

Parameter Estimation and Quantitative Parametric Linkage Analysis with GENEHUNTER-QMOD

Thomas Künzel^{a, b} Konstantin Strauch^{a–c}

^aInstitute of Medical Biometry and Epidemiology, Philipps University Marburg, Marburg, ^bInstitute of Medical Informatics, Biometry and Epidemiology, Chair of Genetic Epidemiology, Ludwig-Maximilians-Universität, Munich, and ^cInstitute of Genetic Epidemiology, Helmholtz Zentrum München – German Research Center for Environmental Health, Neuherberg, Germany

Key Words

Linkage analysis · Likelihood calculation · Quantitative phenotypes · Multipoint analysis · Lander-Green algorithm · GENEHUNTER

Abstract

Objective: We present a parametric method for linkage analysis of quantitative phenotypes. The method provides a test for linkage as well as an estimate of different phenotype parameters. We have implemented our new method in the program GENEHUNTER-QMOD and evaluated its properties by performing simulations. **Methods:** The phenotype is modeled as a normally distributed variable, with a separate distribution for each genotype. Parameter estimates are obtained by maximizing the LOD score over the normal distribution parameters with a gradient-based optimization called PGRAD method. **Results:** The PGRAD method has lower power to detect linkage than the variance components analysis (VCA) in case of a normal distribution and small pedigrees. However, it outperforms the VCA and Haseman-Elston regression for extended pedigrees, nonrandomly ascertained data and non-normally distributed phenotypes. Here, the higher power even goes along with conservativeness, while the VCA has an inflated type I error.

Parameter estimation tends to underestimate residual variances but performs better for expectation values of the phenotype distributions. **Conclusion:** With GENEHUNTER-QMOD, a powerful new tool is provided to explicitly model quantitative phenotypes in the context of linkage analysis. It is freely available at <http://www.helmholtz-muenchen.de/genepi/downloads>.

Copyright © 2012 S. Karger AG, Basel

Introduction

To a large degree, human diseases are influenced or caused by genetic variants. In order to understand the mechanism of the disease and to treat patients in a causative way, a first step is to locate the genetic variants in the human genome. An important tool for achieving this goal is linkage analysis. Linkage analysis uses the fact that alleles at loci that are physically close together are often inherited together. When mapping disease genes, linkage analysis investigates if alleles at a marker locus with a known position are disproportionately often inherited together with the disease. In the era of genome-wide or exome-wide sequencing to screen for rare disease-causing variants, linkage analysis has the important task to

restrict the analysis to those genetic regions that most likely harbor the corresponding genes.

There are various types of statistical methods for linkage analysis. A commonly used test statistic is the logarithm-of-odds score (LOD score). Examples of software packages for linkage analysis are GENEHUNTER [1], MERLIN [2] or ALLEGRO [3] (for an overview, see e.g. [4]). Linkage analysis is possible for both dichotomous and quantitative phenotypes. In the latter, the phenotype is expressed by a measured numeric variable (for a thorough introduction to quantitative phenotypes, see for instance [5] or [6]). A method dealing with quantitative traits is the variance components analysis (VCA), which has been implemented in several software packages, e.g. GENEHUNTER or SOLAR [7]. In addition to the LOD score, extra information regarding the variance components is given. Another linkage analysis method to map quantitative trait loci is the Haseman-Elston regression approach [8]. However, the Haseman-Elston regression only exploits relationships between sib pairs in the analysis and does not provide specific information about the genotype-phenotype relation. The MDE method [9], a more sophisticated method based on the Haseman-Elston regression, uses relative pairs of various types in the analysis. Like VCA and the parametric approach implemented in the LINKAGE package [10], it assumes equal residual phenotypic variances for all three disease locus genotypes.

In this paper, we present a parametric linkage analysis method for quantitative phenotypes called PGRAD optimization, which uses the complete pedigree information rather than merely looking at pairs of relatives. It models the phenotype as a normally distributed variable conditional on the genotype at the trait locus. By maximizing the LOD score over the genotype-specific means and variances, it can both locate the putative disease locus and provide specific information about the genotype-phenotype relation.

Parametric linkage analysis for quantitative traits has been previously implemented into the software packages LINKAGE [10] and PAP [11]. However, both programs use the Elston-Stewart algorithm [12], which restricts the analysis to a small number of markers. This impedes their application to gene mapping projects based on SNPs, which are less informative than microsatellite markers, such that a larger number of SNPs needs to be analyzed jointly.

We have implemented the PGRAD method in the GENEHUNTER software package, which is based on the Lander-Green algorithm [13] and therefore can simulta-

neously handle many markers. GENEHUNTER is a free program which was originally written by Kruglyak et al. [1]. Since then, it has been modified and enhanced several times (see for instance [14–18]). The new extension for PGRAD optimization is called GENEHUNTER-QMOD. Our method allows researchers to model genotype-specific residual variances, dominance effects, and genomic imprinting. By performing simulations, we compare the power of GENEHUNTER-QMOD to that of the VCA and Haseman-Elston regression for a variety of scenarios.

Methods

Overview

In this section, we briefly explain the principles of linkage analysis. Basically, parametric linkage analysis is performed by means of a statistical likelihood ratio test. We start by giving the definition of the LOD score, which is also the final output of the GENEHUNTER calculation.

$$\text{LOD} = \max_{\theta \in [0, 1/2]} \log_{10} \left(\frac{P(\text{data} | \phi, \theta)}{P(\text{data} | \phi, \theta = 1/2)} \right) \quad (1)$$

With 'data', we denote the marker information and the disease phenotypes of the investigated persons for a given pedigree. ϕ comprises the parameters of the genotype-phenotype relation, which usually includes three genotype-specific penetrances in the case of a dichotomous trait, and the disease allele frequency. θ is the recombination fraction between the putative disease locus and the marker, which is related to the genetic distance. $\theta = 0$ means that the locus is located directly at the marker, $\theta = 1/2$ implies that there is no linkage between the disease locus and the marker. Since θ is unknown, the LOD score is maximized over $\theta \in [0, 1/2]$. The recombination fraction yielding the largest LOD score is the maximum likelihood estimate, which corresponds to the distance between the marker and the disease locus. In the multimarker situation, θ is replaced by the genetic position x of the disease locus.

The actual calculation of the LOD score is a complex task. A thorough description can be found in Strauch [19, p. 41]. In the context of the GENEHUNTER software, its heart is the Lander-Green algorithm [13]. The advantage of this algorithm is that its computational cost rises linearly with the number of markers in the data set. Unfortunately, there is an exponential dependency on the number of meioses in the used pedigrees. Each meiosis determines one bit in the inheritance vector. Therefore, GENEHUNTER can be applied to data sets with many markers, but using pedigrees with an inheritance vector of more than 17–18 bits leads to unacceptably high computation time and memory demands on a standard PC (Quad-Core CPU, 3.0 GHz, 8 GB RAM).

