplex (Supplementary Information, Figure S9). In total, 500 populations were simulated ($5 \times 2 \times 50$) and for each, we obtained the phased map using polypmapR and our HMM-based procedure. All simulations were performed using the software PedigreeSim (Voorrips and Maliepaard 2012) considering no preferential

pairing and no quadrivalent formation. For both methods, we recorded the percentage of correctly phased markers, the final map length, and the number of markers inserted in each phased map.

To employ our HMM-based method, we ordered the markers using the unconstrained MDS algorithm (Preedy and Hackett 2016) weighted by LOD^2 with no rounds of marker removal. Given the MDS order, two levels of η were used: 3 and 5; the same levels were used for the multilocus LOD threshold. The phase search was limited to the last 50 markers inserted at the end of the map. To construct maps using polypmapR, we first applied the function cluster_SN_markers

to perform a grid search from LOD Scores 1 to 20 and chose the lowest one that yields six homologs based on simplex markers in coupling linkage phase. In the next step, we assigned double simplex and duplex markers to the linkage group using the function `assign_linkage_group`, and the remaining marker types were assigned using the function `homolog_lg_assignment`. This procedure was performed for both parents assuming LOD Score thresholds of 3 and 5. All remaining pairwise recombination fractions were computed, and the MDS algorithm (Preedy and Hackett 2016) was used to order the markers. Marker positions were estimated using the projection of the MDS result onto a single dimension principal curve. Finally, a phased map was created using the function `create_phased_maplist`. We used `polymapR` version 1.0.19. Differently from the HMM-based method, where the η indicates the LOD threshold from which the multipoint likelihood should be used to chose the best phase configuration, `polymapR` uses LOD thresholds to make decisions whether a marker should be used in a certain mapping context, such as, clustering homologs or to assign it into assembled linkage groups. Thus, they are not directly comparable.

Table S10.1 in Supplementary Information, shows the results obtained using both methods. Overall, both methods recovered the vast majority of the linkage phases of the markers across all simulations. In the presence of a high number of single-dose markers both methods recovered the correct linkage phase from 99.7 to 100% of the included markers with `polymapR` positioning in average 94.3% of the markers when using $LOD = 3.0$ and 85.0% when using $LOD = 5.0$ and our HMM-based method positioned in average 100.0% of the markers when using $LOD = 3.0$ and 99.8% when using $LOD = 5.0$. In cases where the dosages were uniformly assigned, 96.2–100% of the markers were correctly phased when using `polymapR` and from 99.1 to 100% when using the HMM-based method. In this case, the number of positioned markers varied from 16.9 to 77.5% for $LOD = 3$ and from 14.4 to 57.0% for $LOD = 5$ for `polymapR` and from 99.3 to 99.9% for $LOD = 3$ and from 97.8 to 99.9% for $LOD = 5$ when using our HMM-based method. As pointed out by van Geest *et al.* (2017), when using `polymapR`'s method, the presence of a sufficiently large number of simplex markers uniformly distributed throughout the genome is essential to define each homologs and be able to assign multiple dose markers to the framework map. It is important to mention that using modern genotyping technologies, most data sets have high proportions of single dose markers. While this is generally true, single dose markers can be absent in certain chromosome regions or even along entire homologs due to recent duplication events.

Table S10.2, Supplementary Information, shows the average map length and the associated standard deviation obtained in all simulations. In cases with a high number of single-dose markers, the average map length produced by `polymapR` ranged from 87.0 to 89.9 cM, while in the case where the dosages were uniformly assigned, map lengths ranged from 76.7 and 80.1 cM. These results confirm the underestimation tendency in the MDS algorithm when using when using LOD^2 as weighting function, as observed by Preedy and Hackett (2016). In our HMM-based method, considering the MDS order, the maps lengths were highly overestimated, ranging from 143.0 cM to 548.0 cM. Since we did not introduce errors in our simulation procedures, the observed map inflation is exclusively due to local marker misplacement caused by the MDS algorithm. Nonetheless, even with local marker misplacements, the linkage phase configuration was correctly estimated in the vast majority of the cases. To accommodate the local marker misplacements, we used an *ad hoc* global error of 5% in the HMM emission function. As observed in the tetraploid potato analysis, this strategy allowed markers in the wrong order, but in the same

linkage disequilibrium block, being positioned closely together through the HMM estimation process. The resulting maps lengths were close to the simulated 100 cM (Supplementary Information, Table S10.2 and Figure S10.1). Although in general, our method yielded denser maps, it is worthwhile to mention that, `polymapR`'s method is substantially faster than ours, notably when our method uses high values for η , in which case the HMM computations play a significant role in the phasing procedures (Supplementary Information, Table S10.3). Nonetheless, it was precisely the multipoint procedure that allowed our method to position more markers when compared to `polymapR`.

