

Statistical aspects of QTL detection in small samples based on a data set from selective DNA pooling in blue fox

J. SZYDA¹, M. ZATOŃ-DOBROWOLSKA¹, H. WIERZBICKI¹ AND A. RZAŚA²

¹Department of Animal Genetics, Koźuchowska 7, 51-631 Wrocław, ²Department of Veterinary Prevention and Immunology, Grunwaldzki 7, 50-366 Wrocław; Agricultural University of Wrocław, Poland

Summary

We analysed data from a selective DNA pooling experiment consisting of 130 unrelated individuals of blue fox (*Alopex lagopus*), which originated from two different types regarding body size. Association between allele frequency and body size was tested using uni- and multivariate logistic regression approach applying Odds Ratio and test statistics from power divergence family. Unfortunately, due to a small sample size and resulting sparseness of the data table, in hypotheses testing we could not rely on the asymptotic distributions of the tests. Instead, we tried to account for data sparseness by (i) modifying confidence intervals of Odds Ratio, (ii) by using a normal approximation of the asymptotic distribution of the power divergence tests with different approaches for calculating moments of the statistics and (iii) by assessing P-values empirically, based on bootstrap samples. As a result, significant association was observed for markers C03.629 and C05.771 representing dog chromosomes 3 and 5 respectively (map location in the blue fox genome is unknown). Furthermore, using simulations we show that among statistics from the power divergence family – Pearson's goodness-of-fit test has the best asymptotic properties for sparse data, while its normalised transformations are extremely conservative.

Introduction

Although fur animals are not the most important species bred on farms, in some countries (Denmark, Finland, Norway, Poland, Russia) fur production is of economic importance. The strong competition on the international fur market forces breeders to speed up the genetic progress in fur animal populations. In Poland one of the most important fur animals is the arctic fox (*Alopex lagopus*). The arctic fox belongs to the class *Mammalia*, order *Carnivora*, family *Canide*, and genus *Alopex*. Farmed arctic foxes descend from the Alaskan blue fox and the Greenland blue fox. The selection of foxes, both arctic (*Alopex lagopus*) and silver (*Vulpes vulpes*), has always been oriented towards improvement of conformation and coat traits as well as reproductive performance. The most important fur coat traits affecting the pelt price are: pelt size, fur coat quality and colour type (FILISTOWICZ et al. 1999).

Up to now selection of foxes is mainly done on farms, based on a simplified selection index and work is underway to set up a national routine genetic evaluation based on BLUB. A natural further development in selection methods would be the incorporation of the molecular information through marker assisted selection, but before this can be done information on genes responsible for a significant proportion of genetic variation of “fur traits” is required. Unfortunately, such an analysis is hampered, by poor information on the blue fox genome, including the polymorphism and localisation of both, markers and functional genes and on linkage groups (KLUKOWSKA et al. 2002; ROGALSKA-NIŻNIK et al. 2003; SZCZERBAL et al. 2003a, 2003b; ŚWITOŃSKI et al. 2003). In the current paper we (i) assess polymorphism of 20 microsatellite markers (known to be polymorphic in dogs) based on individual genotyping, (ii) test for association between selected markers and body size based on samples from selective DNA pooling, (iii) investigate asymptotic properties of statistics representing tests from the power divergence family, which we applied to test for association.

Material

For the analysis 130 DNA samples from blue fox were available, including individuals from the Norwegian type and the Finnish type. Those types markedly differ in body size, so that the Norwegian type represents larger, while the Finnish type – smaller animals.

A fraction of the whole sample was subjected to individual genotyping in order to assess polymorphism of the following microsatellite markers known to be amplifiable in dogs: REN112I02 (localisation in canine genome - CFA01), C02.342 (CFA02), C03.629 (CFA03), FH2732 (CFA04), C05.771 (CFA05), FH2734 (CFA06), C08.410 (CFA08), G06401 (CFA09),

REN153O12 (CFA12), REN227M12 (CFA13), FH2763 (CFA14), REN275L19 (CFA16), FH3047 (CFA17), REN100J13 (CFA20), REN128E21 (CFA22), LEI002 (CFA27), REN248F14 (CFA30), REN43H24 (CFA31), REN106I07 (CFA36), and REN67C18 (CFA37). The selection of markers was based on information from the dog genome – heterozygosity (possibly high) and localisation (each marker on different canine chromosome).