Since the parameters of the disease model are often unknown, it is conceivable to estimate them. This can be achieved by further maximizing the LOD score function given in equation (1) over the disease model parameters ϕ . Such a maximization of the LOD score is called MOD score analysis, the maximized LOD score is also called MOD score.

$$\text{MOD} = \max_{\phi} \max_{\theta \in [0, 1/2]} \log_{10} \left(\frac{P(\text{data}|\phi, \theta)}{P(\text{data}|\phi, \theta = 1/2)} \right). \quad (2)$$

Besides providing a maximum likelihood estimate of the disease model parameters and the recombination fraction, the MOD score also serves as the likelihood ratio test statistic for the null hypothesis of no linkage. When seen from the perspective of parametric linkage analysis as in equation (2), the MOD score calculation is done by maximizing the likelihood ratio over the trait model parameters ϕ , rather than maximizing each of the two likelihoods separately. For this reason, it has sometimes been argued whether the MOD score is in fact a real linkage test statistic. However, it has been shown for the special case of affected sib pairs [20] as well as for general pedigrees [21] that the likelihood ratio used in a parametric LOD or MOD score analysis is exactly the same as the nonparametric likelihood ratio using the probabilities z_i that a family falls into a certain allele-sharing class. For example, the possible triangle test for affected sib pairs [22] includes the probabilities z_0 , z_1 or z_2 that an affected sib pair shares 0, 1 or 2 alleles identical by descent, respectively. Importantly, the parameters z_i are fixed to their values under H_0 in the denominator likelihood. Therefore, the nonparametric allele sharing-based test statistic, in which a maximization is performed over the z_i , represents a true likelihood ratio test for the null hypothesis of no linkage. Because of the identity of the parametric and the nonparametric likelihood ratio, the same also holds for the parametric MOD score.

In the dichotomous case, the asymptotic distribution of the MOD score under the null hypothesis is known for special pedigree structures [21], and the corresponding p value can be easily obtained. In most cases, and particularly for quantitative phenotypes, the null distribution is unknown and the p value has to be determined empirically. The approach we use to determine p values had already been implemented in GENEHUNTER-MOD-SCORE, which performs MOD score analysis for dichotomous traits (see [18] for a thorough description). The empirical p value is calculated by simulating data sets under the null hypothesis: the original data set is taken, and founder alleles at the marker loci are assigned randomly, taking the corresponding marker allele frequencies into account. The inheritance of alleles to the next generation is simulated with respect to the recombination values for the intermarker distances, also keeping the pedigree structure of the original data set. Then the MOD score of the simulated data is calculated and compared to the MOD score of the original data. This is repeated with many simulated marker data sets. We denote these simulated data sets with H_0 replicates.

The p value is the probability of obtaining the original MOD score, or a larger one if the null hypothesis is valid. It is approximated by the empirical p value:

$$p = \frac{n}{N}, \quad (3)$$

where N is the total number of H_0 replicates, and $n \leq N$ the number of H_0 replicates with a MOD score larger than or equal to the original MOD score. If p is sufficiently small, the null hypothesis of no linkage is rejected.

Approach

So far, parametric linkage analysis with GENEHUNTER has only been possible for dichotomous phenotypes. The disease

model consists of the disease allele frequency p and the penetrances. The latter are the probabilities of suffering from the (dichotomous) disease depending on the genotype. We denote the wild-type alleles with '+' and the mutant alleles that increase the disease risk with 'm'.

Without imprinting effects, we thus have three possible genotypes at the disease locus: homozygous wild-type (+,+), heterozygous (m ,+) and homozygous mutant (m , m).

Our intention was to develop a model for quantitative traits that provides as much information as possible about the phenotype. For instance, the average value of the phenotype of a person with a certain genotype ((+,:), (m ,+), (m , m)) as well as the average variation of phenotype values of persons with the same genotype are of interest. Under the assumption that, given a certain genotype, the phenotype is normally distributed with mean μ_i and standard deviation σ_i , the density function is

$$f_i(x) := \frac{1}{\sqrt{2\pi}\sigma_i} \exp \left(-\frac{(x - \mu_i)^2}{2\sigma_i^2} \right), \quad (4)$$

with $i \in \{(+,:), (m,+), (m,m)\}$. A way to interpret this model would be to assume that one gene locus has a major effect on the phenotype, determining the value μ , depending on the genotype. Other genetic or environmental effects are considered by σ , thus creating 'noise' around the average phenotype value μ . Hence, σ comprises polygenic and environmental variance. It is assumed that these factors causing residual variance are independent between individuals.

The influence of other minor genes may depend on the genotype at the considered disease locus. In addition, environmental factors may also have a different influence, depending on the genotype at the disease locus. Therefore, it can be useful to model a separate 'noise' parameter for every genotype.

All in all, we have three density functions, one for each of the possible genotypes, each with its own expectation value μ_i and its standard deviation σ_i . Each density function describes the distribution of the quantitative phenotype of persons with one specific genotype: (+,:), (m ,+) or (m , m) (see also fig. 1).

These density functions describe the probability of a person with a certain genotype having a certain phenotype. They are thus equivalent to the penetrances in the dichotomous case. Therefore, we adapted the GENEHUNTER version for dichotomous traits to the quantitative case by replacing the penetrances with the density functions.

With this adaptation, it is now possible to calculate the LOD score for every genetic position in the genome for a given quantitative model. Then, the position with the highest LOD score will be our guess for the disease locus.

As in the dichotomous case, the model parameters of the quantitative phenotype (here: μ_i , σ_i and disease allele frequency p) are often unknown and need to be estimated. This is done by maximization of the LOD score, which is a function of the vector w of the phenotype parameters:

$$g: \mathbb{R}^7 \rightarrow \mathbb{R}, \quad w \mapsto g(w), \quad (5)$$

with

$$w = (\mu_{(+,+)}, \mu_{(m,+)}, \mu_{(m,m)}, \sigma_{(+,+)}, \sigma_{(m,+)}, \sigma_{(m,m)}, p). \quad (6)$$

Of course, $g(w)$ is the quantitative LOD score for the given parameter vector w . We also assume that it has already been maxi-

mized over the recombination fraction θ or the genetic position x in the multipoint case.

The best estimate for the parameter vector w is the one that yields the highest LOD score because this best explains the data. We therefore need to maximize g over w . A first idea would be to find the maximum analytically. This would be by far the computationally most efficient and elegant solution. Unfortunately, the arising equations have no analytical solution. This can easily be seen by looking at the algorithm as described by Strauch [19] that is represented by g . Therefore, we need to find a numerical approximation of the maximum of g . Modern mathematics provides many different methods for multidimensional optimization. The point is to choose the right method for the problem.