DISCUSSION

Although the concept of linkage mapping is relatively simple, the combinatorial properties and increasingly missing information that arise from the multiple sets of chromosomes make the construction of genetic maps in high-level autopolyploids challenging. In this work, we frame and solve two fundamental steps toward the construction of such maps, namely multipoint recombination fraction estimation and linkage phase estimation. We showed that, combined with standard grouping and ordering procedures (Preedy and Hackett 2016), these maps could be reliably constructed. Our method can be applied to biallelic codominant markers and, due to the flexibility of the HMM framework upon which it was derived, it is extendable to any codominant molecular marker. The HMM used in this work takes into account the linkage phase configuration of the whole linkage group to estimate the recombination fractions between adjacent markers. An efficient two-point approach was also presented to reduce the search space of linkage phase configurations. As a result, our method provides the likelihood of the model, which can be used as an objective function to compare different map configurations, including linkage phases and marker orders. When considering experimental populations, our method is a generalization, for any even ploidy level, of well established genetic linkage mapping methods. For diploid ($m = 2$) populations derived from biparental crosses, our method is equivalent to the influential Lander and Green algorithm (Lander and Green 1987); considering full-sib phase-unknown crosses, it is equivalent to Wu *et al.* (2002). For tetraploids ($m = 4$) the method is equivalent to Leach *et al.* (2010), disregarding double reduction.

To assess the statistical power of our method, we conducted two simulation studies. In simulation 1, we demonstrated that our model was capable of correctly estimating the majority of parental linkage phase configurations and recombination fractions in a limited number of markers, even for complex linkage phase configurations and high ploidy levels. Since other methods are based on single-dose markers to assemble homology groups, to the best of our knowledge, this is the only method capable of phasing markers in high-dose autopolyploid genomes in small regions. These well-assembled regions could function as multiallelic codominant markers which propagate their information through the HMM to the rest of the chain, improving the quality of the final map. In simulation 2, we analyzed a sequence of 200 markers in combinations of different levels of preferential pairing and rates of quadrivalent formation. In this situation, quadrivalent formation rate had a marginal effect on the phasing procedure, whereas preferential pairing reduced its performance, especially for autohexaploids. The usage of a higher two-point threshold (η) improved the linkage phase estimation in all cases. This fact indicates that the haplotype phasing is more accurate when HMM-based likelihood is used as objective function to evaluate linkage phases. We also observed that quadrivalent formation yield overestimated recombination fractions between adjacent markers located further away from

