

Overview for Windows version of DMLE+2.2
Release 2.2, January 6, 2005.

Input file format

Windows details

Screen Output

File Output

INPUT FILE:

In general, the input file consists of alternating lines of descriptions then data. Use the sample files provided at the program's Web site as templates. Data may over-run a line as long as no new line character is used (i.e. don't use the "Enter" key to make an extra new line). All fields with multiple entries may be delimited with space(s) or tabs. Avoid having any blank lines in the input file. In the following description, each line from the input file (a combination of input descriptors and data) is given, immediately followed by a description of the meaning of the input contained within square brackets [].

The program's output files are named <filename.ext> (where ".ext" will be .dat, .sig, .hap, .log, .mat, .tre, or .hpf; see description in the Output section).

Step-by-step Guide to Input File Entries:

Data as genotypes?:(0=haplotypes, 1=genotypes):

0

[If the patient data are completely haplotyped, use 0.]

Genetic model: (Dominant=0,Recessive=1)

0

[Only used for genotype data. If haplotypes are used, this field is ignored, and it is assumed that all chromosomes included below in "Patient haplo- or genotypes" are disease bearing, whatever the genetic basis for the disease.]

Read old file?: (0=no, 1=yes):

0

[If the tree is to start from a saved tree at the end of a previous run, set this to 1, and change the name of the output matrix file that you want to use (see Output section) from "filename.mat" to "inmat.txt".]

Use fixed random seed?:(0 = no (=random), negative integer = yes (=fixed)):

0

[The random number generator either starts with a random negative integer (if set to 0), or uses the negative integer supplied here if you want to be able to reproduce the same sequence of pseudo-random numbers.]

chromosomes (N):

148

[Number of chromosomes in the disease sample. If the data are genotypes, this must be an even number.]

loci per chromosome (L):

5

[Number of marker loci.]

Numbers of haplotypes in the normal (base) population:

45 1 1 1 1 1

35 1 1 1 1 6

15 1 1 2 1 2

16 1 1 2 2 2

10 1 2 2 3 3

5 2 2 2 3 4

2 2 -1 3 4 5

1 -1 3 -1 4 2

[Each row starts with a frequency (count) of the haplotype, followed by that haplotype's alleles. Unknown markers types are indicated with a -1. For any particular marker, when the control sample and patient sample are considered together (pooled), the allele codes for that marker should use the numbers 1 through n, where n is the total number of different alleles present at that marker. Gaps in the numbering system (e.g. 1,2,3,7, and 8 as the collection of alleles at a marker in the controls + disease-bearing chromosomes) are not allowed and will produce an error message.

The program uses these normal population counts only to check against normal allele counts - genotypes are exactly equivalent to haplotypes. Therefore, if you have genotype data, split each genotype into any of the possible haplotypes. The frequency column must be included even if it is always "1". Haplotypes do not have to be completely grouped, e.g. you can have

23 1 1 2 1 1

14 1 1 2 1 1

(Note that earlier version of the program ended the list of haplotypes with a row reading "-99".)

Map distances:

0.0 0.00013 0.000165 0.000195 0.000205

[The first map distance must be set to 0.0, with the rest given in Morgans (e.g. given that 1 cM $\approx$ 1000 kb in humans, the final entry above is $\sim$ 20.5 kb). Ensure that there are as many entries here as there are markers indicated above. These entries are overridden if a contig file is being read for sequence information (see below).]

Run simulation?:

0

[If analyzing actual data, this should be zero. Set this to be 1 if you wish to simulate chromosomes for an analysis (under a population coalescent process) using the allele frequencies specified in the population of normal chromosomes and the other parameters as specified in the input file. If a simulation is run, the actual location of the mutation (for the simulated data) is read from the variable below. Regardless of whether 0 or 1 is chosen here, the initial tree topology, and the ancestral and internal states, are simulated in an identical fashion. Only the tips (extant chromosomes) are not simulated when this parameter is set to zero.]

Mutation location (only used for simulated data):

0.0003

[Position of the mutation relative to the first marker (used for simulations). This number can be negative. The value of this number is ignored if no simulation is run, but as with all fields in the input file, some value must still be present.]

Mutation's low and high boundaries

-0.05 0.05

[Limits to the range of mutation positions (relative to marker 1, in Morgans) that will be considered in the Markov chain]

simultaneous runs:

1

[Number of chains run simultaneously in the MCMC analysis. If more than one chain is run simultaneously, the program will generate a "square root of R" statistic that can be used to assess whether the multiple runs have converged (currently implemented only for convergence amongst iterated mutation locations, not for age of mutation).]