Based on the heterozygosity observed in individual genotyping as well as on amplification properties (similar primer annealing temperature) markers: C03.629, C05.771, C08.410, REN227M12, REN275L19, and LEI002 were chosen for the further analysis with the DNA pooling. Subsequently, all 130 samples were subjected to DNA pooling (Fig. 1). The experiment was designed in the way that four pools were formed: two from DNA of the Finnish type foxes (41 and 36 individuals) and two from the Norwegian type (34 and 29 individuals). During the analysis it occurred that about 20% of samples did not have DNA of quality sufficient for pooling. Since the available sample was rather small, individuals from the same type were assigned to each of two pools by random sampling with replacement – meaning that DNA from some individuals appeared in both pools.

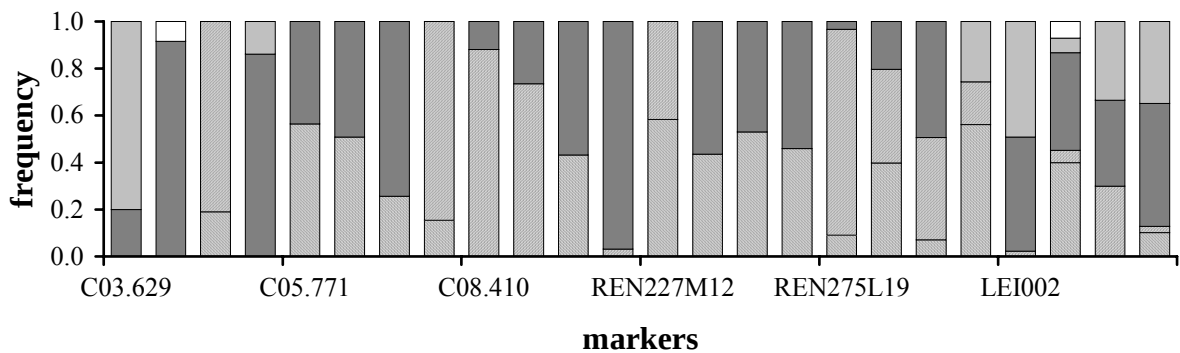


Fig. 1. Allele frequencies observed at each of six markers used in DNA pooling. For a given marker: each of four columns represents one pool - two pools for the Norwegian and two pools for the Finnish type, bars with different backgrounds represent different alleles.

Methods

The statistical analysis of data from DNA pooling comprises the univariate and multivariate logistic regression. While the former approach requires aggregation of data, so that one allele is tested against all the others and allows for the comparisons of two pools at a time, the latter fits a distribution to all alleles observed at a given marker in each of four pools.

Univariate distribution - odds ratio

The hypotheses tested by the univariate model can be expressed by $H_0 : \pi_i = \pi_{i'}$ and $H_1 : \pi_i \neq \pi_{i'}$, where π_i and $\pi_{i'}$ are the allele frequencies of a given marker observed respectively in pools i and i' while $i, i' \in \{1, 2, 3, 4\}$ and $i \neq i'$. A standard test statistic for that case is the natural logarithm of Odds Ratio ($\ln \psi$):

$$\ln \psi = \ln \frac{\frac{n_{iA} + c}{n_{iA'} + c}}{\frac{n_{i'A} + c}{n_{i'A'} + c}} \sim N(0, 1),$$

where, n_{iA} and $n_{i'A}$ represent number of copies of allele A in the i -th and i' -th pool respectively, $n_{iA'}$ and $n_{i'A'}$ are the number of copies of alleles different than A respectively for the i -th and i' -th pool, c represents a constant providing a better agreement of $\ln \psi$ with its asymptotic normal distribution

in case of small allele counts. Here, following HALDANE (1956) and simulation results of AGRESTI (1999), $c=0.5$ was used. Variance of $\ln \psi$ is a function of the observed allele counts:

$$\sigma^2(\ln \psi) = \frac{1}{n_{1A} + c} + \frac{1}{n_{1A^-} + c} + \frac{1}{n_{2A} + c} + \frac{1}{n_{2A^-} + c}$$

Apart from the nominal significance of $\ln \psi$, results of the univariate analysis are presented in the form of 0.01 confidence intervals (CI) of $\ln \psi$ expressed by the standard formula based on the normal approximation of its asymptotic distribution:

$$(\ln \hat{\psi} - z_{0.005} \sigma_{\ln \hat{\psi}}) < \ln \psi < (\ln \hat{\psi} + z_{0.005} \sigma_{\ln \hat{\psi}}).$$

Multivariate model - power divergence statistics

A multivariate model can be applied to the whole data table without the need of allele count aggregation. In such case a set of possible hypotheses can be tested by fitting the following models:

- $H_0 : \pi_{ijk} = \pi$ vs. $H_1 : \pi_{ijk} \neq \pi$ fitting model 1: $\log it(p_{ijk}) = \log\left(\frac{p_{ijk}}{p_{irk}}\right) = \mu + e_{ijk}$,
- $H_0 : \pi_{ijk} = \pi_j$ vs. $H_1 : \pi_{ijk} \neq \pi_j$ fitting model 2: $\log it(p_{ijk}) = \alpha_j + e_{ijk}$,
- $H_0 : \pi_{ijk} = \pi_{jk}$ vs. $H_1 : \pi_{ijk} \neq \pi_{jk}$ fitting model 3: $\log it(p_{ijk}) = \alpha_j + \beta_k + e_{ijk}$,

where, π and p represent respectively estimated and observed allele frequency; μ represents the overall mean, α - allele effect, β - type effect (i.e. Norwegian or Finnish) and e - the residual effect; subscripts i, j, k stand respectively for pool within each type ($i \in \{1, 2\}$), allele $j \in \{1, \dots, r\}$ and type ($k \in \{1, 2\}$). Note that subscript r represents a reference category - which is here the frequency of the last allele.

Comparing fit of the above models can be simplified to the problem of comparison between the observed allele frequencies and the allele frequencies estimated by each of the above models. For that purpose CRESSIE and READ (1984) describe a family of power divergence test statistics, which compare the ratio of observed and expected cell counts scaled by the parameter λ :

$$X(\lambda) = \sum_{i=1}^2 \sum_{j=1}^{n_a} \sum_{k=1}^2 \frac{2n_{ijk}}{\lambda(\lambda+1)} \left[\left(\frac{n_{ijk}}{\hat{n}_{ijk}} \right)^\lambda - 1 \right] - \frac{2}{\lambda+1} (n_{ijk} - \hat{n}_{ijk}) \sim \chi^2_{4(1-n_a)-n_p} \quad (\text{OSIUS and ROJEK 1989}),$$

where, n and $\hat{n}$ represent observed and expected cell counts respectively, n_a is the number of alleles at marker j and n_p is the number of parameters fitted by the model, while other subscripts are as above. In the analysis of fox data we considered three different values of λ , yielding a well known Pearson goodness-of-fit test [$X(1)$], Freeman-Tukey test [$X(-\frac{1}{2})$] and Cressie-Read test [$X(\frac{2}{3})$]. Furthermore, in order to better account for the sparseness of data tables, we applied standard normal transformation to the Pearson goodness-of-fit test with moments derived by OSIUS and ROJEK (1989):

$$Z = \frac{X(1) - \mu_{X(1)}}{\sigma_{X(1)}} \sim N(0,1),$$

comprising:

- A case of “fast increasing harmonic mean”, assuming $\mu_{X(1)} = 4(n_a - 1)$ and $\sigma_{X(1)}^2 = 8(n_a - 1)$, further denoted as Z1.
- A case of “increasing arithmetic mean”, assuming the same mean as above: $\mu_{X(1)} = 4(n_a - 1)$, but a modified variance $\sigma_{X(1)}^2 = 8(n_a - 1) + S_1$ accounting for small cell counts and probabilities, further denoted as Z2.

- A general case, again assuming the same mean: $\mu_{X(1)} = 4(n_a - 1)$ and a reduced variance $\sigma_{X(1)}^2 = 8(n_a - 1) - S_2$, further denoted as Z3.

The Bayesian Information Criterion: $BIC = -2 \ln L + n_p \ln 4n_a$ was calculated as an additional measure of model fit.

Bootstrap

A parametric bootstrap was used in order to obtain empirical significance for the six test statistics, which were applied to the fox data [i.e. for $X(1)$, $X(-\frac{1}{2})$, $X(\frac{2}{3})$, Z1, Z2, Z3]. Following this approach 1000 data tables were generated under each model meaning that allele counts were generated based on marginal sums of alleles at each pool and allele frequencies estimated by a given model (WINKLER 1996).

Simulations

Apart from analysing the particular data set, we were also interested in evaluating asymptotic properties of tests from power divergence family and their normalised versions. For that purpose we simulated data tables assuming various constellations underlying equiprobable as well as skewed null hypotheses. The hypotheses were defined in terms of allele frequencies, so that:

- For the equiprobable null hypothesis the allele frequencies followed a uniform distribution, across possible alleles.
- For the skewed hypothesis frequency of one allele was set to 0.5 while remaining allele frequencies were assigned a uniform distribution.
- Additionally, for skewed hypotheses an erroneous allele assignment was simulated (e.g. allele 140 was wrongly assigned as 144) with two error rates of 5% and 15%, in order to imitate the problem of stuttering bands, which occurs in DNA pooling data.