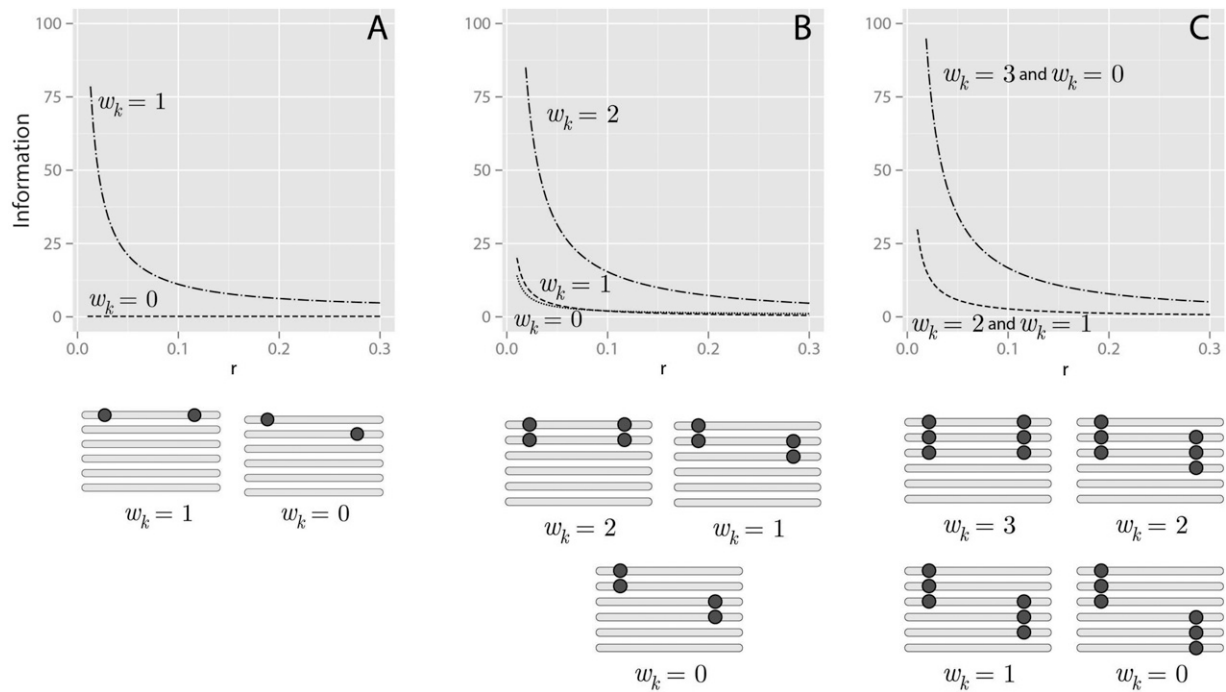


Figure 7 Fisher's information for the two-point maximum likelihood estimators in different combinations of dosages and linkage phases configurations considering one informative hexaploid parent. (I) single-dose markers; alleles share 1 and 0 homologs. (II) double-dose markers; alleles share 2, 1 and 0 homologs. (III) triple dose markers; alleles share 3, 2, 1 and 0 homologs.

the centromere. Interestingly, Bourke *et al.* (2015, 2016) found that higher quadrivalent rates had little effect on the recombination fractions. Most likely, the source of this discrepancy was the different usage of information in both approaches. While in both studies (Bourke *et al.* 2015, 2016) used exclusively two-point recombination fraction estimates, here the information of the markers propagates along the linkage group and the recombination fraction between two adjacent markers depends on the behavior of the whole chromosome. Moreover, the overestimation of the recombination fraction was expected since our model disregards double reduction and, consequently, was not able to correctly assess the number of crossing over events when this phenomenon was present.

Although our model is robust enough to cope with certain levels of preferential pairing and tetravalent rate formation, it is possible to include both phenomena in specific points of its derivation. Preferential pairing can be included in Equation 4 by not considering $\Pr(\psi_j)$ as uniformly distributed. Double reduction can be included in the definition of the genotypic states in the full transition space (Equation 5). These two phenomena add extra layers of complexity to the genetic mapping of polyploid organisms with high ploidy levels and should be addressed in future studies.

We also build a tetraploid map using the ideas presented in this study coupled with standard grouping and ordering procedures. While the choice of the genotype error rate depends on specific characteristics of the data set, we demonstrate that it is possible to use prior information on the HMM framework, including the probability distribution of the marker dosages for each SNP and a global error to avoid map inflation caused by local marker misplacement and genotyping errors. Finally, we compared our HMM-based method to the polyploidR two-point based method and, as already pointed out by (van Geest *et al.* 2017), we concluded that with a number sufficiently large of single-dose markers uniformly distributed across the homologs, both methods performed well. However, when those markers are absent in a specific homologs

or chromosome region, our method was able to build denser maps when compared to polyploidR. Moreover, some autopolyploids with ploidy level higher than six, such as sugarcane (Aitken *et al.* 2007) and garden dahlias (Schie *et al.* 2014), could benefit only from our method, since polyploidR is limited only to tetraploid and hexaploid species. The difficulty in correctly estimating linkage phase configurations where multi-dose allelic variants are spread randomly in all homologous chromosomes lies in two significant aspects of the experiments studied here: (i) the outbred nature of the experimental crosses and (ii) the incomplete information of the markers based on dosage (*i.e.*, by not being multiallelic). In experimental population derived from inbred lines, the origin of the haplotypes can be easily inferred from the genetic design. However, obtaining pure inbred lines in high-level autopolyploids has been proven to be impractical due to the high number of crosses and generations necessary to achieve homozygous genotypes and to the inbreeding depression which some species undergo (Gallais 2003). In our method, the linkage phase configuration is obtained by comparing the likelihood of a set of models with different linkage phase configurations (Equation 18). The capability of estimating the correct configuration is directly related to the information contained in the marker data. Some of these limitations are overcome through the use of HMMs which take into account the information of the linkage group as a whole.