Starting value(s) for mutation location for each simultaneous run (-99 for random):

-99

[The value for the mutation location used in iteration 1 of the Markov chain (not the true location of the mutation). If -99, a random uniform number between the mutation's low and high boundaries (see above) will be generated. If not -99, there must be one number for each simultaneous run, and these must be between the low and high boundaries.]

Population growth rate:

0.085

[Population growth rate, estimated from historical data.]

Proportion of population sampled:

0.001825

[Proportion of disease chromosomes in sample, estimated from disease incidence and number of disease chromosomes sampled.]

Iterate ancestral states, mutation age, mutation location, allele freq. (0=no, 1=yes*):
1 0 1 0

[Choose 0 to hold the parameter fixed, 1 to iterate over the parameter. Internal nodes are always iterated. If the mutation location is known and you want to estimate the mutation's age, use 1 1 0 0 (or 1 1 0 1). If the allele frequencies are iterated, the frequencies are initially set to be equal for all alleles present in the normal and disease sample at each marker. If they are not iterated, frequencies from the normal population are used.

* In version 2.2, there are 3 methods of nominating new mutation locations. The default value "1" (omit the quotes!) uses the standard method. "2" uses slice sampling, which can be tried if the likelihood surface is multimodal and the peaks are relatively far apart. It is slower than the standard method. "3" samples uniformly across the range of the mutation boundaries, and is NOT recommended unless you are having extreme difficulty getting multiple chains to converge to the same locations.]

Flip (potentially) all loci? (0=no, 1=yes):
1

[When the alleles of the internal and ancestral nodes are changed in the MCMC, a single random marker from each node can be (potentially) altered (= 0), or all markers can be (potentially) altered (=1).]

Adjustment level for tree, rootage, mutation_location, ancestral states, internal states, allele freqs., and method:

1.0 3.0 0.005 0.15 0.15 0.1 1.0

[The size of the adjustment factor that determines the range of new values that are nominated at each iteration of the MCMC, for the six variables (tree topology; root age; mutation location; ancestral haplotype alleles; internal node haplotype alleles; and allele frequencies respectively), if they are set to be iterated - see above. The method variable determines how transition probabilities are calculated. 0 uses marginal allele frequencies, 1 uses marginal haplotype frequencies. It is recommended that the default value of 1.0 be used, as simulations over a wide range of input datasets indicate that it will be more accurate roughly 75% of the time. If the control population markers are thought to be in near-perfect linkage equilibrium, allele frequencies can be used, and the program will run faster.]

Burn-in iterations:
1000000

[Number of iterations to perform in the MCMC analysis before the data for the posterior distribution histograms starts to be tabulated.]

Iterations:

1000000

[Number of iterations over which data is tabulated in the MCMC analysis to construct the posterior distribution histograms.]

Screen update and file update intervals, save acceptance proportions, save population haplotype frequencies:

100 100 1 0

[Data is sent to the screen and the output files at intervals corresponding to the first two entries respectively (data from the first and last iterations are also shown). A high rate of screen updates may slow down program execution. Large numbers of file updates produce large data files, and may be difficult to read completely in spreadsheets that can store only limited number of rows, e.g. ~65,000 for Microsoft Excel.

Details about the acceptance rate for changes in the five variables (tree topology, mutation location, ancestral states, internal states, and allele frequencies) will be printed to the output file *.dat (see Output section) if the next parameter is 1. These acceptance proportions are printed to the screen even if this value is 0. The last number determines whether haplotype frequencies over the non burn-in iterations will be saved to the file *.hpf.

Number of histogram bars:

200

[The number of bins into which the mutation locations and/or mutation ages are sorted]

ALPHA level for recdist histogram data:

0.05

[Significance level for the credible set recorded in the *.sig and *.tre output files (see Output section)]

Mutation age:

100

[Estimated age (in generations) of the original mutation in the population. If age is being iterated, this number will be ignored and the initial age will be a random number no more than 100 generations older than the minimum set below.]

Mutation age boundaries:

0 1000

[Sets boundaries for mutation age if it is being iterated.]

Star genealogy (0=no, 1=yes):

0

[If set to 1, the original tree will be given a star phylogeny (all branch lengths equal). If this is combined with an adjustment level of 0.0 for the tree topology (see above), the star topology will remain throughout the run.]