Apart from allele frequencies, other parameters used for data generation were: sample size (40, 70, 200 individuals) and allele numbers (3, 4, 5 alleles). For each set of parameters 1000 data sets were generated. Note that with increasing allele numbers and skewness in allele frequencies, as well as with decreasing sample size the sparseness of the data table increases. Furthermore, for model 1 skewed hypothesis is no longer the null hypothesis.

Results

Testing trait-marker association in real data

As already mentioned above, in the first step of the analysis, for two given pools the comparison was made between the ratio of frequency of one allele against an aggregated frequency of all the other alleles using $\ln \psi$, in order to go around the sparseness of the original (not aggregated data table). Theoretically, one would expect that when two pools belonging to the same type are compared result would be nonsignificant, whereas for comparing pools from two different types significant outcome would occur when an association between body size and a marker allele frequency exists. Considering nominal significance levels depicted in figure 2 it can be seen that cases exist where a small type I error corresponds to comparison of the same types. Nevertheless, for alleles of marker C03.629 differences in allele frequencies between the Finnish and the Norwegian type are significant amounting to $\alpha=0.00168$ and $\alpha=0.00004$, while corresponding comparison between two pools of the Finnish type with $\alpha=0.34643$ remain nonsignificant. Consistent results can also be observed for marker REN227M12, where none of comparisons yields a significant result – indicating no association between this marker and a trait studied.

When analysing estimated CI of $\ln \psi$ given in figures 3 and 4, a nonsignificant (at 0.01 level) outcome is marked by a CI containing zero. Note, that results for C03.629 and REN227M12 are also confirmed on those figures. Additionally, each of tested alleles of C05.771 CI indicates significant differences between different fox types and nonsignificant differences for the same types. However, in case of other markers results are inconsistent with either significance pattern varying among alleles tested or showing significant differences between pools of the same types.

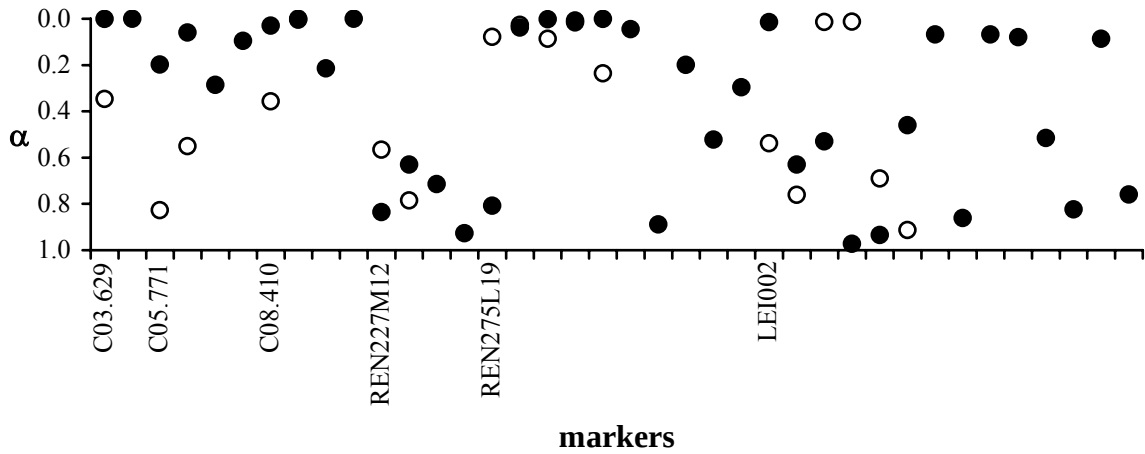


Fig. 2. Nominal type I error rate (α) for $\ln \psi$ comparing allele frequencies belonging to different (●) and the same (○) types of blue fox.