HMMs provide an excellent avenue to assemble genetic maps in complex scenarios, but they are remarkably computational demanding and, in some cases, unfeasible to use. Apart from parallel computing, which can greatly speed up the estimation process and is ubiquitous nowadays, the usage of two-point approaches is a viable option to reduce the dimension of the original problem efficiently. The dimension reduction is achieved by collapsing genotypic states in the full transition space according to the marker information. However, in several cases, the two-point based method can result in low statistical power which is related to the amount of information contained in markers in certain

combinations of allelic dosage and linkage phase configurations. This lack of information is exacerbated as markers get distant from each other. Figure 7 shows nine possible configurations of pairs of markers in one autohexaploid parent. Considering one of the parents non-informative, we computed the Fisher's information equations based on the likelihood Equation (22) (Mather 1957; Ripol *et al.* 1999; Luo *et al.* 2004). The equations were plotted as a function of the recombination fraction. The information profiles are related to the number of different haplotypes present on the parental configuration for a given marker dosage. For instance, for two single-dose markers (Figure 7, panel A), when the alleles share the same homolog ($w_k = 1$), it is always possible to detect if the gamete contains at least one recombinant chromosome. However, when the alleles are in different homologs ($w_k = 0$), the detection of recombination events is limited to meiotic configurations containing a bivalent where these chromosomes paired with each other. Intermediate situations involving multi-dose markers can be observed in the other panels in Figure 7. Additionally, the model proposed here contemplates both parents on the analyses, leading to more complicated linkage phase configurations and information equations. The lack of information for some phase configurations in two-point procedures is essentially caused by the biallelic nature of the dosage-based markers. However, in several situations, genomic and transcriptomic references are available for related diploid species and often provide the physical order of the SNPs in small regions. Our phasing procedure could be applied in these regions to obtain local haplotypes, which could function as multiallelic markers improving the information in a two-point analysis. Moreover, in a multipoint context, when using multiallelic markers, the number of visited states in the Markov model can be significantly reduced, making the HMM procedure much more efficient. Ideally, in a full-sib population, the number of different alleles should be as high as two times the ploidy level (fully informative). In this case, the Markov model would be fully observed and, the task of estimating recombination fraction reduces to count the number of recombinant events given a linkage phase configuration. Since our algorithm does not need the entire transition space to work, only a subset of states should be visited, making the calculation much faster when compared to the biallelic case.

Once the map is assembled, given the HMM framework, it is a trivial exercise to obtain the probability of a specific genotype at any map position, conditioned on the whole linkage group and to compute the probability of any unobserved genotype given the genetic map using this information. These conditional probabilities are the basis for answering a series of fundamental questions about quantitative trait loci analysis in high-level autopolyploids, such as the effect of the dosage level on the variation of quantitative traits, the interaction of the alleles within (dominance effects) and between loci (epistatic effects). Therefore, the present study provides a sound basis for unveiling the complex structure of autopolyploid genomes through genetic mapping.

ACKNOWLEDGMENTS

The authors wish to thank Dr. Guilherme da Silva Pereira and Dr. Zhao-Bang Zeng for their invaluable suggestions for elaboration of the manuscript. We also thank Dr. Peter Bourke for consulting with us on the comparison between polypmapR and our HMM-based method. This work was supported by the Bill and Melinda Gates Foundation [OPP1052983] and is part of the Genomic Tools for Sweetpotato Improvement project (GT4SP); MM was also supported by Fundação de Amparo à Pesquisa do Estado de São Paulo (FAPESP) [2012/17009-8, 2013/12245-8]; AAFG was supported by FAPESP [08/52197-4]; AAFG has a productivity scholarship from CNPq.