Since the density functions f_i are differentiable in μ_i and σ_i , g is differentiable in w . This feature has to be exploited when choosing the right optimization method. Therefore, gradient-based methods might be of use. The optimization should take place on a bounded space because there are some natural restrictions for the phenotype parameters. The affected allele frequency p is a probability and therefore lies in $[0, 1]$. Then, for all genotypes, $\sigma_i > 0$. In general, an upper boundary for σ_i can also be found without restricting the problem. For instance, the difference between the largest and smallest phenotypes in the data set is a reasonable upper limit since the average variation of the phenotype is unlikely to be higher than that. Theoretically, the expectation values μ_i can have any values in $\mathbb{R}$. Again, it is no severe restriction to assume that they lie within the range of the smallest and largest phenotypes or any similar interval. Thus, the optimization takes place in a multidimensional cuboid

$$Q := [l_1, u_1] \times [l_2, u_2] \times \dots \times [l_7, u_7] \subset \mathbb{R}^7. \quad (7)$$

The variables l_j and u_j denote the respective lower and upper boundaries. Then,

$$g: Q \mapsto \mathbb{R}, \quad g \in C^2(Q) \quad (8)$$

is the LOD score function on the phenotype parameter set Q . We are looking for

$$g_{\max} := \max_{w \in Q} g(w), \quad (9)$$

and

$$w_{\max} \in Q \text{ with } g(w_{\max}) = g_{\max}. \quad (10)$$

Then $w_{\max}$ is the estimate of the phenotype parameters.

We know that Q is a compact subset of $\mathbb{R}^7$. Also, g is differentiable, in particular continuous. Therefore, basic analysis tells us that a maximum of g on Q exists.

A requirement of the optimization method is thus that it is able to operate on a bounded subset of $\mathbb{R}^7$. Otherwise, a lot of computational effort would be wasted on testing parameter sets that could have been outruled by simple logic. Another crucial point is the computational cost. As mentioned before, the calculation of the LOD score depends exponentially on the number of persons in the families which are analyzed. Depending on the actual size of the pedigrees, calling the LOD score function g may well take several seconds wallclock time. With these considerations, we can narrow down the amount of optimization methods that come into consideration for our maximization problem. We implemented and tested those which remained, and chose the method that was fastest and converged most often.

Algorithm

It is a convention to formulate minimization rather than maximization problems. However, this is no loss of generality since the maximum of a function h can be found by applying the minimizing algorithm to $-h$. Therefore, we will also stick to this convention.

Gradient-free methods like the downhill simplex algorithm (also known as Nelder-Mead algorithm) almost never converged on our test data sets. Additionally, it did not improve the result of the starting value much. The simulated annealing method needed too many function evaluations to be of practical use. Among gradient-based methods, we tested the Spectral Projected Gradient method [23] and the GENCAN method [24], which had both acceptable results concerning optimization results as well as calculation time. However, these methods were outperformed with regard to both points by the Projected Gradient method (PGRAD). A thorough description can be found in Kelley [25]. A combination of several optimization methods turned out to be infeasible due to computation time.

We therefore chose the PGRAD method. The algorithm is a natural extension of the steepest descent algorithm to problems on simple bounded subsets. The basic idea is to choose an initial point $w_0 \in Q$ that is already close to the minimum. The algorithm then calculates the direction with the lowest derivation ('steepest descent') and follows this direction into the minimum. In the following section, we will give a description of the optimization algorithm. Afterwards, we will present how we chose the starting value w_0 .

Optimization

As already mentioned, we have to choose a starting value w_0 . From this point, we construct a sequence w_i that hopefully converges into the minimum. In order to describe the construction, we need to explain the following terms: by

$$\nabla h(w) = \left(\frac{\partial h}{\partial w_{(1)}}(w), \dots, \frac{\partial h}{\partial w_{(7)}}(w) \right) \quad (11)$$

we denote the gradient of h in w . It describes the direction of the steepest ascent of h in w . $w_{(j)}$ is the j -th component of the vector w , not the j -th element of the sequence. Then we have that

$$P(w)_j = \begin{cases} l_j & w_{(j)} \leq l_j \\ w_{(j)} & l_j < w_{(j)} < u_j \\ u_j & w_{(j)} \geq u_j \end{cases} \quad (12)$$

is the projection onto Q . This maps w to the nearest point in Q . Given an iterate w_i , the new iterate is

$$w_{i+1} = P(w_i - \lambda_i \nabla h(w_i)). \quad (13)$$

We call $\lambda_i > 0$ the steplength parameter. Since $-\nabla h(w_i)$ determines the direction of the steepest descent, we can interpret equation (13) as a step of length $\lambda_i \|\nabla h(w_i)\|$ in the direction of the minimum. λ_i is determined by the so-called sufficient decrease condition. We define

$$w(\lambda_i) = P(w_i - \lambda_i \nabla h(w_i)). \quad (14)$$

Then λ_i satisfies the condition of sufficient decrease if

$$h(w(\lambda_i)) - h(w) \leq \frac{-\alpha}{\lambda_i} \|w - w(\lambda_i)\|^2. \quad (15)$$

α is a parameter that can be adapted to the particular function one wishes to minimize. Typically, it is set to 10^{-4} .

Next, we need to find a termination criterion. In the unconstrained case, a necessary condition for a local minimum is that $\|\nabla h\|$ is zero. With constrained problems, this is not necessarily the case. A natural substitution is to end the iteration if the difference of consecutive iterates is small: $\|w_{i+1} - w_i\| < \varepsilon$. In order to prevent the method from running on and on when no minimum can be found, an additional termination criterion is set: the optimization is stopped after $n_{\max}$ iterations. When applying the algorithm to the LOD score optimization problem, we set $h = -g$.

Starting Values

A disadvantage of gradient-based methods is that, in general, convergence to the global minimum cannot be guaranteed. The method will iterate from the starting point to the next minimum, which might well be a local minimum and not the global one. Another problem is that the duration and the computational cost of the iteration depend on the initial value and its closeness to the minimum.

We deal with the first issue by repeating the iteration with several different starting values. In this way, we increase the probability to find the global minimum since not all starting values will lie close to the same local minimum. In order to obtain a starting value that is already close to the minimum, we roughly guess the parameters of the phenotype and use them as a first starting value. In general, carrying one mutant allele ($m, +$) leads to an average increase or decrease of the phenotype, compared to the homozygous wild-type genotype ($+, +$). A second mutant allele (m, m) leads to a further increase or decrease (see also fig. 1). We can use this for the following approach: $q_1, \dots, q_m$ denotes the ordered set of all phenotypes in the data set such that $q_i \leq q_{i+1} \forall i$. We assume that the mutant allele m increases the phenotype. We initially allocate the lowest third of the (ordered) phenotypes to the ($+, +$) genotype, the second third to the ($m, +$) genotype, and the highest third to the (m, m) genotype. We then calculate the empirical mean and standard deviation for the respective genotype:

$$\mu_{(+, +)} = \frac{3}{m} \sum_{i=1}^{\frac{m}{3}} q_i, \quad \sigma_{(+, +)}^2 = \frac{3}{m} \sum_{i=1}^{\frac{m}{3}} (q_i - \mu_{(+, +)})^2. \quad (16)$$

The other parameters are calculated with the respective thirds of the phenotypes. The disease allele frequency p is set to any random number in $[0, 1]$, preferably some value close to zero.

Other starting values can be obtained by multiplying the starting values in equation (16) with a positive random number r . It should not vary too much from 1.0 in order to keep the basic idea of equation (16) but still provide a different spot to start with, so that the optimization will avoid local minima. We suggest $r \in [0.5, 1.5]$. We refrain from using different random numbers of r_i for the starting values in equation (16) since this could destroy the assumed inequality $\mu_{(+, +)} \leq \mu_{(m, +)} \leq \mu_{(m, m)}$.