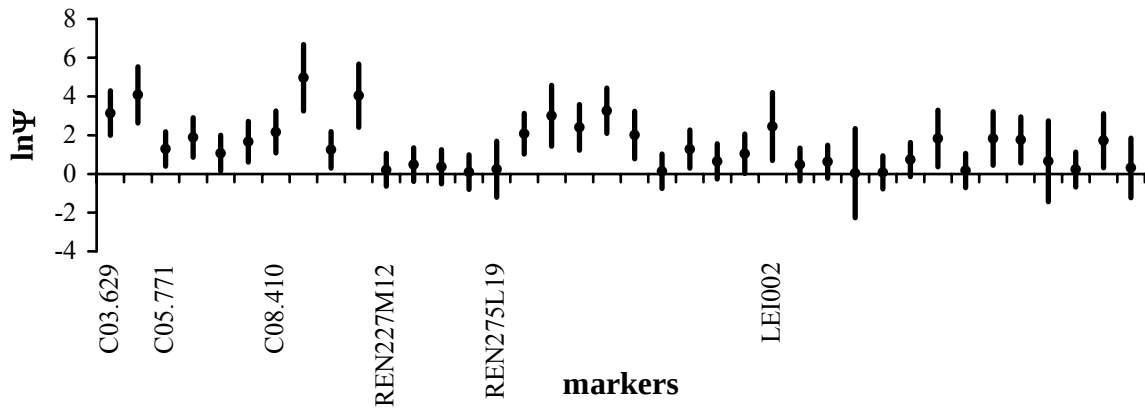


Fig. 3. CI for $\ln \psi$ estimated assuming $\alpha=0.01$ for comparison of different types of blue fox.

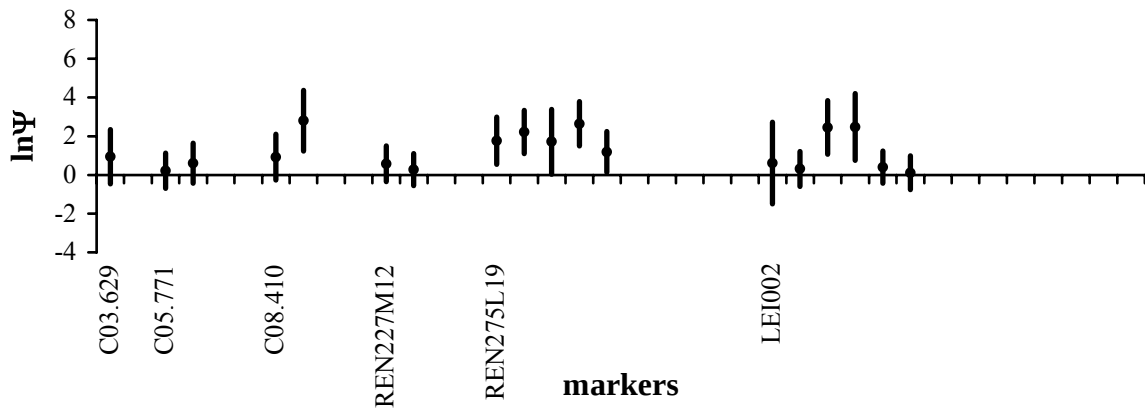


Fig. 4. CI for $\ln \psi$ estimated assuming $\alpha=0.01$ for comparison of the same types of blue fox.

The next step comprised fitting models 1-3 to the original (not aggregated) data, what on one hand provides means for testing more precise hypotheses, but on the other hand sparseness of the data causes that the reliability of inferences suffers from poor agreement with asymptotic conditions. Table 1 summarizes information on the quality of model fit as well as asymptotic (based on asymptotic test distribution) and empirical (based on parametric bootstrap) P-values. Considering values of BIC and asymptotic P-values it is evident that none of the applied models fit the data. Even if separate allele frequencies are assigned to each type, a common parameter does not suffice to describe variation in allele frequencies within pools of the same type. However when P-values based on bootstrap are considered one can conclude that for markers REN227M12 and

REN275L19 model fitting the same allele frequency to all the pools is sufficient (indicating no association) while for the remaining four markers *type-specific* allele frequencies are required.

Table 1. Information on fitting multivariate logistic models to blue fox data. Models 1-3 are described in text, „Obs.” denotes the full model fitted to the observed cell counts. LnL is the natural logarithm of model likelihood.