LITERATURE CITED

- Aitken, K., P. Jackson, and C. McIntyre, 2007 Construction of a genetic linkage map for *Saccharum officinarum* incorporating both simplex and duplex markers to increase genome coverage. *Genome* 50: 742–756. <https://doi.org/10.1139/G07-056>
- Arizio, C. M., S. M. Costa Tártara, and M. M. Manifesto, 2014 Carotenoids gene markers for sweetpotato (*Ipomoea batatas* L. lam): applications in genetic mapping, diversity evaluation and cross-species transference. *Mol. Genet. Genomics* 289: 237–251. <https://doi.org/10.1007/s00438-013-0803-3>
- Bargary, N., J. Hinde, and A. A. F. Garcia, 2014 Finite mixture model clustering of SNP data, pp. 139–157 in *Statistical Modeling in Biostatistics and Bioinformatics*. Springer, Switzerland.
- Behrouzi, P. and E. C. Wit, 2017 De novo construction of polyploid linkage maps using discrete graphical models. *arXiv*.
- Birchler, J. A., 2012 Genetic Consequences of Polyploidy in Plants, pp. 21–32 in *Polyploidy and Genome Evolution*, edited by Soltis, P. S., and D. E. Soltis. Springer-Verlag, Berlin, Germany. https://doi.org/10.1007/978-3-642-31442-1_2
- Bourke, P. M., G. van Geest, R. E. Voorrips, J. Jansen, T. Kranenburg *et al.*, 2018 polypmapR—linkage analysis and genetic map construction from F1 populations of outcrossing polyploids. *Bioinformatics* 34: 3496–3502. Erratum: 540. <https://doi.org/10.1093/bioinformatics/bty371>
- Bourke, P. M., R. E. Voorrips, T. Kranenburg, J. Jansen, R. G. F. Visser *et al.*, 2016 Integrating haplotype-specific linkage maps in tetraploid species using SNP markers. *Theor. Appl. Genet.* 129: 2211–2226. <https://doi.org/10.1007/s00122-016-2768-1>
- Bourke, P. M., R. E. Voorrips, R. G. Visser, and C. Maliepaard, 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>
- Broman, K., and S. Sen, 2009 *A Guide to QTL Mapping with R/qtl*, Springer, New York. <https://doi.org/10.1007/978-0-387-92125-9>
- Buetow, K. H., and A. Chakravarti, 1987 Multipoint gene mapping using seriation. I. General methods. *Am. J. Hum. Genet.* 41: 180–188.
- Burnham, C. R., 1962 *Discussions in cytogenetics*, Burgess Publishing, Minneapolis.
- Comai, L., 2005 The advantages and disadvantages of being polyploid. *Nat. Rev. Genet.* 6: 836–846. <https://doi.org/10.1038/nrg1711>
- Doerge, R. W., 1996 Constructing genetic maps by rapid chain delineation. *J. Quant. Trait Loci* 2: 1–14.
- Doerge, R. W., and B. A. Craig, 2000 Model selection for quantitative trait locus analysis in polyploids. *Proc. Natl. Acad. Sci. USA* 97: 7951–7956. <https://doi.org/10.1073/pnas.97.14.7951>
- Fisher, R. A., 1947 The theory of linkage in polysomic inheritance. *Philos. Trans. R. Soc. Lond. B Biol. Sci.* 233: 55–87. <https://doi.org/10.1098/rstb.1947.0006>
- Fisher, R. A., 1943 Allowance for double reduction in the calculation of genotype frequencies with polysomic inheritance. *Ann. Eugen.* 12: 169–171. <https://doi.org/10.1111/j.1469-1809.1943.tb02320.x>
- Gallais, A., 2003 *Quantitative genetics and breeding methods in autopolyploids plants*, INRA, Paris.
- Garcia, A. A. F., M. Mollinari, T. G. Marconi, O. R. Serang, R. R. Silva *et al.*, 2013 SNP genotyping allows an in-depth characterisation of the genome of sugarcane and other complex autopolyploids. *Sci. Rep.* 3: 3399. <https://doi.org/10.1038/srep03399>
- Grandke, F., S. Ranganathan, N. van Bers, J. R. de Haan, and D. Metzler, 2017 PERGOLA: Fast and deterministic linkage mapping of polyploids. *BMC Bioinformatics* 18: 12. <https://doi.org/10.1186/s12859-016-1416-8>
- Hackett, C. A., 2001 A comment on Xie and Xu: 'Mapping quantitative trait loci in tetraploid species'. *Genet. Res.* 78: 187–189. <https://doi.org/10.1017/S0016672301005262>
- Hackett, C., J. E. Bradshaw, and G. J. Bryan, 2014 QTL mapping in autotetraploids using snp dosage information. *Theor. Appl. Genet.* 127: 1885–1904. <https://doi.org/10.1007/s00122-014-2347-2>
- Hackett, C. A., J. E. Bradshaw, and J. W. McNicol, 2001 Interval mapping of quantitative trait loci in autotetraploid species. *Genetics* 159: 1819–1832.