Phenotype Models

In special cases, it can be useful to adapt the model formulation described in the section Approach. For instance, if it is sus-

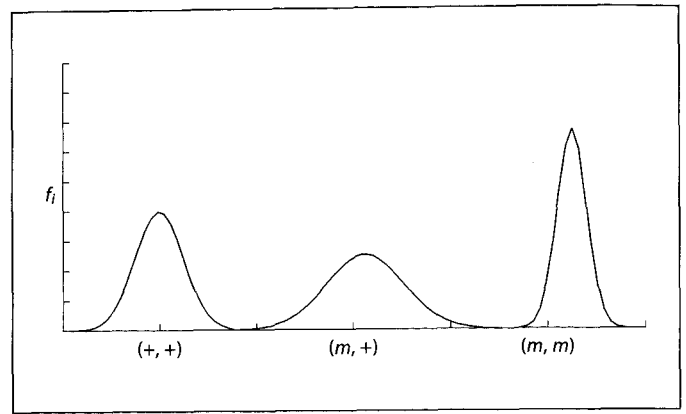


Fig. 1. Density functions of the phenotype depending on the genotype.

pected that the disease is determined by imprinting effects, it is useful to include this in our model as a separate parameter. In this case, ($+, m$) means that the wild-type allele is transmitted by the father and the mutant allele by the mother. On the other hand, a person with a ($m, +$) genotype has a paternal mutant allele and a maternal wild-type allele. Accordingly, the quantitative model in GENEHUNTER-QMOD can be extended with a fourth normal distribution for the ($+, m$) genotype. Then, we have two additional parameters, the expectation value $\mu_{(+, m)}$ and the standard deviation $\sigma_{(+, m)}$. The parameter vector w (see equation (6)) is extended by these two parameters, and the optimization process is carried out analogously to the nonimprinting case.

We may also know that the residual variation is independent of the phenotype, i.e. all standard deviations of the phenotype distributions in equation (4) are equal. In this case, we can adapt our model in GENEHUNTER-QMOD by setting

$$\sigma_{(+, +)} = \sigma_{(m, +)} = \sigma_{(m, m)}. \quad (17)$$

Again, the optimization is carried out analogously to the ordinary case.

Another adaptation is the modeling of additivity, i.e. the absence of dominance effects. This means that the average phenotype of heterozygous persons is exactly between the average phenotype of homozygous mutant and wild-type individuals:

$$\mu_{(m, +)} = \frac{1}{2} (\mu_{(+, +)} + \mu_{(m, m)}). \quad (18)$$

Since this is a characteristic of many inherited traits, we decided to implement the purely additive model in GENEHUNTER-QMOD, too.

Validation

In order to assess the performance of PGRAD optimization using the quantitative MOD score as test statistic, we evaluated the type I error and power of the method as implemented in GENEHUNTER-QMOD. Furthermore, the estimates of the parameters of the phenotype model were of interest, i.e. the expectation values μ_i and standard deviations σ_i for the three genotypes (see fig. 1).

For all data sets we analyzed, we used ten different starting values as described in section Starting Values. The maximum number of iterations of the PGRAD algorithm was set to 500. A higher number of iterations led only to improvements on the sixth decimal place of the MOD score if the PGRAD method had not converged before.

Type I Error Rate Estimation

For the simulations, we used the program SIBSIM developed by Franke et al. [26], which allows us to create pedigree files with any given number of families and any kind of family structure. The position of the disease locus can also be specified. We made some small modifications to the software so that we could explicitly specify the genotype-specific expectation values μ_i and the standard deviations σ_i of the phenotype. SIBSIM then creates the appropriate pedigree file and assigns to each person a random quantitative phenotype according to the respective density functions (fig. 1). The genotypes of the founders are set randomly under the assumption of Hardy-Weinberg equilibrium.

In order to estimate the type I error of a certain scenario, we simulated 500 pedigree files under the null hypothesis of no linkage. Each pedigree file contained 300 sib pair pedigrees, i.e. families with two parents with two children. We used one highly informative marker with 20 equipotent alleles. In the analysis, the genotypes for both parents and siblings were available, as were all phenotypes. We set the genetic distance between marker and disease locus to 1,000 cM, which equates to a recombination fraction θ which deviates from 1/2 by less than 10^{-90} (Haldane mapping function), so that we can safely assume no linkage in the data. We then calculated the MOD score of each data set. Since the asymptotic distribution of the MOD score is unknown in the case of quantitative traits, we calculated the p value of each simulated data set empirically by performing a further round of simulations.

Fortunately, the previous program extension GENEHUNTER-MODSCORE, which calculates MOD scores for dichotomous traits [18], was already able to perform the p value calculation. The MOD score distribution was simulated by creating pedigrees with the same structure and the same phenotypes as the original data set, randomly generating marker data under the null hypothesis of no linkage. These replicates of the inner simulation loop are therefore called H_0 replicates (see section Overview). For each of the 5,000 H_0 replicates which were generated in this manner, the MOD score was calculated and compared to the MOD score of the original simulated data set. The p value equals the relative amount of MOD scores reaching or surpassing the MOD score of the original data, i.e. the empirical probability under the null hypothesis to obtain an equal or larger test statistic.

We decided to use a nominal α level of 5%. Therefore, we counted all the data sets simulated by SIBSIM with a p value ≤ 0.05 . The relative amount of these data sets is the empirical probability to reject the null hypothesis if the null hypothesis is valid. Therefore, it is the type I error rate.

Power Estimation

The estimation of the power was done in a similar way as the type I error estimation. We simulated data with SIBSIM under the assumption of complete linkage. Therefore, we set the genetic distance between marker and disease locus to 0 cM. The number of families and the family structure remained as introduced in section Type I Error Rate Estimation. The MOD score and the cor-

responding p value were calculated by means of an inner loop of simulations under H_0 , as before. This time, a p value ≤ 0.05 denoted a justly rejected null hypothesis. Therefore, the relative amount of data sets with p values ≤ 0.05 is the estimate of the power of the scenario.

Sample Method

In order to obtain a precise estimate of the p value, many H_0 replicates have to be created, and the corresponding MOD score has to be calculated. This is a huge computational effort. To calculate a p value by using 5,000 H_0 replicates may easily take more than 10 days (standard PC), which is too long for practical application. Fortunately, a good approximation of the p value can be obtained without performing a full MOD score calculation for every H_0 replicate. When calculating the p value, we only need to know whether the MOD score of the H_0 replicate is larger than the MOD score of the original data or not. Simulations showed that this information is highly correlated with the average LOD score value g on Q . In particular, we took 500 randomly chosen function evaluations of g for a H_0 replicate and calculated the arithmetic mean. Then we did the same for the original data. If the arithmetic mean of the H_0 replicate was higher/lower than the corresponding arithmetic mean of the original data, the same was true for the MOD scores. Since the evaluation of such a random sample takes only a fraction of the time of the PGRAD optimization, we can calculate the p value of a H_0 replicate much faster. We call this technique sample method.