Model	lnL	BIC	Power divergence statistics											
			Asymptotic P-values						<i>Empirical P-values</i>					
			$X(1)$		$X(-\frac{1}{2})$		$X(\frac{2}{3})$		Z1		Z2		Z3	
C03.629														
1	-443.83	890	534.7		567.0		498.3		91.7		90.4		92.2	
			0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
2	-403.20	817	402.5		490.7		388.6		68.3		66.8		68.6	
			0.000	0.006	0.000	0.000	0.000	0.000	0.000	0.006	0.000	0.006	0.000	0.006
3	-296.06	614	160.8		269.4		167.7		25.6		23.5		25.6	
			0.000	1.000	0.000	1.000	0.000	0.814	0.000	1.000	0.000	1.000	0.000	1.000
Obs.	-192.77	430												
C05.771														
1	-320.11	642	209.5		311.6		211.4		50.4		50.7		51.5	
			0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
2	-315.15	634	232.6		279.3		222.2		56.1		56.5		56.5	
			0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
3	-274.08	556	118.3		209.0		124.0		27.6		26.2		26.8	
			0.000	0.984	0.000	0.935	0.000	0.000	0.000	0.984	0.000	0.988	0.000	0.987
Obs.	-195.27	407												
C08.410														
1	-309.12	620	289.0		339.1		279.6		70.2		70.6		71.2	
			0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
2	-292.13	588	244.2		303.0		240.9		59.1		59.1		59.5	
			0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
3	-186.47	381	42.1		62.8		43.7		8.5		7.9		8.4	
			0.204	0.543	0.057	1.000	0.187	1.000	0.014	1.000	0.026	1.000	0.016	1.000
Obs.	-159.94	337												
REN227M12														
1	-335.95	674	136.6		276.2		149.1		32.1		32.4		33.0	
			0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
2	-296.45	597	109.0		145.9		106.7		25.3		25.2		25.4	
			0.000	0.541	0.000	0.474	0.000	0.016	0.000	0.541	0.000	0.543	0.000	0.541
3	-277.19	563	58.4		93.6		60.9		12.6		12.0		12.4	
			0.076	1.000	0.000	1.000	0.063	1.000	0.000	1.000	0.000	1.000	0.000	1.000
Obs.	-256.52	530												
REN275L19														
1	-389.02	781	193.3		224.9		189.6		37.0		37.0		39.3	
			0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
2	-372.17	752	168.4		179.7		164.0		31.9		31.9		32.1	
			0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
3	-352.60	720	112.6		155.9		116.0		20.5		19.8		21.9	
			0.000	0.328	0.000	0.184	0.000	0.002	0.000	0.328	0.000	0.346	0.000	0.300
Obs.	-284.07	598												
LEI002														
1	-470.62	944	157.5		228.7		163.2		25.0		25.0		25.2	
			0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
2	-428.65	868	113.0		117.6		108.7		17.2		17.0		17.3	
			0.000	0.458	0.000	0.438	0.000	0.223	0.000	0.458	0.000	0.461	0.000	0.458
3	-417.75	858	78.0		117.6		79.7		11.0		10.8		11.0	
			0.047	0.992	0.000	0.986	0.043	0.784	0.001	0.992	0.001	0.993	0.001	0.992
Obs.	-371.29	787												

Empirical distribution of power divergence statistics

A comparison of the empirical distributions of tests from the power divergence family is shown on figure 5. Distribution of the X statistics varies with the degrees of sparseness, expressed by the number of alleles and sample size. Although somewhat conservative for sparse data, X statistics remain in good agreement with their asymptotic distributions for denser data. The robustness of the Z statistics towards sparseness of data is manifested by the fact, that their distribution does not depend much on the sample size and number of alleles, but unfortunately, their type I error rates are very low for different testing conditions – making those tests extremely conservative.

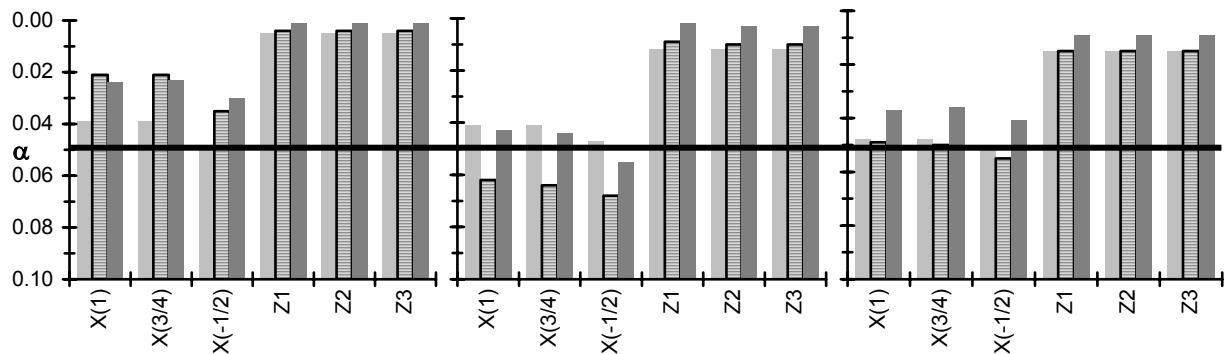


Fig. 5. P-values realised for a 0.05 critical value corresponding to the asymptotic distribution. Three bars for each test statistics represent respectively (from left to right) simulated designs with 3, 4 and 5 marker alleles. Three graphs represent respectively (from left to right) simulated designs with 40, 70 and 200 individuals.