- Hackett, C. A., and L. B. Broadfoot, 2003 Effects of genotyping errors, missing values and segregation distortion in molecular marker data on the construction of linkage maps. *Heredity* 90: 33–38. <https://doi.org/10.1038/sj.hdy.6800173>
- Hackett, C. A., K. McLean, and G. J. Bryan, 2013 Linkage analysis and QTL mapping using SNP dosage data in a tetraploid potato mapping population. *PLoS One* 8: e63939. <https://doi.org/10.1371/journal.pone.0063939>
- Haldane, J., 1930 Theoretical Genetics of Autopolyploids. *J. Genet.* 22: 359–372. <https://doi.org/10.1007/BF02984197>
- Jiang, C., and Z.-B. Zeng, 1997 Mapping quantitative trait loci with dominant and missing markers in various crosses from two inbred lines. *Genetica* 101: 47–58. <https://doi.org/10.1023/A:1018394410659>
- Krieger, A., J. C. Cervantes, K. Burg, R. O. Mwanga, and D. Zhang, 2003 A genetic linkage map of sweetpotato [*Ipomoea batatas*(L.) lam.] based on aflp markers. *Mol. Breed.* 11: 169–185. <https://doi.org/10.1023/A:1022870917230>
- Lander, E. S., and P. Green, 1987 Construction of multilocus genetic linkage maps in humans. *Proc. Natl. Acad. Sci. USA* 84: 2363–2367. <https://doi.org/10.1073/pnas.84.8.2363>
- Lander, E. S., P. Green, J. Abrahamson, A. Barlow, M. J. Daly *et al.*, 1987 Mapmaker: An interactive computer package for constructing primary genetic linkage maps of experimental and natural populations. *Genomics* 1: 174–181. [https://doi.org/10.1016/0888-7543\(87\)90010-3](https://doi.org/10.1016/0888-7543(87)90010-3)
- Leach, L. J., L. Wang, M. J. Kearsey, and Z. Luo, 2010 Multilocus tetrasomic linkage analysis using hidden markov chain model. *Proc. Natl. Acad. Sci. USA* 107: 4270–4274. <https://doi.org/10.1073/pnas.0908477107>
- Lemmon, Z. H., and J. F. Doebley, 2014 Genetic dissection of a genomic region with pleiotropic effects on domestication traits in maize reveals multiple linked QTL. *Genetics* 198: 345–353. <https://doi.org/10.1534/genetics.114.165845>
- Lewin, H. A., D. M. Larkin, J. Pontius, and S. J. O'Brien, 2009 Every genome sequence needs a good map. *Genome Res.* 19: 1925–1928. <https://doi.org/10.1101/gr.094557.109>
- Li, J., K. Das, G. Fu, C. Tong, Y. Li *et al.*, 2010 EM algorithm for mapping quantitative trait loci in multivalent tetraploids. *Int. J. Plant Genomics* 2010: 216547. <https://doi.org/10.1155/2010/216547>
- Luo, M.-C., Y. Q. Gu, F. M. You, K. R. Deal, Y. Ma *et al.*, 2013 A 4-gigabase physical map unlocks the structure and evolution of the complex genome of *Aegilops tauschii*, the wheat d-genome progenitor. *Proc. Natl. Acad. Sci. USA* 110: 7940–7945. <https://doi.org/10.1073/pnas.1219082110>
- Luo, Z. W., R. M. Zhang, and M. J. Kearsey, 2004 Theoretical basis for genetic linkage analysis in autotetraploid species. *Proc. Natl. Acad. Sci. USA* 101: 7040–7045. <https://doi.org/10.1073/pnas.0304482101>
- Luo, Z. W., Z. Zhang, L. Leach, R. M. Zhang, J. E. Bradshaw *et al.*, 2006 Constructing genetic linkage maps under a tetrasomic model. *Genetics* 172: 2635–2645. <https://doi.org/10.1534/genetics.105.052449>
- Massa, A. N., N. C. Manrique-Carpintero, J. J. Coombs, D. G. Zarka, A. E. Boone *et al.*, 2015 Genetic linkage mapping of economically important traits in cultivated tetraploid potato (*Solanum tuberosum* L.). *G3: Genes, Genomes, Genetics* 5: 2357–2364.
- Mather, K., 1936 Segregation and linkage in autotetraploids. *J. Genet.* 32: 287–314. <https://doi.org/10.1007/BF02982683>
- Mather, K., 1957 *The measurement of linkage in heredity*, Methuen & Co, London.
- Mollinari, M., G. R. A. Margarido, R. Vencovsky, and A. A. F. Garcia, 2009 Evaluation of algorithms used to order markers on genetic maps. *Heredity* 103: 494–502. <https://doi.org/10.1038/hdy.2009.96>
- Mollinari, M., and O. Serang, 2015 Quantitative SNP Genotyping of Polyploids with MassARRAY and Other Platforms, pp. 215–241 in *Plant genotyping: methods and protocols*, edited by Batley, J. Springer, New York. https://doi.org/10.1007/978-1-4939-1966-6_17
- Muller, H. J., 1914 A New Mode of Segregation in Gregory's Tetraploid *Primulas*. *Am. Nat.* 48: 508–512. <https://doi.org/10.1086/279426>