Parameter Estimation

We estimated the genotype-phenotype relation by using the parameter set that yielded the PGRAD-maximized LOD score, i.e. the MOD score. We compared the distribution of parameter estimates obtained for the simulated replicates with the true generating value.

Comparison to Other Methods

In addition, the type I error rate and power for all scenarios have also been analyzed with the VCA as implemented in GENEHUNTER (see [27] for more information). LOD scores obtained by the VCA, when calculated without a dominance component as we did here, asymptotically follow a 50:50 mixture of a χ^2 distribution with one degree of freedom and a point mass at zero. In addition, all scenarios have been analyzed with the traditional Haseman-Elston method as implemented in GENEHUNTER.

Results

The results are given in four parts. First, the type I error and power results for different scenarios with randomly ascertained data are described, followed by the results for scenarios with nonrandom ascertainment. Then, the results of the phenotype parameter estimation are given. In the end, we present a brief application of the GENEHUNTER-QMOD to a real data set, namely the investigation of human family data regarding sensitization to house dust mite allergens.

Type I Error and Power

In this section, we only investigate scenarios with random ascertainment. First, we wanted to know how our method dealt with a normally distributed phenotype. Scenario 1 addresses the ‘simple case’ without dominance effects and equal variances (table 1). As for the simulations, the analysis with GENEHUNTER-QMOD was performed assuming no dominance and the same residual variances for all three genotypes. The type I error rate estimation showed that the sample method rejects H_0 with a probability of 0.06 and the VCA with 0.062 (table 1). Therefore, both methods are slightly above the nominal type I error rate of 0.05. However, based on the number of replicates we used for the calculation of the type I error (500), this is not statistically different from the nominal level. Haseman-Elston regression shows a type I error rate of 0.05 and is therefore perfectly accurate.

Since our model was specifically designed for phenotypes with different variances, another scenario was also created without dominance effects, but with genotype-dependent residual variances. This time, the analysis with GENEHUNTER-QMOD was performed assuming no dominance but different residual variances for all genotypes. Here, the VCA and Haseman-Elston regression reject H_0 with a probability of 0.04 and 0.038, respectively. However, given that 500 replicates were simulated, these values are not significantly different from the expected value of 0.05. The sample method shows a true type I error rate of 0.024 and is therefore conservative.

Since not all quantitative phenotypes are normally distributed, we investigated how GENEHUNTER-QMOD, VCA and Haseman-Elston regression perform with data for which the assumption of normality does not hold. For the next scenario, we therefore simulated the phenotypes for each genotype with lognormal distributions:

$$f_i(x) = \begin{cases} \frac{1}{\sqrt{2\pi}\sigma_i x} \exp\left(-\frac{(\ln x - \mu_i)^2}{2\sigma_i^2}\right) & x > 0 \\ 0 & x \leq 0 \end{cases} \quad (19)$$

That is, the phenotype, given genotype i , had a specific lognormal distribution, with parameters μ_i and σ_i^2 . The sample method has a type I error rate of 0.036, and Haseman-Elston regression a type I error rate of 0.034 (see table 1, scenario 3). Again, both type I error rates do not differ significantly from the expected value of 0.05. The type I error rate of the VCA is 0.134 and therefore exceeds by far the nominal level.

Table 1. True type I error rate comparison of the sample method, the VCA and Haseman-Elston regression

Scenario	$\sigma^2_{(+,+)}$	$\sigma^2_{(m,+)}$	$\sigma^2_{(m,m)}$	α_S	α_{VCA}	α_{HE}
1	200.0	200.0	200.0	0.060	0.062	0.050
2	100.0	225.0	400.0	0.024	0.04	0.038
3	2.25	6.25	9.0	0.036	0.134	0.034
4	350.0	350.0	350.0	0.046	0.042	0.042

α_S , α_{VCA} and α_{HE} denote the true type I error rate of the sample method, the VCA and Haseman-Elston regression, respectively. Note that the parameters of scenario 3 correspond to lognormal distributions. The expectation values $\mu_{(+,+)}$, $\mu_{(m,+)}$ and $\mu_{(m,m)}$ were set to 20.0, 40.0 and 60.0, respectively, in scenario 1, 2 and 4. In scenario 3, they were set to 20.0, 100.0 and 180.0, respectively. Note that each family in scenario 4 consists of two parents and four siblings instead of two. The disease allele frequency p was set to 0.4 in scenarios 1, 2 and 3. In scenario 4, we have $p = 0.1$. In each scenario, the type I error was calculated using 500 replicates simulated under the null hypothesis of no linkage.

Table 2. Power of the sample method, the VCA and Haseman-Elston regression

Scenario	P_S	P_{VCA}	P_{HE}
1	0.82	0.98	0.915
2	0.874	0.95	0.745
3	0.645	0.295	0.13
4	0.82	0.77	0.225

P_S is the power of the sample method, P_{VCA} is the power of the VCA and P_{HE} denotes the power of Haseman-Elston regression. The parameters of the scenarios are as in table 1. In each scenario, the power was calculated using 200 replicates simulated under the assumption of complete linkage.

Next, we examined scenarios with families consisting of two parents and four siblings. In order to avoid having a power of 1 for all types of analysis, we increased the residual variance σ_i^2 to 350.0 for all three genotypes in the simulation. All three methods yield an acceptable type I error rate, although with a slight nonsignificant tendency towards conservativeness. The sample method comes closest to the theoretical value of 0.05 (table 1, scenario 4).

When we analyzed the power, we created 200 data sets for each scenario under complete linkage. Again, GENEHUNTER-QMOD calculated 5,000 H_0 replicates in order to obtain the p value for each original simulated data set. All results are shown in table 2. In scenario 1 with equal

Table 3. True type I error rate and power of the sample method, the VCA and Haseman-Elston regression when using an ascertainment condition

Scenario	α_S	α_{VCA}	α_{HE}	P_S	P_{VCA}	P_{HE}
5	0.066	0.0	0.06	0.895	0.04	0.585
6	0.032	0.088	0.052	0.94	0.89	0.03
7	0.042	0.0	0.052	0.975	0.0	0.495

α_S , α_{VCA} and α_{HE} denote the true type I error rate of the sample method, the VCA and Haseman-Elston regression, respectively. P_S , P_{VCA} and P_{HE} denote the empirical power of the three methods. The expectation values in scenarios 5, 6 and 7 were set to $\mu_{(+,+)} = 20.0$, $\mu_{(m,+)} = 40.0$ and $\mu_{(m,m)} = 60.0$, and the variances to $\sigma^2_{(+,+)} = \sigma^2_{(m,+)} = \sigma^2_{(m,m)} = 200.0$. The disease allele frequency in the generated data sets from which the families were drawn was set to 0.1. In each scenario, the type I error was calculated using 500 replicates simulated under the null hypothesis of no linkage, and the power was calculated using 200 replicates simulated under the assumption of complete linkage.

variances, the power of the VCA is by 0.16 higher compared to the sample method. The power difference between the sample method and the VCA is reduced to 0.076 when analyzing scenario 2 with different variances. Note that the VCA assumes equal variances, while the sample method/PGRAD optimization does not. On the other hand, in the lognormal case (scenario 3), the sample method has 0.35 power gain compared to the VCA method. In the case of different variances, the sample method has a moderately higher power compared to the scenario with equal residual variances. Haseman-Elston regression performs well for normally distributed data, although not as well as the VCA. Different variances in the phenotype result in a power loss of 17% (91.5 vs. 74.5%). However, Haseman-Elston regression has the lowest power of the three methods when non-normally distributed data are analyzed, i.e. being 0.13 in scenario 3.