Loci for the ancestral state (-99 for random):

1 1 2 1 1

[The contents of this field are only important if the “iterate ancestral states” variable (see above) is set to 0.]

Patient haplo- or genotypes:

79 1 1 1 1 1

31 1 1 2 1 1

18 2 1 1 1 1

5 1 2 2 1 2

5 1 1 1 1 2

4 2 1 2 2 4

2 2 1 2 2 1

1 1 1 1 2 -1

1 -1 1 1 1 1

1 1 2 1 2 1

1 1 2 2 3 2

[These entries can be haplotypes or genotypes. Missing alleles (-1) are allowed in both cases.

1) Haplotypes (shown above). The first column is the frequency of the haplotype, followed by the allele at each marker. Columns are delimited by spaces or tabs. If simulated data is being used, you may fill in the above with arbitrary haplotypes or genotypes with the constraint that they must add up to the total number of chromosomes indicated in the # chromosomes field above (in this case the haplotypes you provide are not used, but their frequencies are needed to know when to stop reading this field and move on to the next one). E.g.

148 1 1 1 1 1

would be fine. The total number of rows in this table of haplotypes does not matter. As with the normal population haplotypes, all identical disease haplotypes do not have to be grouped.

2) Genotypes. There is no first column for frequencies. Alleles may be separated with “/”, “\”, or “;”. E.g.

1/2 1/1 1/3 2/4 1/1

3/2 4/3 4/2 5/1 2/1

etc.

Note that for genotype data, if the number of chromosomes is 148, there should only be 74 rows of genotype data, whether dominant or recessive.]

Use sequence weights?

0

[When this is set to 1, priors will be assigned to the DNA sequence depending on whether the region is an exon, intron, or non-gene (see below). If set to 0, a uniform prior is used.]

Weights for exons, introns, non-genes

1 0.17 0.02

[These weights are only used if the above sequence weight parameter is set to 1. The first entry is for exons, the second for introns, and the last for non-genic regions. These are followed immediately by either:]

A)

-0.00208767 e
-0.00208449 i
-0.00206302 e
etc.

[The first column indicates the starting position, relative to the first marker (0.0), of the initial element in the DNA sequence in the area of interest. The second column indicates what type of element starts at this location (e = exon, i = intron, n = non-genic).]

or B)

Input file:

filename

301500

301700

303234

311111

312456

[Where filename is the filename for a Genbank/NCBI contig file, for instance NT_029406, saved in text format (see sample file). This file must be placed in the same directory as DMLE+. Following the filename are the positions of your markers in the contig (which starts at base #1).](Note that earlier versions of the program had a line that read "99999 Input file:")

WINDOWS DETAILS:

After double-clicking DMLE+.exe, the program's main window appears. To run the program, an input file must be selected from the File Open menu. After a file is opened, click the right-pointing arrow at the far left of the toolbar to start running the Markov chain(s). On the main window, a printout of the program's interpretation of the input file is listed. This list should be checked to insure that the program is reading the input as you intended. If the program crashes without printing this information to the screen, check the <filename.log> file to look at the input.

The DATA1234 button on the toolbar brings up the data window.

The mutation location graph button (horizontal arrow over a DNA double helix) shows a graph that tracks the location of the mutation in each chain, updated according to the screenupdate parameter (see Input file format). This graph can be zoomed-in on, by holding down the left mouse button while dragging a rectangle within the graph screen. The axes of the graph can be shifted by holding down the right mouse button and dragging. To zoom back out to the original full-sized graph, hold down the left mouse button and drag a rectangle that extends past the left side of the graph surface. A zoomed graph will not scroll to keep up with iterations that run off the right side of the screen - only the full sized graph will continue to scroll.

The mutation age button (vertical arrow next to a tree) shows a graph that tracks the location of the mutation in each chain, updated according to the screenupdate parameter (see Input file format). Zooming-in works as for the mutation location button.

The pause button (arrow head and two horizontal lines separated by a single vertical line) pauses and restarts (through toggling) program execution. When the program is paused, a message appears on the toolbar.

As the program executes, progress indicators on the menu bar show how far the program is into either the burn-in or "real" iterations. When the program is done, "Finished" will appear on the toolbar.

The panel below the toolbar gives updates (at a rate determined to the screenupdate parameter - see Input file format) for a number of variables. These are:

- 1) iteration number (Iter.)
- 2) log-likelihood for the entire tree (LL)
- 3) the "square root of R" statistic (sqR) if more than one chain is being run.