- Osborn, T. C., J. C. Pires, J. A. Birchler, D. L. Auger, Z. J. Chen *et al.*, 2003 Understanding mechanisms of novel gene expression in polyploids. *Trends Genet.* 19: 141–147. [https://doi.org/10.1016/S0168-9525\(03\)00015-5](https://doi.org/10.1016/S0168-9525(03)00015-5)
- Otto, S. P., and J. Whitton, 2000 Polyploid incidence and evolution. *Annu. Rev. Genet.* 34: 401–437. <https://doi.org/10.1146/annurev.genet.34.1.401>
- Preedy, K. F., and C. A. Hackett, 2016 A rapid marker ordering approach for high-density genetic linkage maps in experimental autotetraploid populations using multidimensional scaling. *Theor. Appl. Genet.* 129: 2117–2132. <https://doi.org/10.1007/s00122-016-2761-8>
- Rabiner, L., 1989 A tutorial on hidden markov models and selected applications in speech recognition. *Proc. IEEE* 77: 257–286. <https://doi.org/10.1109/5.18626>
- Rak, K., P. C. Bethke, and J. P. Palta, 2017 Qtl mapping of potato chip color and tuber traits within an autotetraploid family. *Mol. Breed.* 37: 15. <https://doi.org/10.1007/s11032-017-0619-7>
- Rehmsmeier, M., 2013 A computational approach to developing mathematical models of polyploid meiosis. *Genetics* 193: 1083–1094. <https://doi.org/10.1534/genetics.112.145581>
- Ripol, M. I., G. A. Churchill, J. A. G. D. Silva, and M. Sorrells, 1999 Statistical aspects of genetic mapping in autopolyploids. *Gene* 235: 31–41. [https://doi.org/10.1016/S0378-1119\(99\)00218-8](https://doi.org/10.1016/S0378-1119(99)00218-8)
- Schie, S., R. Chaudhary, and T. Debener, 2014 Analysis of a Complex Polyploid Plant Genome using Molecular Markers: Strong Evidence for Segmental Allooctoploidy in Garden Dahlias. *Plant Genome* 7. <https://doi.org/10.3835/plantgenome2014.01.0002>
- Serang, O., M. Mollinari, and A. A. Garcia, 2012 Efficient exact maximum a posteriori computation for bayesian snp genotyping in polyploids. *PLoS One* 7: e30906. <https://doi.org/10.1371/journal.pone.0030906>
- Sharma, S. K., D. Bolser, J. de Boer, M. Sønderkær, W. Amorós *et al.*, 2013 Construction of Reference Chromosome-Scale Pseudomolecules for Potato: Integrating the Potato Genome with Genetic and Physical Maps. *G3: Genes, Genomes, Genetics* 3: 2031–2047.
- Shirasawa, K., M. Tanaka, Y. Takahata, D. Ma, Q. Cao *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>
- Soltis, D. E., M. C. Segovia-Salcedo, I. Jordon-Thaden, L. Majure, N. M. Miles *et al.*, 2014a Are polyploids really evolutionary dead-ends (again)? A critical reappraisal of Mayrose *et al.* (*J. New Phytol.* 202: 1105–1117. <https://doi.org/10.1111/nph.12756>
- Soltis, D. E., P. S. Soltis, and L. H. Rieseberg, 1993 Molecular Data and the Dynamic Nature of Polyploidy. *Crit. Rev. Plant Sci.* 12: 243–273. <https://doi.org/10.1080/07352689309701903>
- Soltis, D. E., C. J. Visger, and P. S. Soltis, 2014b The polyploidy revolution then... and now: Stebbins revisited. *Am. J. Bot.* 101: 1057–1078. <https://doi.org/10.3732/ajb.1400178>
- Sorrells, M. E., 1992 Development and application of rflps in polyploids. *Crop Sci.* 32: 1086. <https://doi.org/10.2135/cropsci1992.0011183X003200050003x>
- Syngma, J., 1975 *Meiotic configurations*, Springer, Berlin. <https://doi.org/10.1007/978-3-642-80960-6>
- van Geest, G., P. M. Bourke, R. E. Voorrips, A. Marasek-Ciolakowska, Y. Liao *et al.*, 2017 An ultra-dense integrated linkage map for hexaploid chrysanthemum enables multi-allelic QTL analysis. *Theor. Appl. Genet.* 130: 2527–2541. <https://doi.org/10.1007/s00122-017-2974-5>
- Van Ooijen, J. W., and J. Jansen, 2013 *Genetic Mapping in Experimental Populations*, Cambridge University Press, Cambridge. <https://doi.org/10.1017/CBO9781139003889>
- Van Os, H., P. Stam, R. G. Visser, and H. J. Van Eck, 2005 Record: a novel method for ordering loci on a genetic linkage map. *Theor. Appl. Genet.* 112: 30–40. <https://doi.org/10.1007/s00122-005-0097-x>
- Voorrips, R. E., G. Gort, and B. Vosman, 2011 Genotype calling in tetraploid species from bi-allelic marker data using mixture models. *BMC Bioinformatics* 12: 172. <https://doi.org/10.1186/1471-2105-12-172>
- Voorrips, R. E., and C. Maliepaard, 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>
- Wang, H., F. A. van Eeuwijk, and J. Jansen, 2016 The potential of probabilistic graphical models in linkage map construction. *Theor. Appl. Genet.* 130: 1–12.