As shown on figure 6 the reason for their conservativeness is that the expression for $\mu_{X(1)}$ as proposed by OSIUS and ROJEK (1989) seems to overestimate the true mean of $X(1)$ regardless of testing conditions. On the other hand the standard deviations are always underestimated. As can be further recognised from table 1 and figures 5-6 not much difference is observed between three different types of the normalised tests. Summarising the simulation results we observe that all statistics tend to be conservative with the exception of X tests for sparse samples (figures 6-7).

Conclusions

The investigations of the asymptotic properties of test statistics our simulations showed that the Pearson's goodness-of-fit test has reasonable good asymptotic properties and unless a modified formulation of mean is derived, the normalised statistics are not recommended. Our findings remain in good agreement with simulation results presented by GARCIA-PÉREZ and NÚÑEZ-ANTÓN (2001).

Genetically, some evidence of the association between marker allele frequency and body size observed for fox genome regions corresponding to canine chromosomes 3 and 5 is observed. Since no linkage and association studies have been carried in blue for - out findings indicate that those two chromosomes could be the first choice for further analysis in order to verify our preliminary result.

Using power divergence tests for extremely sparse data causes problems with their asymptotic distribution and, if possible, exact or empirical type I error rates are recommended. The best asymptotic performance was observed for the classical test – Pearson's goodness-of-fit test, while its normalised versions failed to keep the assumed significance level.

References

- AGRESTI, A., 1999: On logit confidence intervals for the odds ratio with small samples. *Biometrics* 55: 597-602.
- CRESSIE, N. A. C.; READ, T. R. C., 1984: Multinomial goodness-of-fit tests. *Journal of the Royal Statistics Society Ser. B*, 46: 440-464.
- CRESSIE, N. A. C.; READ, T. R. C., 1988: *Goodness-of-Fit Statistics for Discrete Multivariate Data*, Springer, New York.
- FILISTOWICZ, A.; ŻUK, B.; ŚLAWOŃ, J., 1999: Evaluation of factors determining prices of Polish arctic fox skins at the Helsinki International Auction. *Animal Science Papers and Reports* 17: 209-219.
- GARCIA-PÉREZ, M. A.; NÚÑEZ-ANTÓN, V., 2001: Small-sample comparisons for power divergence goodness-of-fit statistics for symmetric and skewed simple null hypotheses. *Journal of Applied Statistics* 28: 855-874.

- HALDANE, J. B. S., 1956: The estimation and significance of the logarithm of a ratio of frequencies. *Annals of Human Genetics* 20: 309-311.
- KLUKOWSKA, J.; SZYDŁOWSKI, M.; ŚWITOŃSKI, M., 2002: Linkage of the canine-derived microsatellites in the red fox (*Vulpes vulpes*) and arctic fox (*Alopex lagopus*) genomes. *Hereditas* 137: 234-236.
- OSIUS, G.; ROJEK, D., 1989: Normal goodness-of-fit tests for parametric multinomial models with large degrees of freedom. *Fahbereich Mathematik/Informatik, Universitaet Bremen. Mathematik Arbeitspapiere* 36.
- ROGALSKA-NIŻNIK, N.; SZCZERBAL, I.; DOLF, G.; SCHLÄPFER, J.; SCHELLING, C.; ŚWITOŃSKI, M., 2003: Canine-derived cosmid probes containing microsatellites can be used in physical mapping of arctic fox (*Alopex lagopus*) and chinese racoon dog (*Nyctereutes procyonoides procyonoides*) genomes. *Journal of Heredity* 94: 89-93.
- SZCZERBAL, I.; ROGALSKA-NIŻNIK, N.; KLUKOWSKA, J.; SCHELLING, C.; DOLF, G.; ŚWITOŃSKI, M., 2003: Comparative chromosomal localization of the canine-derived BAC clones containing LEP and IGF1 genes in four species of the family Canidae. *Cytogenetic and Genome Research* 102: 264-266.
- SZCZERBAL, I.; ROGALSKA-NIŻNIK, N.; SCHELLING, C.; SCHLÄPFER, J.; DOLF, G.; ŚWITOŃSKI, M., 2003a: Development of a cytogenetic map for the Chinese racoon dog (*Nyctereutes procyonoides procyonoides*) and the arctic fox (*Alopex lagopus*) genomes, using canine-derived microsatellite probes. *Cytogenetic and Genome Research* 102: 267-271.
- ŚWITOŃSKI, M.; ROGALSKA-NIŻNIK, N.; SZCZERBAL, I.; BÄR, M. 2003: Chromosome polymorphism and karyotype evolution of four canids: the dog, red fox, arctic fox and racoon dog. *Caryologia* 56: 375-385.
- WINKLER, B., 1996: Bootstrapping goodness of fit statistics in loglinear Poisson models. *Discussion Papers of the Institute of Statistics at the University of Munich*.