When pedigrees with four instead of two sibs are analyzed, the sample method outperforms both the VCA and Haseman-Elston regression. The power of the sample method in scenario 4 remained at 0.82, whereas the power of the VCA dropped to 0.77, and the power of Haseman-Elston regression even to 0.225.

Ascertainment

Since the recruitment of families is often carried out with a certain selection scheme, we also investigated scenarios with nonrandom ascertainment. Again, we simulated scenarios with 200 data sets for each scenario for

power analysis and 500 data sets for type I error analysis. Family structure, marker and disease locus as well as number of H_0 replicates for p value calculation remained as in scenarios 1, 2 and 3. Along the lines of Kleensang et al. [28], we imposed the following ascertainment conditions. Scenario 5 contains only families with at least one sibling in the highest quartile of the original distribution (single-proband selection). In scenario 6, families were only selected if siblings were either both in the highest or both in the lowest quartile. Scenario 7 consists of families with one sibling in the highest and the other sibling in the lowest quartile (double-proband selection). Families were simulated using SIBSIM and drawn until each data set contained 300 families with the corresponding structure according to the ascertainment criterion. The results are shown in table 3.

The results show that the estimate of the type I error for the sample method tends to be liberal under the single-proband selection (scenario 5), and to be conservative under the double-proband selection (scenarios 6 and 7). However, deviations from the expected value of 0.05 cannot be deemed significant given the number of replicates (500) used to assess the empirical type I error rate. The VCA is extremely conservative under the single- and double-proband selection with both siblings in opposite quartiles. However, it is too liberal under the double-proband selection with both siblings in the same quartile. By contrast, Haseman-Elston regression has acceptable type I error rates in all ascertainment scenarios. The sample method shows good power in all ascertainment scenarios. It is especially strong under the double-proband selection. The power of the VCA is unacceptably low under the single- and double-proband selection with siblings in different quartiles (scenarios 5 and 7). It should be noted that these are the scenarios for which the VCA is also strongly conservative. However, it performs well for the double-proband selection with siblings in the same quartiles (scenario 6), although at the price of being liberal, and it still does not quite reach the power of the sample method.

The power of Haseman-Elston regression shows a reversed picture compared to the VCA. While it still performs moderately in scenarios 5 and 7, it has almost no power in scenario 6. However, even in scenarios 5 and 7, it fails to reach the power of the sample method by far.

Parameter Estimation

Another intention of the new method implemented in GENEHUNTER-QMOD was to gather specific information about the genotype-phenotype relation. Therefore, the next point we focused on was the estimation of the

Table 4. Parameter estimates for scenarios 1 and 6

Parameter	True value	Mean	Standard deviation	Error, %
Scenario 1				
$\mu_{(+,+)}$	20.0	15.1	10.7	24.5
$\mu_{(m,+)}$	40.0	37.4	8.8	13.0
$\mu_{(m,m)}$	60.0	59.6	17.6	2.0
σ	14.1	7.7	2.74	45.4
Scenario 6				
$\mu_{(+,+)}$	20.0	12.8	7.0	36.0
$\mu_{(m,+)}$	40.0	36.5	8.2	17.5
$\mu_{(m,m)}$	60.0	60.3	16.5	1.5
σ	14.1	5.4	3.24	61.7

σ is the standard deviation of the genotype-specific density functions: $\sigma = \sigma_{(+,+)} = \sigma_{(m,+)} = \sigma_{(m,m)}$, see table 1.

parameters of the quantitative trait model; that is, the expectation values $\mu_{(+,+)}$, $\mu_{(m,+)}$ and $\mu_{(m,m)}$ and the standard deviations $\sigma_{(+,+)}$, $\sigma_{(m,+)}$ and $\sigma_{(m,m)}$ of the respective genotype-specific distribution of the phenotype. As already mentioned before, the parameter estimates were the arguments of the model-maximized LOD score function.

We ran the PGRAD method for each of the 200 data sets simulated in scenarios 1 and 6. The results of the parameter estimation are given in table 4. The first column shows the true value for the corresponding parameter. The second column gives the empirical mean of all 200 estimates for the corresponding parameter. The third column shows the empirical standard deviation. For the expectation parameters, the error term (last column) describes the difference of the true value and the mean estimate in relation to the difference of the phenotype means μ_i (in this case 20.0). In this way, we avoid error terms which are small due to high phenotypes and not due to precise estimates.

In order to give the reader an idea of the distribution of the parameters, we have plotted the distribution of the expectation value of the homozygous mutant genotype in scenario 1 (fig. 2). We can see a clear peak at the true value $\mu_{(m,m)} = 60.0$.

Application to Real Data

We applied GENEHUNTER-QMOD to human pedigree data regarding sensitization to house dust mite allergens. The data set consists of sib pairs as well as more complex pedigrees, and includes several European populations. Here, we analyzed German and English families.

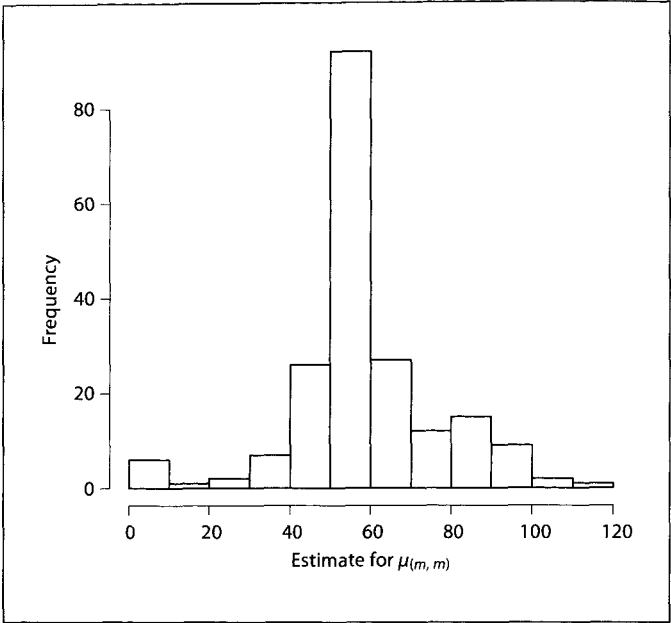


Fig. 2. Histogram of 200 parameter estimates of $\mu_{(m,m)}$ (true value = 60.0) for scenario 1.

The German population is made up of 44 families with 187 individuals, thereby being 33 families with two sibs and 11 families with three sibs. The English population consists of 19 families with 122 individuals, including 7 families with two sibs, 3 families with three sibs, 5 families with four affected sibs, and 4 families with an extended structure.