- Wang, J., B. Roe, S. Macmil, Q. Yu, J. E. Murray *et al.*, 2010 Microcollinearity between autopolyploid sugarcane and diploid sorghum genomes. *BMC Genomics* 11: 261. <https://doi.org/10.1186/1471-2164-11-261>
- Wu, K. K., W. Burnquist, M. E. Sorrells, T. L. Tew, P. H. Moore *et al.*, 1992 The detection and estimation of linkage in polyploids using single-dose restriction fragments. *Theor. Appl. Genet.* 83: 294–300. <https://doi.org/10.1007/BF00224274>
- Wu, R., C.-X. Ma, and G. Casella, 2004 A bivalent polyploid model for mapping quantitative trait loci in outcrossing tetraploids. *Genetics* 166: 581–595. <https://doi.org/10.1534/genetics.166.1.581>
- Wu, R., C.-X. Ma, I. Painter, and Z.-B. Zeng, 2002 Simultaneous Maximum Likelihood Estimation of Linkage and Linkage Phases in Outcrossing Species. *Theor. Popul. Biol.* 61: 349–363. <https://doi.org/10.1006/tpbi.2002.1577>
- Wu, Y., P. R. Bhat, T. J. Close, and S. Lonardi, 2008 Efficient and accurate construction of genetic linkage maps from the minimum spanning tree of a graph. *PLoS Genet.* 4: e1000212. <https://doi.org/10.1371/journal.pgen.1000212>
- Xu, F., Y. Lyu, C. Tong, W. Wu, X. Zhu *et al.*, 2013 A statistical model for QTL mapping in polysomic autotetraploids underlying double reduction. *Brief. Bioinform.* 15: 1044–1056. <https://doi.org/10.1093/bib/bbt073>
- Zheng, C., R. E. Voorrips, J. Jansen, C. A. Hackett, J. Ho *et al.*, 2016 Probabilistic multilocus haplotype reconstruction in outcrossing tetraploids. *Genetics* 203: 119–131. <https://doi.org/10.1534/genetics.115.185579>

Communicating editor: J. Holland