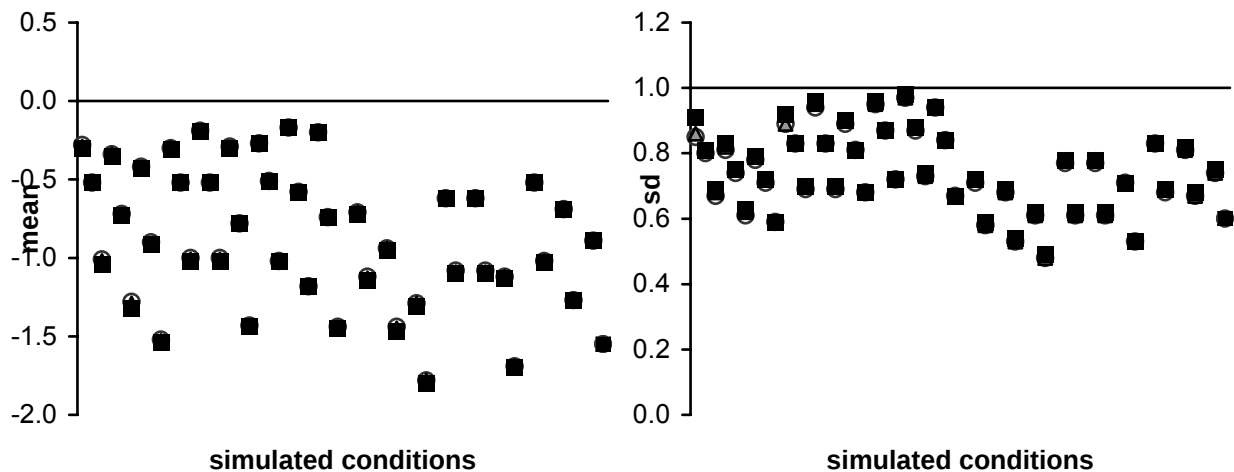


Fig. 6. Realised distribution of Z statistics across various sets of simulated conditions. The left and the right diagrams show respectively means and standard deviations (sd) of test statistic calculated from 1000 realisations. Z1 is represented by circles, Z2 by triangles, and Z3 by squares.

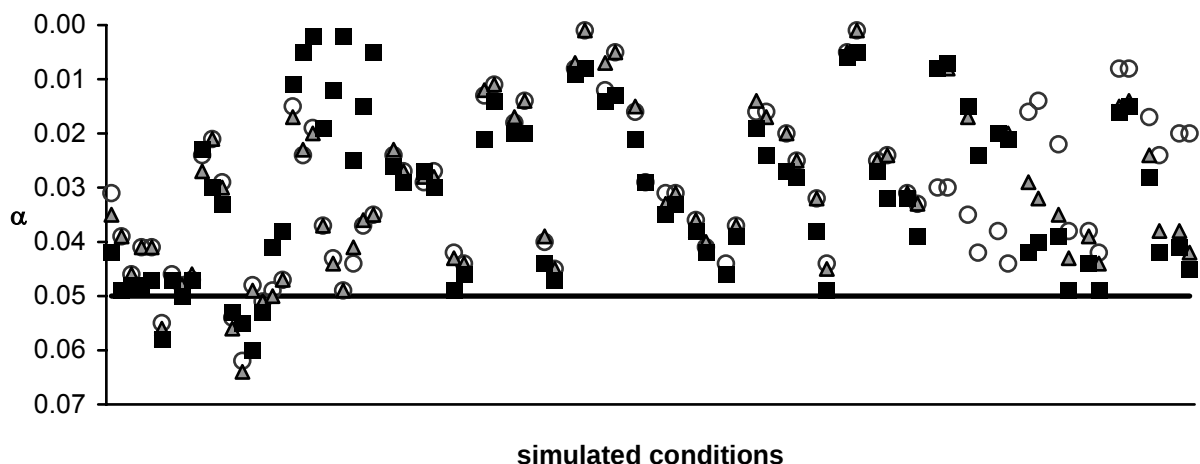


Fig. 7. Realised P-values of X statistics realised for a 0.05 asymptotic distribution cut-off, across various sets of simulated conditions. $X(1)$ is represented by circles, $X(-\frac{1}{2})$ by squares, and $X(\frac{2}{3})$ by triangles.