- 4) acceptance proportion for the tree topology (AT)
- 5) acceptance proportion for the mutation location (AM)
- 6) acceptance proportion for changes in the ancestral node (AA)
- 7) acceptance proportion for changes in the internal nodes (AI)

(Note that the ancestral node is the haplotype on which the mutation arose, not the root node of the tree.)

- 8) acceptance proportion for allele frequency changes (AALL).

After the burn-in iterations have finished, some buttons will appear to the right of the progress bars. These produce histograms of the iterations up to that point in time, excluding the burn-in iterations. “M” stands for mutation location, “T” for time of original mutation. Buttons with arrows cycle through chains, those without arrows update the current chain. How many buttons you see depends on what is being iterated, and whether there is more than one chain.

When the program is finished, these buttons disappear and two extra buttons appear to the right of the DATA1234 button. Clicking the histogram button cycles through histograms of mutation location for each chain. Clicking the tree button cycles through histograms of age of mutation for each chain. Each chain includes data for all the non-burn-in iterations, regardless of the updatescreen or update parameters set in the input file. The histogram bars in blue (or green for the mutation age) fall within the credible set determined by the ALPHA parameter (see Input file format), while those in red do not.

In the File menu, Printer Setup allows the choice of printer. File Print prints whichever graph screen is currently visible. These are low resolution graphs. For better quality graphs, use the data saved to the *.hist output files (see Output). If the visible graph is zoomed in on, only the zoomed area will be printed. Graphs can also be saved to .pdf's (through Printer Setup) if the appropriate software is installed on your computer, but again, the quality will be poor. Terminate and Save terminates the program but first saves whatever data has been produced so far, and also produces histograms up to the current generation (again, only for iterations after the burn-in).

When finished, the program can be shut down either with the “x” button in the top right corner of the screen, or through File Exit

Screen output:

Each time the screen is updated (at an interval determined by the screenupdate parameter - see Input file format section), either 7 or 8 numbers are displayed.

1: the iteration number.

2: the log-likelihood for the entire tree (topology, mutation location, internal and ancestral states of marker alleles).
3: This column is only present if the number of simultaneous runs is set to greater than 1. It is the value for the “square root of R” statistic.
4: the proportion of iterations in which the change in tree topology has been accepted.
5: the proportion of iterations in which the change in mutation location has been accepted.

6: the proportion of iterations in which a change in the original mutation’s haplotype has been accepted.

7: the proportion of possible internal node haplotype changes in which one (or more) changes in an internal node’s haplotype has been accepted (i.e. total number of changed haplotypes/(iterations x internal nodes)).

If more than 1 chain is being run, the first 7 of the above numbers refer to the statistics from chain #1. Values for the other chains are not printed to screen, but are printed to file output (see below).

8: the proportion of possible allele frequency changes that have been accepted.

When the program is done, it will print “Finished”. In a large sample, there may be a small delay between the screen output for the final iteration and the Finished signal, due to writing potentially large files to disk.

File output:

Six output files are created on each run of the program. Each has the same filename (+ the extension if any) of the input file, followed by an extension of either .dat, .sig, .hap, .log, .tre., or .mat. (For example, an input file called test1 will produce test1.dat, test1.sig, test1.hap, test1.log, test1.tre, and test1.mat.) The contents of these files are as described below. The columns of these files are space delimited, and are therefore easily parsed in a spreadsheet. If the output files from a previous run are present in the directory from which the program is run, and the new input file has the same name as one previously used, the original contents of the old output files will be overwritten at the start of the new run.

In addition to these six files, if a contig file is used for sequence information, an additional file (exons.txt) is created, which shows the exon/intron/non-genic boundaries, in Morgans, translated from the contig file. If the “save haplotypes” (see above) option is set to 1, an additional file *.hpf will be created, which stores the haplotype frequencies of the disease bearing chromosomes over the non burn-in iterations.

*.dat - This file stores information from both the burn-in and “good” iterations. The frequency at which the data is saved to this file is set by the fileupdate parameter. The columns correspond to the 7 or 8 numbers described for the screen output. If there are multiple chains, these seven columns will be repeated for each run. In addition, the value

of the “square root of R” statistic will be recorded in an extra column at the end. If the “acceptance details” parameter is set to 0, columns 4 through 8 will not be printed to the *.dat file.