For each nonfounder, a RAST (Radio-Allergo-Sorbent Test) class test was carried out by measuring specific IgE antibodies to crude extract of *Dermatophagoides pteronyssinus* (house dust mite). According to the value of the IgE level, the individual was assigned a RAST class, ranging from 0 to 6. This value was taken as the quantitative phenotype. As for marker information, 604 microsatellite markers were typed on all chromosomes. The results of the original analyses for the dichotomous phenotype mite sensitization have been described previously [29, 30].

The analysis with GENEHUNTER-QMOD yielded a MOD score of 1.019 on chromosome 4, at marker D4S194, for the German population, with a corresponding p value of 0.00124. The p value was calculated by generating 50,000 replicates under the null hypothesis.

For the English population, a MOD score of 1.193 occurred at chromosome 5, at marker D5S486. The corresponding p value is 0.001. It was calculated by generating 1,000 replicates under the null hypothesis.

These two genetic loci were already identified in a previous study with dichotomous phenotypes [15]. Here, peak MOD scores occurred for the German population at D4S430 (which is directly next to D4S194) and for the English population at D5S416 (which is also next to D5S486).

Discussion

Our intention was to develop a statistical method for quantitative linkage analysis which provides as much information about the phenotype as possible. To this end, we developed the PGRAD optimization method and implemented it in the Lander-Green-based software package GENEHUNTER-QMOD.

For randomly ascertained data that match the assumption of normally distributed phenotypes, the power of GENEHUNTER-QMOD stays behind that of the VCA when sib pairs are analyzed. However, when larger families with four children are analyzed, the sample method outperforms both Haseman-Elston regression and the VCA. In our example, the sample method was able to compensate a stronger noise parameter (i.e. the residual variances) with the additional information given in families with four siblings, while the VCA and especially Haseman-Elston regression had a substantial power drop. It seems that the latter methods do not exploit the additional information that is given in extended pedigrees to the same degree as GENEHUNTER-QMOD. In particular, Haseman-Elston regression only looks at one pair of sibs at a time, even when larger families are analyzed.

Also, when non-normal data are analyzed, GENEHUNTER-QMOD turns out to be more robust than the other methods. Our simulations showed that the power of the scenario with lognormally distributed data was by 0.35 higher when analyzed with our new method compared to the VCA. This is remarkable given that our model specifically models a normally distributed trait. At this point, it should be noted that the assumption of a normal distribution also underlies the VCA. It seems that, due to the variety of modeling possibilities for the data, our method provides a certain robustness with regard to violations of the normality assumption that the VCA method is lacking. When no ascertainment condition was applied, the power of the traditional Haseman-Elston method always stayed behind the power of the VCA, being higher than that of the sample method only in the 'standard' case of normally distributed data with equal variances. Therefore, the VCA should be preferred to Hase-

man-Elston regression when one can be sure of normally distributed data if unselected, population-based samples are analyzed.

For non-normally distributed data, GENEHUNTER-QMOD has both higher power and a smaller type I error rate than the VCA. Therefore, if one cannot safely assume normally distributed phenotypes, GENEHUNTER-QMOD might be preferred to the VCA as well as to Haseman-Elston regression.

Simulations under the null hypothesis showed that GENEHUNTER-QMOD roughly keeps the α level. This holds even if we use non-normally distributed data, whereas the VCA is only able to control the type I error if normality holds. Modifications or extensions of the VCA method are conceivable, e.g. applying a transformation, using a robust sandwich estimator for the variance, or employing a method of simulating markers under the null hypothesis of no linkage similar to the one used in this work, as done by Chen et al. [31]. We have found that Haseman-Elston regression shows a slight tendency towards conservativeness for 'difficult' data, meaning non-normally distributed phenotypes or phenotypes with genotype-dependent residual variances. However, given that only 500 replicates could be simulated due to otherwise excessive computational demands (since the p value for each replicate needs to be determined by another round of simulations), these differences from the expected value of 0.05 are not significant. Haseman-Elston regression yields a correct type I error for the standard case with a normally distributed phenotype and equal variances.

For data obtained with an ascertainment condition, the sample method shows superior power in all scenarios. The VCA fails completely in the case of single- and double-proband selection with siblings in opposite quartiles. The power is acceptable only under the double-proband selection with siblings in the same quartiles, but this comes at the price of an inflated type I error. Therefore, one should refrain from the use of the VCA under these conditions. In all ascertainment scenarios, the power of Haseman-Elston regression stays clearly behind the power of the sample method, although having an accurate type I error. Thus, the sample method should be preferred to Haseman-Elston regression, too.

With regard to type I error and power analysis, our results of the VCA and Haseman-Elston regression are, to a large part, consistent with the findings of Kleensang et al. [28]. In that simulation study, the true type I error for the VCA also often deviates from the nominal type I error and is unacceptably liberal when the normality as-

sumption does not hold. Imposing ascertainment conditions also strongly influences the type I error in both directions. A low power under normality violations is described as well. Haseman-Elston regression is shown to be conservative under non-normally distributed data [28], as we can confirm with our results. The power is described as inferior to the power of the VCA when no ascertainment condition is applied, which is also consistent with our findings. At this point, it should be noted that the power analysis for every method in this paper has been based on a type I error level of 0.05. As has been shown by Morris and Elston [32], relative power comparisons of two different methods at $\alpha = 0.05$ do not necessarily lead to the same conclusion as when performed at lower type I error levels. This should be kept in mind when generalizing our findings.

An important feature of GENEHUNTER-QMOD is the parameter estimation. By assuming a separate normal distribution for each genotype, GENEHUNTER-QMOD estimates the variance and the mean of each distribution. Our simulations showed that parameter estimation works well for the means of the different normal distributions. Especially the estimated parameter for the (m,m) genotype had a small relative error. Unfortunately, the estimates of the standard deviations turned out to be too low. One should consider this when interpreting the results. These characteristics did not change substantially when ascertainment conditions were applied.

Conclusion

Linkage analysis continues to be an important method to identify and locate genes that cause or contribute to inherited diseases, especially in the era of large-scale sequencing. By detecting the corresponding genetic variant and identifying how it influences the trait, the disease mechanism can be elucidated, opening the possibility of a causal treatment. Also, a disease might be diagnosed even before its manifestation. In this way, preventive measures can be applied in good time.

With GENEHUNTER-QMOD, we add a method for the analysis of quantitative phenotypes to the variety of linkage tools. It provides both an inferential test for linkage and estimates of the phenotype parameters. The phenotype is modeled as a normally distributed variable, with parameters depending on the genotype. The estimates include both the means and the residual dispersions, specific for every genotype. If further information about the inheritance mode of the disease is available, one

can specify the disease model appropriately: imprinting, no dominance and equal residual variances for all genotypes can be taken into account.

Besides the identification of the gene per se, the gathered information might be useful in several ways. If the average phenotype changes substantially with the genotype, patients with different genotypes might require individualized treatments, such as a different dose of medication. The estimates of the disease model parameters might give further clues to the disease mechanism. If the variance changes with the genotype, the mutation at the disease locus might trigger other reactions that influence the phenotype, for instance other genes or environmental effects that now have a stronger impact. That is, heterogeneity of residual variances may hint to gene-gene or gene-environment interactions.