- 1: the iteration number.
- 2: the mutation location.
- 3: the log-likelihood for the entire tree (topology, mutation location, internal and ancestral states of marker alleles).
- 4: the current age of the mutation.
- 5: the log-likelihood for the tree structure itself, excluding priors.
- 6: the proportion of iterations in which the change in tree topology has been accepted.
- 7: the proportion of iterations in which the change in mutation location has been accepted.
- 8: the proportion of iterations in which a change in the original mutation’s haplotype has been accepted.
- 9: the proportion of possible internal node haplotype changes in which one (or more) changes in an internal node’s haplotype has been accepted (i.e. total number of changed haplotypes/(iterations x internal nodes)).
- 10: the proportion of possible allele frequency changes that have been accepted.

These ten columns are repeated for each chain. If more than 1 chain is run, a final column gives the value for the “square root of R” statistic.

*.sig - Histogram data for the posterior distribution of the mutation locations over only the “good” iterations, i.e. with the burn-in results excluded. The data for this histogram is recorded every iteration, and is not affected by either of the “update” parameters. The file is sorted by column 1. If more than one chain is run, the histogram data for each chain are printed in the same file, separated by a blank line. The columns in the output are as follows:

Column 1: the bin midpoints for the mutation locations over the “good” iterations. The number of bins is from an input parameter (see Input file format). The total range between the lowest and highest mutation locations from these iterations is divided by the number of intervals to determine the bin width.

Column 2: the frequency for each bin.

Column 3: the actual counts for each bin.

Column 4: significance (1=sig, 0 = non-sig) based on the ALPHA value (see Input file format). The sum of the frequencies of the non-significant bins is $\leq$ ALPHA.

If the mutation location is not being iterated, this file will be all zeroes.

*.hap - The frequencies (in actual numbers) of all haplotypes present in the sample. If the data are not being simulated, these haplotype frequencies are those from the input file. If the data were simulated, these values will (generally) be different from the input

haplotype frequencies (which are ignored). If multiple chains are run simultaneously, only the haplotypes from chain #1 run are used. (With simulated data, the haplotypes would generally be different for all chains, whereas the haplotypes are identical if real data is used.)

*.log - This file contains a list of what the program has stored as the input parameters. This file should be checked to make sure the program is correctly reading in the input. It is written and closed at the start of program execution, and can therefore be checked while the program is running. Program crashes are usually a result of incorrect input format. The first line is the random number seed used, which can be useful for debugging.

*.tre - Histogram data for the posterior distribution of the time since the original mutation, over only the “good” iterations. Columns are as in the *.sig output file described above. If the depth of the mutation is not being iterated, this file will be all zeroes.

.mat - This file contains a matrix representation of the tree at the end of the “good” iterations. Such trees may then be loaded into the start of a new run. This might be done if it is suspected that the posterior distribution of the mutation locations had not reached stationarity by the start of the good iterations. If an old tree is to be used, “.mat” must be renamed “inmat.txt” and placed in the same directory as the executable. In addition, the “use oldfile?” input parameter must be set to 1. This matrix is sorted by the first column. If multiple chains are run simultaneously, only the matrix from the first run is saved. The columns in the output are as follows:

Column 1: the node ID number.

Column 2: the node’s left child. If there is no left child (i.e. the node is a leaf [= extant chromosome]), this is coded as -1.0.

Column 3: the node’s right child. Coded as -1.0 as above if no right child.

Column 4: the node’s parent. Coded as -1.0 if there is no parent (i.e. the node is the root node).

Column 5: the node’s coalescence time (in generations). Leaves have a coalescence time of 0.0.

Column 6: the length of time (in generations) between the coalescence time of the node and that of its parent. If the node is the root, this is the time between the root’s coalescence time and the age of the mutation.

Column 7: the transition probability (log-likelihood) associated with that node’s current state, given the state of its parent (or that of the ancestral haplotype in the case of the root node).

Column 7a: In the recessive model, this column pairs chromosomes by individual.

Column 8 to column 8 + (number of markers) -1: the alleles at each marker locus. In the dominant model, there are an extra (number of markers) columns, indicating that individual’s alleles on the haplotype not currently in the disease tree.

The final row in the file records the haplotype of the root haplotype.

If there were multiple runs of the chain, only the tree from chain #1 run is saved.

*.hpf - This file contains haplotype frequencies for the non burn-in iterations of the first chain (results from additional chains, if present, are not stored). The haplotypes for these results are only recorded every “file update” generations (see Input file above).

.