Because GENEHUNTER-QMOD is based on the Lander-Green algorithm, it can simultaneously use many markers in the analysis, which makes it applicable to gene mapping projects based on SNP arrays. GENEHUNTER-QMOD is freely available on the web at <http://www.helmholtz-muenchen.de/genepi/downloads>.

Acknowledgement

We would like to thank Dr. Thorsten Kurz and Prof. Dr. Klaus Deichmann, University Children's Hospital Freiburg, for providing the data on house dust mite allergy. In addition, we are grateful to the anonymous reviewers for their comments on a previous version of the manuscript. The simulations have been performed on the MARC cluster of the Philipps University Marburg, Marburg, Germany.

Funding

This work was supported by grant Str643/4-1 of the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation).

References

- 1 Kruglyak L, Daly MJ, Reeve-Daly MP, Lander ES: Parametric and nonparametric linkage analysis: a unified multipoint approach. *Am J Hum Genet* 1996;58:1347–1363.
- 2 Abecasis GR, Cherny SS, Cookson WO, Cardon LR: Merlin – rapid analysis of dense genetic maps using sparse gene flow trees. *Nat Genet* 2002;30:97–101.
- 3 Gudbjartsson DF, Jonasson K, Frigge M, Kong A: Allegro, a new computer program for multipoint linkage analysis. *Nat Genet* 2000;25:12–13.

- 4 Dudbridge F: A survey of current software for linkage analysis. *Hum Genomics* 2003;1: 63–65.
- 5 Falconer DS, Mackay FC: *Introduction to Quantitative Genetics*. Harlow, Pearson Education, 1996.
- 6 Camp NJ, Cox A: *Quantitative Trait Loci – Methods and Protocols*. Totowa, humanapress, 2002.
- 7 Almasy L, Blangero J: Multipoint quantitative-trait linkage analysis in general pedigrees. *Am J Hum Genet* 1998;62:1198–1211.
- 8 Haseman JK, Elston RC: The investigation of linkage between a quantitative trait and a marker locus. *Behav Genet* 1972;2:3–19.
- 9 Ziegler A, Kastner C: A Minimum distance estimation approach to estimate the recombination fraction from a marker locus in robust linkage analysis for quantitative traits. *Biom J* 1997;39:765–775.
- 10 Lathrop GM, Lalouel JM: Easy calculations of lod scores and genetic risks on small computers. *Am J Hum Genet* 1984;36:460–465.
- 11 Hasstedt SJ, Cartwright P: PAP-Pedigree Analysis Package. Technical Report No. 13. Technical report, Department of Medical Biophysics and Computing, University of Utah, 1981.
- 12 Elston RC, Stewart J: A general model for the genetic analysis of pedigree data. *Hum Hered* 1971;21:523–542.
- 13 Lander ES, Green P: Construction of multi-locus genetic linkage maps in humans. *Proc Natl Acad Sci USA* 1987;84:2363–2367.
- 14 Kong A, Cox NJ: Allele-sharing models: LOD scores and accurate linkage tests. *Am J Hum Genet* 1997;61:1179–1188.
- 15 Strauch K, Fimmers R, Kurz T, Deichmann KA, Wienker TF, Baur MP: Parametric and nonparametric multipoint linkage analysis with imprinting and two-locus-trait models: application to mite sensitization. *Am J Hum Genet* 2000;66:1945–1957.
- 16 Dietter J, Spiegel A, an Mey D, Pflug HJ, Al-Kateb H, Hoffmann K, Wienker TF, Strauch K: Efficient two-trait-locus linkage analysis through program optimization and parallelization: application to hypercholesterolemia. *Eur J Hum Genet* 2004;12: 542–550.
- 17 Dietter J, Mattheisen M, Fürst R, Rüschendorf F, Wienker TF, Strauch K: Linkage analysis using sex-specific recombination fractions with GENEHUNTER-MODSCORE. *Bioinformatics* 2007;23:64–70.
- 18 Mattheisen M, Dietter J, Knapp M, Baur MP, Strauch K: Inferential testing for linkage with GENEHUNTER-MODSCORE: the impact of the pedigree structure on the null distribution of multipoint MOD scores. *Genet Epidemiol* 2008;32:73–83.
- 19 Strauch K: Kopplungsanalyse bei genetisch komplexen Erkrankungen mit genomischem Imprinting und Zwei-Genort-Krankheitsmodellen. *Medizinische Informatik, Biometrie und Epidemiologie, Urban und Vogel*, 2002, vol 87.
- 20 Knapp M, Seuchter SA, Baur MP: Linkage analysis in nuclear families. 2: Relationship between affected sib-pair tests and lod score analysis. *Hum Hered* 1994;44:44–51.
- 21 Strauch K: MOD-score analysis with simple pedigrees: an overview of likelihood-based linkage methods. *Hum Hered* 2007;64:192–202.
- 22 Holmans P: Asymptotic properties of affected-sib-pair linkage analysis. *Am J Hum Genet* 1993;52:362–374.
- 23 Birgin EG, Martinez JM, Raydan M: Non-monotone spectral projected gradient methods on convex sets. *SIAM J Optim* 2000;10: 1196–1211.
- 24 Birgin EG, Martinez JM: Large-scale active-set box-constrained optimization method with spectral projected gradients. *Comput Optim Appl* 2002;23:101–125.
- 25 Kelley CT: *Iterative Methods for Optimization*. Philadelphia, SIAM Society for Industrial and Applied Mathematics, 1999.
- 26 Franke D, Kleensang A, Ziegler A: SIBSIM – quantitative phenotype simulation in extended pedigrees. *GMS Med Inform Biom Epidemiol* 2006;2:Doc02.
- 27 Pratt SC, Daly MJ, Kruglyak L: Exact multipoint quantitative-trait linkage analysis in pedigrees by variance components. *Am J Hum Genet* 2000;66:1153–1157.
- 28 Kleensang A, Franke D, Alcais A, Abel L, Müller-Myshok B, Ziegler A: An extensive comparison of quantitative trait loci mapping methods. *Hum Hered* 2010;69:202–211.
- 29 Kurz T, Strauch K, Heinzmann A, Braun S, Jung M, Rüschendorf F, Moffatt MF, Cookson WOCM, Inacio F, Ruffilli A, Nordskov-Hansen G, Peltre G, Forster J, Kuehr J, Reis A, Wienker TF, Deichmann KA: A European study on the genetics of mite sensitization. *J Allergy Clin Immunol* 2000;106:925–932.
- 30 Kurz T, Altmueller J, Strauch K, Rüschendorf F, Heinzmann A, Moffatt MF, Cookson WO, Inacio F, Nürnberg P, Stassen HH, Deichmann KA: A genome-wide screen on the genetics of atopy in a multiethnic European population reveals a major atopy locus on chromosome 3q21.3. *Allergy* 2005;60:192–199.
- 31 Chen WM, Broman KW, Liang KY: Power and robustness of linkage tests for quantitative traits in general pedigrees. *Genet Epidemiol* 2005;28:11–23.
- 32 Morris N, Elston R: A note on comparing the power of test statistics at low significance levels. *Am Stat* 2011;65:164–166.